

Universal Matrix Sparsifiers and Fast Deterministic Algorithms for Linear Algebra

Rajarshi Bhattacharjee* Gregory Dexter† Cameron Musco* Archan Ray*
 Sushant Sachdeva ‡ David P Woodruff§

Abstract

Let $\mathbf{S} \in \mathbb{R}^{n \times n}$ be any matrix satisfying $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$, where $\mathbf{1}$ is the all ones matrix and $\|\cdot\|_2$ is the spectral norm. It is well-known that there exists $\mathbf{S}$ with just $O(n/\epsilon^2)$ non-zero entries achieving this guarantee: we can let $\mathbf{S}$ be the scaled adjacency matrix of a Ramanujan expander graph. We show that, beyond giving a sparse approximation to the all ones matrix, $\mathbf{S}$ yields a *universal sparsifier* for any positive semidefinite (PSD) matrix. In particular, for any PSD $\mathbf{A} \in \mathbb{R}^{n \times n}$ which is normalized so that its entries are bounded in magnitude by 1, we show that $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon n$, where $\circ$ denotes the entrywise (Hadamard) product. Our techniques also yield universal sparsifiers for non-PSD matrices. In this case, we show that if $\mathbf{S}$ satisfies $\|\mathbf{1} - \mathbf{S}\|_2 \leq \frac{\epsilon^2 n}{c \log^2(1/\epsilon)}$ for some sufficiently large constant c , then $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$, where $\|\mathbf{A}\|_1$ is the nuclear norm. Again letting $\mathbf{S}$ be a scaled Ramanujan graph adjacency matrix, this yields a sparsifier with $\tilde{O}(n/\epsilon^4)$ entries. We prove that the above universal sparsification bounds for both PSD and non-PSD matrices are tight up to logarithmic factors.

Since $\mathbf{A} \circ \mathbf{S}$ can be constructed *deterministically* without reading all of $\mathbf{A}$, our result for PSD matrices derandomizes and improves upon established results for randomized matrix sparsification, which require sampling a random subset of $O(\frac{n \log n}{\epsilon^2})$ entries and only give an approximation to any fixed $\mathbf{A}$ with high probability. We further show that any randomized algorithm must read at least $\Omega(n/\epsilon^2)$ entries to spectrally approximate general $\mathbf{A}$ to error ϵn , thus proving that these existing randomized algorithms are optimal up to logarithmic factors. We leverage our deterministic sparsification results to give the first deterministic algorithms for several problems, including singular value and singular vector approximation and positive semidefiniteness testing, that run in faster than matrix multiplication time. This partially addresses a significant gap between randomized and deterministic algorithms for fast linear algebraic computation.

Finally, if $\mathbf{A} \in \{-1, 0, 1\}^{n \times n}$ is PSD, we show that a spectral approximation $\tilde{\mathbf{A}}$ with $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon n$ can be obtained by deterministically reading $\tilde{O}(n/\epsilon)$ entries of $\mathbf{A}$. This improves the $1/\epsilon$ dependence on our result for general PSD matrices by a quadratic factor and is information-theoretically optimal up to a logarithmic factor.

*University of Massachusetts Amherst, {rbhattacharj, cmusco, ray}@cs.umass.edu

†Purdue University, {gdexter}@purdue.edu

‡University of Toronto, {sachdeva}@cs.toronto.edu

§Carnegie Mellon University, {dwoodruff}@cs.cmu.edu

1 Introduction

A common task in processing large matrices is *element-wise sparsification*. Given an $n \times n$ matrix $\mathbf{A}$, the goal is to choose a small subset S of coordinates in $[n] \times [n]$, where $[n] = \{1, 2, \dots, n\}$, such that $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2$ is small, where $\circ$ denotes the Hadamard (entrywise) product and $\mathbf{S}$ is a sampling matrix which equals $\frac{n^2}{|S|}$ on the entries in S , but is 0 otherwise. As in previous work, we consider operator norm error, where for a matrix $\mathbf{B}$, $\|\mathbf{B}\|_2 \stackrel{\text{def}}{=} \sup_{\mathbf{x} \in \mathbb{R}^n} \frac{\|\mathbf{B}\mathbf{x}\|_2}{\|\mathbf{x}\|_2}$. Elementwise sparsification has been widely studied [AM07, DZ11, AKL13, BKKS21] and has been used as a primitive in several applications, including low-rank approximation [AM07, Kun15], approximate eigenvector computation [AHK06, AM07], semi-definite programming [AHK05, d'A11], and matrix completion [CR12a, CT10]. Without loss of generality, one can scale the entries of $\mathbf{A}$ so that the maximum entry is bounded by 1 in absolute value, and we refer to such matrices as having *bounded entries*. With this normalization, it will be convenient to consider the task of finding a small subset S with corresponding sampling matrix $\mathbf{S}$ such that, for a given error parameter $\epsilon \in (0, 1)$,

$$\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon \cdot n. \quad (1)$$

One can achieve the error guarantee in (1) for any bounded entry matrix $\mathbf{A}$ with high probability by uniformly sampling a set of $O\left(\frac{n \log n}{\epsilon^2}\right)$ entries of $\mathbf{A}$. See, e.g., Theorem 1 of [DZ11].¹ However, there are no known lower bounds for this problem, even if we consider the harder task of *universal sparsification*, which requires finding a fixed subset S such that (1) *holds simultaneously for every bounded entry matrix $\mathbf{A}$* . The existence of such a fixed subset S corresponds to the existence of a *deterministic sublinear query algorithm* that constructs a spectral approximation to any $\mathbf{A}$ by forming $\mathbf{A} \circ \mathbf{S}$ (which requires reading just $|S|$ entries of $\mathbf{A}$). As we will see, such algorithms have applications to fast deterministic algorithms for linear algebraic computation. We ask:

What is the size of the smallest set S that achieves (1) simultaneously for every bounded entry matrix $\mathbf{A}$?

Previously, no bound on the size of S better than the trivial $O(n^2)$ was known.

1.1 Our Results

Our first result answers the above question for the class of *symmetric positive semidefinite* (PSD) matrices, i.e., those matrices $\mathbf{A}$ for which $\mathbf{x}^T \mathbf{A} \mathbf{x} \geq 0$ for all vectors $\mathbf{x} \in \mathbb{R}^n$, or equivalently, whose eigenvalues are all non-negative. PSD matrices arise e.g., as covariance matrices, kernel similarity matrices, and graph Laplacians, and significant work has considered efficient algorithms for approximating their properties [WLZ16, CW17, XG10, Cha11, ACK⁺16, MMMW21, SKO21]. We summarize our results below and give a comparison to previous work in Table 1.1.

We show that there exists a set S with $|S| = O(n/\epsilon^2)$ that achieves (1) simultaneously for all bounded entry PSD matrices. This improves the best known randomized bound of $O\left(\frac{n \log n}{\epsilon^2}\right)$ for algorithms which only succeed on a *fixed matrix* with high probability.

Theorem 1 (Universal Sparsifiers for PSD Matrices). *There exists a subset S of $s = O(n/\epsilon^2)$ entries of $[n] \times [n]$ such that, letting $\mathbf{S} \in \mathbb{R}^{n \times n}$ have $\mathbf{S}_{ij} = \frac{n^2}{s}$ for $(i, j) \in S$ and $\mathbf{S}_{ij} = 0$ otherwise, simultaneously for all PSD matrices $\mathbf{A} \in \mathbb{R}^{n \times n}$ with bounded entries, $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon n$.*

¹To apply Thm. 1 of [DZ11] we first rescale $\mathbf{A}$ so that all entries are $\leq 1/n$ in magnitude, and note that $\|\mathbf{A}\|_F^2 \leq 1$ in this case. Thus, they sample $s = O\left(\frac{n \log n}{\epsilon^2}\right)$ entries. Rescaling their error guarantee analogously gives (1).

Matrix Type	Approx. Type	Randomized Error	Randomized Sample Complexity	Deterministic Error	Deterministic Sample Complexity
$\ \mathbf{A}\ _\infty \leq 1$ $\mathbf{A}$ is PSD	sparse	ϵn	$\Theta(n/\epsilon^2)$ (Thms 1, 3)	ϵn	$\Theta(n/\epsilon^2)$ (Thms 1, 3)
$\ \mathbf{A}\ _\infty \leq 1$	sparse	ϵn	$O\left(\frac{n \log n}{\epsilon^2}\right)$ [AM07] $\Omega(n/\epsilon^2)$ (Thm 3)	$\epsilon \max(n, \ \mathbf{A}\ _1)$	$O\left(\frac{n \log^4(1/\epsilon)}{\epsilon^4}\right)$ (Thm 4) $\Omega(n/\epsilon^4)$ (Thm 6)
$\ \mathbf{A}\ _\infty \leq 1$ $\mathbf{A}$ is PSD	any	ϵn	$O\left(\frac{n \log(1/\epsilon)}{\epsilon}\right)$ [MM17]	ϵn	$O\left(\frac{n \log n}{\epsilon}\right)$ (Thm 9) for $\mathbf{A} \in \{-1, 0, 1\}^{n \times n}$ $\Omega(n/\epsilon)$ (Thm 10)
$\ \mathbf{A}\ _\infty \leq 1$	any	ϵn	$\Omega(n/\epsilon^2)$ (Thm 11)	$\epsilon \max(n, \ \mathbf{A}\ _1)$	$\Omega(n/\epsilon^2)$ (Thm 7)

Table 1: Summary of results. Algorithms in the first two rows output a spectral approximation that is sparse, as in (1). Our deterministic algorithms in this case follow directly from our universal sparsification results. Algorithms in the last two rows can output an approximation of any form.

Theorem 1 can be viewed as a significant strengthening of a classic spectral graph expander guarantee [Mar73, Alo86]; indeed, letting $\mathbf{A} = \mathbf{1}$ be the all ones matrix, we have that if $\mathbf{A} \circ \mathbf{S}$ satisfies (1), then it matches the near-optimal spectral expansion of Ramanujan graphs, up to a constant factor.² In fact, we prove Theorem 1 by proving a more general claim: that any matrix $\mathbf{S}$ that sparsifies the all ones matrix also sparsifies every bounded entry PSD matrix. In particular, Theorem 1 follows as a direct corollary of:

Theorem 2 (Spectral Expanders are Universal Sparsifiers for PSD Matrices). *Let $\mathbf{1} \in \mathbb{R}^{n \times n}$ be the all ones matrix. Let $\mathbf{S} \in \mathbb{R}^{n \times n}$ be any matrix such that $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$. Then for any PSD matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ with bounded entries, $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon n$.*

To prove Theorem 1 given Theorem 2, we let S be the edge set of a Ramanujan spectral expander graph with degree $d = O(1/\epsilon^2)$ and adjacency matrix $\mathbf{G}$. We have $s = nd = O(n/\epsilon^2)$ and $\mathbf{S} = \frac{n^2}{s} \mathbf{G} = \frac{n}{d} \mathbf{G}$. Thus, the top eigenvector of $\mathbf{S}$ is the all ones vector with eigenvalue $\lambda_1(\mathbf{S}) = \frac{n}{d} \lambda_1(\mathbf{G}) = n$. All other eigenvalues are bounded by $|\lambda_i(\mathbf{S})| = \frac{n}{d} |\lambda_i(\mathbf{G})| = O(\frac{n}{d} \cdot \sqrt{d}) = O(\epsilon n)$. Combined, after adjusting ϵ by a constant factor, this shows that $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$, as required.

Sparsity Lower Bound. We show that Theorem 1 is tight. Even if we only seek to approximate the all ones matrix (rather than all bounded PSD matrices), $\mathbf{S}$ requires $\Omega(n/\epsilon^2)$ non-zero entries.

Theorem 3 (Sparsity Lower Bound – PSD Matrices). *Let $\mathbf{1} \in \mathbb{R}^{n \times n}$ be the all-ones matrix. Then, for any $\epsilon \in (0, 1/2)$ with $\epsilon \geq c/\sqrt{n}$ for large enough constant c , any $\mathbf{S} \in \mathbb{R}^{n \times n}$ with $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$ must have $\Omega(n/\epsilon^2)$ nonzero entries.*

Theorem 3 resolves an open question of [BKKS21], which asked whether a spectral approximation must have $\Omega\left(\frac{\text{ns}(\mathbf{A}) \text{sr}(\mathbf{A})}{\epsilon^2}\right)$ non-zeros, where $\text{sr}(\mathbf{A}) = \frac{\|\mathbf{A}\|_F^2}{\|\mathbf{A}\|_2^2}$ is the stable rank and $\text{ns}(\mathbf{A}) = \max_i \frac{\|\mathbf{a}_i\|_1^2}{\|\mathbf{a}_i\|_2^2}$, where $\mathbf{a}_i$ is the i^{th} row of $\mathbf{A}$, is the ‘numerical sparsity’. They give a lower bound when $\text{sr}(\mathbf{A}) = \Theta(n)$ but ask if this can be extended to $\text{sr}(\mathbf{A}) = o(n)$. For the all ones matrix, $\text{sr}(\mathbf{A}) = 1$ and $\text{ns}(\mathbf{A}) = n$, and so Theorem 3 resolves this question. Further, by applying the theorem to a block diagonal matrix with r disjoint $k \times k$ blocks of all ones, we resolve the question for integer $\text{sr}(\mathbf{A})$ and $\text{ns}(\mathbf{A})$.

²If we let $d = s/n = O(1/\epsilon^2)$ be the average number of entries per row of $\mathbf{S}$, and let $\mathbf{G} \in \{0, 1\}^{n \times n}$ be the binary adjacency matrix of the graph with edges in S (i.e., $\mathbf{G} = \frac{s}{n^2} \cdot \mathbf{S}$), then by (1) applied with $\mathbf{A} = \mathbf{1}$, $\lambda_1(\mathbf{G}) = \frac{s}{n^2} \lambda_1(\mathbf{S}) = \frac{s}{n^2} \lambda_1(\mathbf{1} \circ \mathbf{S}) = \frac{s}{n^2} (\lambda_1(\mathbf{1}) \pm \epsilon n) = \Theta(s/n) = \Theta(d)$, while for $i > 1$, $|\lambda_i(\mathbf{G})| \leq \frac{s}{n^2} (|\lambda_i(\mathbf{1})| + \epsilon n) \leq \epsilon d = \Theta(\sqrt{d})$.

Non-PSD Matrices. The techniques used to prove Theorem 1 also give nearly tight universal sparsification bounds for general bounded entry (not necessarily PSD) matrices. In this case, we show that a subset S of $O\left(\frac{n \log^4(1/\epsilon)}{\epsilon^4}\right)$ entries suffices to achieve spectral norm error depending on the nuclear norm $\|\mathbf{A}\|_1$, which is the sum of singular values of $\mathbf{A}$.

Theorem 4 (Universal Sparsifiers for Non-PSD Matrices). *There exists a subset S of $s = O\left(\frac{n \log^4(1/\epsilon)}{\epsilon^4}\right)$ entries of $[n] \times [n]$ such that, letting $\mathbf{S} \in \mathbb{R}^{n \times n}$ have $\mathbf{S}_{ij} = \frac{n^2}{s}$ for $(i, j) \in S$ and $\mathbf{S}_{ij} = 0$ otherwise, simultaneously for all symmetric matrices $\mathbf{A} \in \mathbb{R}^{n \times n}$ with bounded entries,*

$$\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1). \quad (2)$$

Note that for a bounded entry PSD matrix $\mathbf{A}$, we have $\|\mathbf{A}\|_1 = \text{tr}(\mathbf{A}) = O(n)$ and so for PSD matrices, (1) and (2) are equivalent.

Remark 1. *Although Theorem 4 is stated for symmetric $\mathbf{A}$, one can first symmetrize $\mathbf{A}$ by considering $\mathbf{B} = [\mathbf{0}, \mathbf{A}; \mathbf{A}^T, \mathbf{0}]$. Then for any $\mathbf{z} = (\mathbf{x}, \mathbf{y}) \in \mathbb{R}^{2n}$, we have $\mathbf{z}^T \mathbf{B} \mathbf{z} = 2\mathbf{x}^T \mathbf{A} \mathbf{y}$, and $\|\mathbf{B}\|_1 = 2\|\mathbf{A}\|_1$, so applying the above theorem to $\mathbf{B}$ gives us a spectral approximation to $\mathbf{A}$.*

As with Theorem 1, Theorem 4 follows from a general claim which shows that sparsifying the all ones matrix to small enough error suffices to sparsify any bounded entry symmetric matrix.

Theorem 5 (Spectral Expanders are Universal Sparsifiers for Non-PSD Matrices). *Let $\mathbf{1} \in \mathbb{R}^{n \times n}$ be the all ones matrix. Let $\mathbf{S} \in \mathbb{R}^{n \times n}$ be any matrix such that $\|\mathbf{1} - \mathbf{S}\|_2 \leq \frac{\epsilon^2 n}{c \log^2(1/\epsilon)}$ for some large enough constant c . Then for any symmetric $\mathbf{A} \in \mathbb{R}^{n \times n}$ with bounded entries, $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$.*

Theorem 4 follows by applying Theorem 5 where $\mathbf{S}$ is taken to be the scaled adjacency matrix of a Ramanujan expander graph with degree $d = O(\log^4(1/\epsilon)/\epsilon^4)$.

Lower Bounds for Non-PSD Matrices. We can prove that Theorem 4 is tight up to logarithmic factors. Our lower bound holds even for the easier problem of top singular value (spectral norm) approximation and against a more general class of algorithms, which non-adaptively and deterministically query entries of the input matrix. The idea is simple: since the entries read by the deterministic algorithm are fixed, we can construct two very different input instances on which the algorithm behaves identically: one which is the all ones matrix, and the other which is one only on the entries read by the algorithm and zero everywhere else. We show that if the number of entries s that are read is too small, the top singular value of the second instance is significantly smaller than the first (as compared to their Schatten-1 norms), which violates the desired error bound of $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. To obtain a tight bound, we apply the above construction on just a subset of the rows of the matrix on which the algorithm does not read too many entries. Here, we critically use non-adaptivity, as this subset of rows can be fixed, independent of the input instance.

Theorem 6 (Non-Adaptive Query Lower Bound for Deterministic Spectral Approximation of Non-PSD Matrices). *For any $\epsilon \in (1/n^{1/4}, 1/4)$, any deterministic algorithm that queries entries of a bounded entry matrix $\mathbf{A}$ non-adaptively and outputs $\tilde{\sigma}_1$ satisfying $|\sigma_1(\mathbf{A}) - \tilde{\sigma}_1| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ must read at least $\Omega\left(\frac{n}{\epsilon^4}\right)$ entries. Here $\sigma_1(\mathbf{A}) = \|\mathbf{A}\|_2$ is the largest singular value of $\mathbf{A}$.*

Observe that Theorem 4 implies the existence of a non-adaptive deterministic algorithm for the above problem with query complexity $O\left(\frac{n \log^4(1/\epsilon)}{\epsilon^4}\right)$, since $\mathbf{A} \circ \mathbf{S}$ can be computed with non-adaptive deterministic queries and since $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ implies via Weyl's inequality that

$|\sigma_1(\mathbf{A}) - \sigma_1(\mathbf{A} \circ \mathbf{S})| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Also recall that while Theorem 4 is stated for bounded entry symmetric matrices, it applies to general (possibly asymmetric) matrices by Remark 1.

Note that Theorem 6 establishes a separation between universal sparsifiers and randomized sparsification since randomly sampling $O\left(\frac{n \log n}{\epsilon^2}\right)$ entries of any $\mathbf{A} \in \mathbb{R}^{n \times n}$ achieves error ϵn by [DZ11]. In fact, for the problem of just approximating $\sigma_1(\mathbf{A})$ to error $\pm \epsilon n$, randomized algorithms using just $\text{poly}(\log n, 1/\epsilon)$ samples are known [BCJ20, BDD⁺23]. Theorem 6 shows that universal sparsifiers for general bounded entry matrices (and in fact all non-adaptive deterministic algorithms for spectral approximation) require a worse $1/\epsilon^4$ dependence and error bound of $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. We can extend Theorem 6 to apply to general deterministic algorithms that query $\mathbf{A}$ possibly *adaptively*, however the lower bound weakens to $\Omega(n/\epsilon^2)$. Understanding if this gap between adaptive and non-adaptive query deterministic algorithms is real is an interesting open question.

Theorem 7 (Adaptive Query Lower Bound for Deterministic Spectral Approximation of Non-PSD Matrices). *For any $\epsilon \in (1/\sqrt{n}, 1/4)$, any deterministic algorithm that queries entries of a bounded entry matrix $\mathbf{A}$ (possibly adaptively) and outputs $\tilde{\sigma}_1$ satisfying $|\sigma_1(\mathbf{A}) - \tilde{\sigma}_1| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ must read at least $\Omega\left(\frac{n}{\epsilon^2}\right)$ entries. Here $\sigma_1(\mathbf{A}) = \|\mathbf{A}\|_2$ is the largest singular value of $\mathbf{A}$.*

1.1.1 Applications to Fast Deterministic Algorithms for Linear Algebra

Given sampling matrix $\mathbf{S}$ such that $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$, one can use $\mathbf{A} \circ \mathbf{S}$ to approximate various linear algebraic properties of $\mathbf{A}$. For example, by Weyl's inequality [Wey12], the eigenvalues of $\mathbf{A} \circ \mathbf{S}$ approximate those of $\mathbf{A}$ up to additive error $\pm \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Thus, our universal sparsification results (Theorems 1, 4) immediately give the first known deterministic algorithms for approximating the eigenspectrum of a bounded entry matrix up to small additive error with *sublinear entrywise query complexity*. Previously, only randomized sublinear query algorithms were known [WS00, MM17, MW17, TYUC17, BCJ20, BDD⁺23].

Further, our results yield the first deterministic algorithms for several problems that run in $o(n^\omega)$ time, where $\omega \approx 2.373$ is the matrix multiplication exponent [AW21]. Consider e.g., approximating the top singular value $\sigma_1(\mathbf{A})$ (the spectral norm) of $\mathbf{A}$. A $(1 + \epsilon)$ -relative error approximation can be computed in $\tilde{O}(n^2/\sqrt{\epsilon})$ time with high probability using $O(\log(n/\epsilon)/\sqrt{\epsilon})$ iterations of the Lanczos method with a random initialization [KW92, MM15]. This can be further accelerated e.g., via randomized entrywise sparsification [DZ11], allowing an additive $\pm \epsilon n$ approximation to $\sigma_1(\mathbf{A})$ to be computed in $\tilde{O}(n/\text{poly}(\epsilon))$ time. However, prior to our work, no fast deterministic algorithms were known, even with coarse approximation guarantees. The fastest approach was just to perform a full SVD of $\mathbf{A}$, requiring $\Theta(n^\omega)$ time. In fact, this gap between randomized and deterministic methods exists for many linear algebraic problems, and resolving it is a central open question.

By combining our universal sparsification results with a derandomized power method that uses a full subspace of vectors for initialization, and iteratively approximates the singular values of $\mathbf{A} \circ \mathbf{S}$ via ‘deflation’ [Saa11], we give to the best of our knowledge, the first $o(n^\omega)$ time deterministic algorithm for approximating *all singular values* of a bounded entry matrix $\mathbf{A}$ to small additive error.

Theorem 8 (Deterministic Singular Value Approximation). *Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a bounded entry symmetric matrix with singular values $\sigma_1(\mathbf{A}) \geq \sigma_2(\mathbf{A}) \geq \dots \geq \sigma_n(\mathbf{A})$. Then there exists a deterministic algorithm that, given $\epsilon \in (0, 1)$, reads $\tilde{O}\left(\frac{n}{\epsilon^4}\right)$ entries of $\mathbf{A}$, runs in $\tilde{O}\left(\frac{n^2}{\epsilon^8}\right)$ time, and returns singular value approximations $\tilde{\sigma}_1(\mathbf{A}), \dots, \tilde{\sigma}_n(\mathbf{A})$ satisfying for all i ,*

$$|\sigma_i(\mathbf{A}) - \tilde{\sigma}_i(\mathbf{A})| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1).$$

Further, for all $i \leq 1/\epsilon$, the algorithm returns a unit vector $\mathbf{z}_i$ such that $|\|\mathbf{A}\mathbf{z}_i\|_2 - \sigma_i(\mathbf{A})| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ and the returned vectors are all mutually orthogonal.

The runtime of Theorem 8 is stated assuming we have already constructed a sampling matrix $\mathbf{S}$ that samples $\tilde{O}(\frac{n}{\epsilon^4})$ entries of $\mathbf{A}$ and satisfies Theorem 4. It is well known that, if we let $\mathbf{S}$ be the adjacency matrix of a Ramanujan expander graph, it can indeed be constructed deterministically in $\tilde{O}(n/\text{poly}(\epsilon))$ time [Alo21]. This is lower order as compared to the runtime of $\tilde{O}(\frac{n^2}{\epsilon^8})$ unless $\epsilon < 1/n^c$ for some large enough constant c . Further, for a fixed input size n (and in fact, for a range of input sizes), $\mathbf{S}$ needs to be constructed only once – see Section 2 for further discussion.

Recall that if we further assume $\mathbf{A}$ to be PSD, then the error in Theorem 8 is bounded by $\epsilon \cdot \max(n, \|\mathbf{A}\|_1) \leq \epsilon n$. The sample complexity and runtime also improve by $\text{poly}(1/\epsilon)$ factors due to the tighter universal sparsifier bound for PSD matrices given in Theorem 1. Also observe that while Theorem 8 gives additive error approximations to all of $\mathbf{A}$'s singular values, these approximations are only meaningful for singular values larger than $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$, of which there are at most $1/\epsilon$. Similar additive error guarantees have been given using randomized algorithms [BDD⁺23, WS23]. Related bounds have also been studied in work on randomized methods for spectral density estimation [WWAF06, LSY16, BKM22].

We further leverage Theorem 8 to give the first $o(n^\omega)$ time deterministic algorithm for testing if a bounded entry matrix is either PSD or has at least one large negative eigenvalue $\leq -\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Recent work has focused on optimal randomized methods for this problem [BCJ20, NSW22]. We also show that, under the assumption that $\sigma_1(\mathbf{A}) \geq \alpha \cdot \max(n, \|\mathbf{A}\|_1)$ for some $\alpha \in (0, 1)$, one can deterministically compute $\tilde{\sigma}_1(\mathbf{A})$ with $|\sigma_1(\mathbf{A}) - \tilde{\sigma}_1(\mathbf{A})| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ in $\tilde{O}\left(\frac{n^2 \log(1/\epsilon)}{\text{poly}(\alpha)}\right)$ time. That is, one can compute a highly accurate approximation to the top singular value in roughly linear time in the input matrix size. Again, this is the first $o(n^\omega)$ time deterministic algorithm for this problem, and matches the runtime of the best known randomized methods for high accuracy top singular value computation, up to a $\text{poly}(\log(n, 1/\alpha))$ factor.

1.1.2 Beyond Sparse Approximations

It is natural ask if it is possible to achieve better than $O(n/\epsilon^2)$ sample complexity for spectral approximation by using an algorithm that does not output a sparse approximation to $\mathbf{A}$, but can output more general data structures, allowing it to avoid the sparsity lower bound of Theorem 3. Theorems 6 and 7 already rule this out for both non-adaptive and adaptive deterministic algorithms for non-PSD matrices. However, it is known that this *is possible* with randomized algorithms for PSD matrices. For example, following [AKM⁺17], one can apply Theorem 3 of [MM17] with error parameter $\lambda = \epsilon n$. Observing that all ridge leverage score sampling probabilities are bounded by $1/\lambda$ (e.g., via Lemma 6 of the same paper and the bounded entry assumption), one can show that a Nyström approximation $\tilde{\mathbf{A}}$ based on $\tilde{O}(1/\epsilon)$ uniformly sampled columns satisfies $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon n$ with high probability. Further $\tilde{\mathbf{A}}$ can be constructed by reading just $\tilde{O}(1/\epsilon)$ columns and thus $\tilde{O}(n/\epsilon)$ entries of $\mathbf{A}$, giving a linear rather than quadratic dependence on $1/\epsilon$ as compared to Theorem 1. Unfortunately, derandomizing such a column-sampling-based approach seems difficult – any deterministic algorithm must read entries in $\Omega(n)$ columns of $\mathbf{A}$, as otherwise it will fail when $\mathbf{A}$ is entirely supported on the unread columns.

Nevertheless, in the special case where $\mathbf{A}$ is PSD and has entries in $\{-1, 0, 1\}^{n \times n}$, we show that a spectral approximation can be obtained by deterministically reading just $\tilde{O}(n/\epsilon)$ entries of $\mathbf{A}$.

Theorem 9 (Deterministic Spectral Approximation of Binary Magnitude PSD Matrices). *Let $\mathbf{A} \in \{-1, 0, 1\}^{n \times n}$ be PSD. Then for any $\epsilon \in (0, 1)$, there exists a deterministic algorithm which reads $O\left(\frac{n \log n}{\epsilon}\right)$ entries of $\mathbf{A}$ and returns PSD $\tilde{\mathbf{A}} \in \{-1, 0, 1\}^{n \times n}$ such that $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon n$.*

Using Turán's theorem, we show that Theorem 9 is information theoretically optimal up to a $\log n$ factor, even for the potentially much easier problem of eigenvalue approximation:

Theorem 10 (Deterministic Spectral Approximation of Binary PSD Matrices – Lower Bound). *Let $\mathbf{A} \in \{0, 1\}^{n \times n}$ be PSD. Then for any $\epsilon \in (0, 1)$, any possibly adaptive deterministic algorithm which approximates all eigenvalues of $\mathbf{A}$ up to ϵn additive error must read $\Omega\left(\frac{n}{\epsilon}\right)$ entries of $\mathbf{A}$.*

An interesting open question is if $\tilde{O}(n/\epsilon)$ sample complexity can be achieved for deterministic spectral approximation of *any bounded entry PSD matrix*, with a matrix $\tilde{\mathbf{A}}$ of any form, matching what is known for randomized algorithms.

Finally, we show that the PSD assumption is critical to achieve $o(n/\epsilon^2)$ query complexity, for both randomized and deterministic algorithms. Without this assumption, $\Omega(n/\epsilon^2)$ queries are required, even to achieve (1) with constant probability for a single input $\mathbf{A}$. For bounded entry matrices, the randomized element-wise sparsification algorithms of [AM07, DZ11] read just $\tilde{O}(n/\epsilon^2)$ entries of $\mathbf{A}$, and so our lower bounds are the first to show near-optimality of the number of entries read of these algorithms, which may be of independent interest.

Theorem 11 (Lower Bound for Randomized Spectral Approximation). *Any randomized algorithm that (possibly adaptively) reads entries of a binary matrix $\mathbf{A} \in \{0, 1\}^{n \times n}$ to construct a data structure $f : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ that, for $\epsilon \in (0, 1)$ satisfies with probability at least 99/100,*

$$|f(\mathbf{x}, \mathbf{y}) - \mathbf{x}^T \mathbf{A} \mathbf{y}| \leq \epsilon n, \text{ for all unit } \mathbf{x}, \mathbf{y} \in \mathbb{R}^n,$$

must read $\Omega\left(\frac{n}{\epsilon^2}\right)$ entries of $\mathbf{A}$ in the worst case, provided $\epsilon = \Omega\left(\frac{\log n}{\sqrt{n}}\right)$.

1.1.3 Relation to Spectral Graph Sparsification

We remark that while our work focuses on general bounded entry symmetric (PSD) matrices, when $\mathbf{A}$ is a PSD graph Laplacian matrix, it is possible to construct a sparsifier $\tilde{\mathbf{A}}$ with just $\tilde{O}(n/\epsilon^2)$ entries, that achieves a *relative error* spectral approximation guarantee that can be much stronger than the additive error guarantee of (1) [BSST13]. In particular, one can achieve $(1 - \epsilon)\mathbf{x}^T \mathbf{A} \mathbf{x} \leq \mathbf{x}^T \tilde{\mathbf{A}} \mathbf{x} \leq (1 + \epsilon)\mathbf{x}^T \mathbf{A} \mathbf{x}$ for all $\mathbf{x} \in \mathbb{R}^n$. To achieve such a guarantee however, it is not hard to see that the set of entries sampled in $\tilde{\mathbf{A}}$ must depend on $\mathbf{A}$, and cannot be universal. Fast randomized algorithms for constructing $\tilde{\mathbf{A}}$ have been studied extensively [ST04, SS08, BSST13]. Recent work has also made progress towards developing fast deterministic algorithms [CGL⁺20].

1.2 Road Map

We introduce notation in Section 2. In Section 3, we prove our main universal sparsification results for PSD and non-PSD matrices (Theorems 1 and 4). In Section 4 we prove nearly matching lower bounds (Theorems 3 and 6). We also prove Theorems 7 and 11 which give lower bounds for general deterministic (possibly adaptive) spectral approximation algorithms, and general randomized algorithms, respectively. In Section 5, we show how to give tighter deterministic spectral approximation results for PSD matrices with entries in $\{-1, 0, 1\}$. Finally, in Section 6, we discuss applications of our results to fast deterministic algorithms for singular value and vector approximation.

2 Notation and Preliminaries

We start by defining notation used throughout. For any integer n , let $[n]$ denote the set $\{1, 2, \dots, n\}$.

Matrices and Vectors. Matrices are represented with bold uppercase literals, e.g., $\mathbf{A}$. Vectors are represented with bold lowercase literals, e.g., $\mathbf{x}$. $\mathbf{1}$ and $\mathbf{0}$ denote all ones (resp. all zeros) matrices

or vectors. The identity matrix is denoted by $\mathbf{I}$. The size of these matrices vary based on their applications. For a vector $\mathbf{x}$, $\mathbf{x}(j)$ denotes its j^{th} entry. For a matrix $\mathbf{A}$, $\mathbf{A}_{ij}$ denotes the entry in the i^{th} row and j^{th} column. For a vector $\mathbf{x}$ (or matrix $\mathbf{A}$), $\mathbf{x}^T$ (resp. $\mathbf{A}^T$) denotes its transpose. For two matrices $\mathbf{A}, \mathbf{B}$ of the same size, $\mathbf{A} \circ \mathbf{B}$ denotes the entrywise (Hadamard) product.

Matrix Norms and Properties. For a vector $\mathbf{x}$, $\|\mathbf{x}\|_2$ denotes its Euclidean norm and $\|\mathbf{x}\|_1 = \sum_{i=1}^n |\mathbf{x}(i)|$ denotes its ℓ_1 norm. We denote the eigenvalues of a symmetric matrix $\mathbf{A}$ as $\lambda_1(\mathbf{A}) \geq \lambda_2(\mathbf{A}) \geq \dots \geq \lambda_n(\mathbf{A})$ in decreasing order. A symmetric matrix is positive semidefinite (PSD) if $\lambda_i \geq 0$ for all $i \in [n]$. The singular values of a matrix $\mathbf{A}$ are denoted as $\sigma_1(\mathbf{A}) \geq \sigma_2(\mathbf{A}) \geq \dots \geq \sigma_n(\mathbf{A}) \geq 0$ in decreasing order. We let $\|\mathbf{A}\|_2 = \max_x \frac{\|\mathbf{A}\mathbf{x}\|_2}{\|\mathbf{x}\|_2} = \sigma_1(\mathbf{A})$ denote the spectral norm, $\|\mathbf{A}\|_\infty$ denote the largest magnitude of an entry, $\|\mathbf{A}\|_F = (\sum_{i,j} \mathbf{A}_{ij}^2)^{1/2}$ denote the Frobenius norm, and $\|\mathbf{A}\|_1 = \sum_{i=1}^n \sigma_i(\mathbf{A})$ denote the Schatten-1 norm (also called the trace norm or nuclear norm).

Expander Graphs. Our universal sparsifier constructions are based on Ramanujan expander graphs [LPS88, Mar73, Alo86], defined below.

Definition 2.1 (Ramanujan Expander Graphs [LPS88]). *Let G be a connected d -regular, unweighted and undirected graph on n vertices. Let $\mathbf{G} \in \{0, 1\}^{n \times n}$ be the adjacency matrix corresponding to G and λ_i be its i^{th} eigenvalue, such that $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$. Then G is called a Ramanujan graph if:*

$$|\lambda_i| \leq 2\sqrt{d-1}, \text{ for all } i > 1.$$

Equivalently, letting $\mathbf{1}$ be the $n \times n$ all ones matrix, $\|\mathbf{1} - \frac{n}{d}\mathbf{G}\|_2 \leq \frac{n}{d} \cdot 2\sqrt{d-1}$.

Efficient Construction of Ramanujan graphs. Significant work has studied efficient constructions for Ramanujan Graphs [Mar73, LPS88, Mor94, Coh16], and nearly linear time constructions, called *strongly explicit* constructions [Alo21], have been proposed. E.g., by Proposition 1.1 of [Alo21] we can reconstruct for any n and d , a graph on $n(1+o(1))$ vertices with second eigenvalue at most $(2+o(1))\sqrt{d}$ in $\tilde{O}(nd + \text{poly}(d))$ time. Additionally, in our applications, the expander just needs to be constructed once and can then be used for any input of size n . In fact, a single expander can be used for a range of input sizes, by the argument below.

Though the size of the expander above is $(1+o(1))n$ instead of exactly n , and its expansion does not exactly hit the tight Ramanujan bound of $2\sqrt{d-1}$, we can let our input matrix $\mathbf{A}$ be the top $n \times n$ principal submatrix of a slightly larger $n(1+o(1)) \times n(1+o(1))$ matrix $\mathbf{A}'$ that is zero everywhere else. Then, the top $n \times n$ principal submatrix of a spectral approximation to $\mathbf{A}'$ is a spectral approximation to $\mathbf{A}$. Thus, obtaining a spectral approximation to $\mathbf{A}'$ via a Ramanujan sparsifier on $n(1+o(1))$ vertices suffices. Similarly, in our applications, we can adjust $d = 1/\text{poly}(\epsilon)$ by at most a constant factor to account for the $(2+o(1))\sqrt{d}$ bound on the second eigenvalue, rather than the tight Ramanujan bound of $2\sqrt{d-1}$.

3 Universal Sparsifier Upper Bounds

We now prove our main results on universal sparsifiers for PSD matrices (Theorem 1) and general symmetric matrices (Theorem 4). Both theorems follow from general reductions which show that any sampling matrix $\mathbf{S}$ that sparsifies the all ones matrix to sufficient accuracy (i.e., is a sufficiently good spectral expander) yields a universal sparsifier. We prove the reduction for the PSD case in Section 3.1, and then extend it to the non-PSD case in Section 3.2.

3.1 Universal Sparsifiers for PSD Matrices

In the PSD case, we prove the following:

Theorem 2 (Spectral Expanders are Universal Sparsifiers for PSD Matrices). *Let $\mathbf{1} \in \mathbb{R}^{n \times n}$ be the all ones matrix. Let $\mathbf{S} \in \mathbb{R}^{n \times n}$ be any matrix such that $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$. Then for any PSD matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ with bounded entries, $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon n$.*

Theorem 1 follows directly from Theorem 2 by letting $\mathbf{S}$ be the scaled adjacency matrix of a Ramanujan expander graph with degree $d = O(1/\epsilon^2)$ (Definition 2.1).

Proof of Theorem 2. To prove the theorem it suffices to show that for any $\mathbf{x} \in \mathbb{R}^n$ with $\|\mathbf{x}\|_2 = 1$, $|\mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{x}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| \leq \epsilon n$. Let $\mathbf{v}_1, \dots, \mathbf{v}_n$ and $\lambda_1, \dots, \lambda_n$ be the eigenvectors and eigenvalues of $\mathbf{A}$ so that $\mathbf{A} = \sum_{i=1}^n \lambda_i \mathbf{v}_i \mathbf{v}_i^T$. Then we can expand out this error as:

$$\begin{aligned} |\mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{x}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| &= \left| \sum_{i=1}^n \lambda_i \cdot \mathbf{x}^T \mathbf{v}_i \mathbf{v}_i^T \mathbf{x} - \sum_{i=1}^n \lambda_i \mathbf{x}^T (\mathbf{v}_i \mathbf{v}_i^T \circ \mathbf{S}) \mathbf{x} \right| \\ &= \left| \sum_{i=1}^n \lambda_i \mathbf{x}^T [\mathbf{v}_i \mathbf{v}_i^T \circ (\mathbf{1} - \mathbf{S})] \mathbf{x} \right|, \end{aligned}$$

where we use that for any matrix $\mathbf{M}$, $\mathbf{M} \circ \mathbf{1} = \mathbf{M}$. Now observe that if we let $\mathbf{D}_i \in \mathbb{R}^{n \times n}$ be a diagonal matrix with the entries of $\mathbf{v}_i$ on its diagonal, then we can write $\mathbf{v}_i \mathbf{v}_i^T \circ (\mathbf{1} - \mathbf{S}) = \mathbf{D}_i (\mathbf{1} - \mathbf{S}) \mathbf{D}_i$. Plugging this back in we have:

$$\begin{aligned} |\mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{x}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| &= \left| \sum_{i=1}^n \lambda_i \mathbf{x}^T \mathbf{D}_i (\mathbf{1} - \mathbf{S}) \mathbf{D}_i \mathbf{x} \right| \\ &\leq \sum_{i=1}^n \lambda_i \|\mathbf{1} - \mathbf{S}\|_2 \cdot \mathbf{x}^T \mathbf{D}_i^2 \mathbf{x} \\ &\leq \epsilon n \cdot \sum_{i=1}^n \lambda_i \mathbf{x}^T \mathbf{D}_i^2 \mathbf{x}. \end{aligned} \tag{3}$$

In the second line we use that $\mathbf{A}$ is PSD and thus λ_i is non-negative for all i . We also use that $|\mathbf{x}^T \mathbf{D}_i (\mathbf{1} - \mathbf{S}) \mathbf{D}_i \mathbf{x}| \leq \|\mathbf{1} - \mathbf{S}\|_2 \cdot \|\mathbf{D}_i \mathbf{x}\|_2^2 = \|\mathbf{1} - \mathbf{S}\|_2 \cdot \mathbf{x}^T \mathbf{D}_i^2 \mathbf{x}$. In the third line we bound $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$ using the assumption of the theorem statement.

Finally, writing $\mathbf{x}^T \mathbf{D}_i^2 \mathbf{x} = \sum_{j=1}^n \mathbf{x}(j)^2 \mathbf{v}_i(j)^2$ and plugging back into (3), we have:

$$\begin{aligned} |\mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{x}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| &\leq \epsilon n \cdot \sum_{i=1}^n \sum_{j=1}^n \lambda_i \mathbf{x}(j)^2 \mathbf{v}_i(j)^2 \\ &= \epsilon n \cdot \sum_{j=1}^n \mathbf{x}(j)^2 \sum_{i=1}^n \lambda_i \mathbf{v}_i(j)^2 \\ &= \epsilon n \cdot \sum_{j=1}^n \mathbf{x}(j)^2 \mathbf{A}_{jj} \\ &\leq \epsilon n, \end{aligned}$$

where in the last step we use that $\mathbf{A}_{jj} \leq 1$ by our bounded entry assumption, and that $\mathbf{x}$ is a unit vector. This completes the theorem. $\square$

Remark 2. Note that we can state a potentially stronger bound for Theorem 2 by observing that in the second to last step, $|\mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{x}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| \leq \epsilon n \cdot \sum_{j=1}^n \mathbf{x}(j)^2 \mathbf{A}_{jj} = \epsilon n \cdot \mathbf{x}^T \mathbf{D} \mathbf{x}$, where $\mathbf{D}$ is a diagonal matrix containing the diagonal elements of $\mathbf{A}$. That is, letting $\mathbf{M} \preceq \mathbf{N}$ denote that $\mathbf{N} - \mathbf{M}$ is PSD, we have the following spectral approximation bound: $-\epsilon n \cdot \mathbf{D} \preceq \mathbf{A} - (\mathbf{A} \circ \mathbf{S}) \preceq \epsilon n \cdot \mathbf{D}$.

3.2 Universal Sparsifiers for Non-PSD Matrices

We next extend the above approach to give a similar reduction for universal sparsification of general symmetric matrices.

Theorem 5 (Spectral Expanders are Universal Sparsifiers for Non-PSD Matrices). *Let $\mathbf{1} \in \mathbb{R}^{n \times n}$ be the all ones matrix. Let $\mathbf{S} \in \mathbb{R}^{n \times n}$ be any matrix such that $\|\mathbf{1} - \mathbf{S}\|_2 \leq \frac{\epsilon^2 n}{c \log^2(1/\epsilon)}$ for some large enough constant c . Then for any symmetric $\mathbf{A} \in \mathbb{R}^{n \times n}$ with bounded entries, $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$.*

Observe that as compared to the PSD case (Theorem 2), here we require that $\mathbf{S}$ gives a stronger approximation to the all ones matrix. Theorem 4 follows directly by letting $\mathbf{S}$ be the scaled adjacency matrix of a Ramanujan expander graph with degree $d = O\left(\frac{\log^4(1/\epsilon)}{\epsilon^4}\right)$ (Definition 2.1).

Proof of Theorem 5. To prove the theorem, it suffices to show that for any $\mathbf{x} \in \mathbb{R}^n$ with $\|\mathbf{x}\|_2 = 1$, $|\mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{x}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. We will split the entries of $\mathbf{x}$ into level sets according to their magnitude. In particular, let $\bar{\epsilon} = \frac{\epsilon}{\log_2(1/\epsilon)}$ and let $\ell = \log_2(1/\bar{\epsilon})$. Then we can write $\mathbf{x} = \sum_{i=0}^{\ell+1} \mathbf{x}_i$, such that for any $t \in [n]$:

$$\mathbf{x}_i(t) = \begin{cases} \mathbf{x}(t) & \text{if } i \in \{1, 2, \dots, \ell\} \text{ and } |\mathbf{x}(t)| \in \left(\frac{2^{i-1}}{\sqrt{n}}, \frac{2^i}{\sqrt{n}}\right] \\ \mathbf{x}(t) & \text{if } i = 0 \text{ and } |\mathbf{x}(t)| \in \left[0, \frac{1}{\sqrt{n}}\right] \\ \mathbf{x}(t) & \text{if } i = \ell + 1 \text{ and } |\mathbf{x}(t)| \in \left(\frac{1}{\bar{\epsilon}\sqrt{n}}, 1\right] \\ 0 & \text{otherwise.} \end{cases} \quad (4)$$

Via triangle inequality, we have:

$$|\mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{x}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| \leq \sum_{i=0}^{\ell+1} |\mathbf{x}_i^T \mathbf{A} \mathbf{x} - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}|. \quad (5)$$

We will bound each term in (5) by $O(\bar{\epsilon} \cdot \max(n, \|\mathbf{A}\|_1)) = O\left(\frac{\epsilon \cdot \max(n, \|\mathbf{A}\|_1)}{\log(1/\epsilon)}\right) = O\left(\frac{\epsilon \cdot \max(n, \|\mathbf{A}\|_1)}{\ell+2}\right)$. Summing over all $\ell+2$ terms we achieve a final bound of $O(\epsilon \cdot \max(n, \|\mathbf{A}\|_1))$. The theorem then follows by adjusting ϵ by a constant factor.

Fixing $i \in \{0, \dots, \ell\}$, let $\mathbf{x}_L(t) = \mathbf{x}(t)$ if $|\mathbf{x}(t)| \leq \frac{1}{2^i \bar{\epsilon} \sqrt{n}}$ and $\mathbf{x}_L(t) = 0$ otherwise. Let $\mathbf{x}_H(t) = \mathbf{x}(t)$ if $|\mathbf{x}(t)| > \frac{1}{2^i \bar{\epsilon} \sqrt{n}}$ and $\mathbf{x}_H(t) = 0$ otherwise. In the edge case, for $i = \ell + 1$, let $\mathbf{x}_H = \mathbf{x}$ and $\mathbf{x}_L = \mathbf{0}$. Writing $\mathbf{x} = \mathbf{x}_H + \mathbf{x}_L$ and applying triangle inequality, we can bound:

$$|\mathbf{x}_i^T \mathbf{A} \mathbf{x} - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| \leq |\mathbf{x}_i^T \mathbf{A} \mathbf{x}_L - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_L| + |\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_H|. \quad (6)$$

In the following, we separately bound the two terms in (6). Roughly, since $\mathbf{x}_L$ has relatively small entries and thus is relatively well spread, we will be able to show, using a similar approach to Theorem 2, that the first term is small since sparsification with $\mathbf{S}$ approximately preserves $\mathbf{x}_i^T \mathbf{A} \mathbf{x}_L$. On the otherhand, since $\mathbf{x}_H$ has relatively large entries and thus is relatively sparse, we can show that the second term is small since $\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H$ is small (and $\mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_H$ cannot be much larger).

Term 1: Well-Spread Vectors. We start by bounding $|\mathbf{x}_i^T \mathbf{A} \mathbf{x}_L - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_L|$. Let $\mathbf{v}_1, \dots, \mathbf{v}_n$ and $\lambda_1, \dots, \lambda_n$ be the eigenvectors and eigenvalues of $\mathbf{A}$ so that $\mathbf{A} = \sum_{k=1}^n \lambda_k \mathbf{v}_k \mathbf{v}_k^T$. Then we write:

$$\begin{aligned} |\mathbf{x}_i^T \mathbf{A} \mathbf{x}_L - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_L| &= \left| \sum_{k=1}^n \lambda_k \cdot \mathbf{x}_i^T \mathbf{v}_k \mathbf{v}_k^T \mathbf{x}_L - \sum_{k=1}^n \lambda_k \mathbf{x}_i^T (\mathbf{v}_k \mathbf{v}_k^T \circ \mathbf{S}) \mathbf{x}_L \right| \\ &\leq \sum_{k=1}^n |\lambda_k| \cdot |\mathbf{x}_i^T [\mathbf{v}_k \mathbf{v}_k^T \circ (\mathbf{1} - \mathbf{S})] \mathbf{x}_L|. \end{aligned}$$

As in the proof of Theorem 2, if we let $\mathbf{D}_k \in \mathbb{R}^{n \times n}$ be a diagonal matrix with the entries of $\mathbf{v}_k$ on its diagonal, then we have $\mathbf{v}_k \mathbf{v}_k^T \circ (\mathbf{1} - \mathbf{S}) = \mathbf{D}_k (\mathbf{1} - \mathbf{S}) \mathbf{D}_k$. Further, $\mathbf{x}_i^T \mathbf{D}_k (\mathbf{1} - \mathbf{S}) \mathbf{D}_k \mathbf{x}_L = \mathbf{v}_k^T \mathbf{D}_i (\mathbf{1} - \mathbf{S}) \mathbf{D}_L \mathbf{v}_k$, where $\mathbf{D}_i, \mathbf{D}_L$ are diagonal with the entries of $\mathbf{x}_i$ and $\mathbf{x}_L$ on their diagonals. Plugging this back in we have:

$$\begin{aligned} |\mathbf{x}_i^T \mathbf{A} \mathbf{x}_L - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_L| &\leq \sum_{k=1}^n |\lambda_k| \cdot |\mathbf{v}_k^T \mathbf{D}_i (\mathbf{1} - \mathbf{S}) \mathbf{D}_L \mathbf{v}_k| \\ &\leq \|\mathbf{1} - \mathbf{S}\|_2 \cdot \sum_{k=1}^n |\lambda_k| \cdot \|\mathbf{v}_k^T\|_2 \cdot \|\mathbf{D}_i\|_2 \cdot \|\mathbf{D}_L\|_2 \cdot \|\mathbf{v}_k\|_2 \\ &\leq \|\mathbf{1} - \mathbf{S}\|_2 \cdot \sum_{k=1}^n |\lambda_k| \cdot \|\mathbf{D}_i\|_2 \cdot \|\mathbf{D}_L\|_2. \end{aligned} \tag{7}$$

Finally, observe that by definition, for any $i \in \{0, \dots, \ell\}$, $\|\mathbf{D}_i\|_2 = \max_{t \in [n]} |\mathbf{x}_i(t)| \leq \frac{2^i}{\sqrt{n}}$ and $\|\mathbf{D}_L\|_2 = \max_{t \in [n]} |\mathbf{x}_L(t)| \leq \frac{1}{2^i \bar{\epsilon} \sqrt{n}}$. Thus, $\|\mathbf{D}_i\|_2 \cdot \|\mathbf{D}_L\|_2 \leq \frac{1}{\bar{\epsilon} n}$. For $i = \ell + 1$, $\mathbf{x}_L = \mathbf{0}$ by definition and so, the bound $\|\mathbf{D}_i\|_2 \cdot \|\mathbf{D}_L\|_2 \leq \frac{1}{\bar{\epsilon} n}$ holds vacuously.

Also, by the assumption of the theorem, $\|\mathbf{1} - \mathbf{S}\|_2 \leq \frac{\epsilon^2 n}{c \log^2(1/\epsilon)} \leq \bar{\epsilon}^2 n$. Plugging back into (7),

$$|\mathbf{x}_i^T \mathbf{A} \mathbf{x}_L - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_L| \leq \bar{\epsilon}^2 n \cdot \frac{1}{\bar{\epsilon} n} \cdot \|\mathbf{A}\|_1 \leq \bar{\epsilon} \cdot \|\mathbf{A}\|_1, \tag{8}$$

as required.

Term 2: Sparse Vectors. We next bound the second term of (6): $|\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_H|$. We write $\mathbf{x}_i = \mathbf{x}_{i,P} + \mathbf{x}_{i,N}$, where $\mathbf{x}_{i,P}$ and $\mathbf{x}_{i,N}$ contain its positive and non-positive entries respectively. Similarly, write $\mathbf{x}_H = \mathbf{x}_{H,P} + \mathbf{x}_{H,N}$. We can then bound via triangle inequality:

$$\begin{aligned} |\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_H| &\leq |\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H| + |\mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_H| \\ &\leq |\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H| + |\mathbf{x}_{i,P}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_{H,P}| + |\mathbf{x}_{i,P}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_{H,N}| \\ &\quad + |\mathbf{x}_{i,N}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_{H,P}| + |\mathbf{x}_{i,N}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_{H,N}|. \end{aligned} \tag{9}$$

We will bound each term in (9) by $O(\bar{\epsilon} n)$, giving that overall $|\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_H| = O(\bar{\epsilon} n)$. We first observe that

$$|\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H| \leq \|\mathbf{x}_i\|_1 \cdot \|\mathbf{x}_H\|_1 \cdot \|\mathbf{A}\|_\infty.$$

By assumption $\|\mathbf{A}\|_\infty \leq 1$. Further, for $i \in \{1, \dots, \ell + 1\}$, since $|\mathbf{x}_i(t)| \geq \frac{2^{i-1}}{\sqrt{n}}$ for all t and since $\|\mathbf{x}_i\|_2 \leq 1$, $\mathbf{x}_i$ has at most $\frac{n}{2^{2i-2}}$ non-zero entries. Thus, $\|\mathbf{x}_i\|_1 \leq \frac{\sqrt{n}}{2^{i-1}}$. For $i = 0$, this bound holds trivially since $\|\mathbf{x}_i\|_1 \leq \sqrt{n} \cdot \|\mathbf{x}_i\|_2 \leq \sqrt{n} \leq \frac{\sqrt{n}}{2^{i-1}}$.

Similarly, for $i \in \{0, \dots, \ell\}$, since $\|\mathbf{x}_H\|_2 \leq 1$ and since $|\mathbf{x}_H(t)| \geq \frac{1}{2^i \bar{\epsilon} \sqrt{n}}$ by definition, $\mathbf{x}_H$ has at most $2^{2i} \cdot \bar{\epsilon}^2 n$ non-zero entries. So $\|\mathbf{x}_H\|_1 \leq 2^i \bar{\epsilon} \cdot \sqrt{n}$. For $i = \ell + 1 = \log_2(1/\bar{\epsilon}) + 1$ this bound holds trivially since $\|\mathbf{x}_H\|_1 \leq \sqrt{n} \cdot \|\mathbf{x}_H\|_2 \leq 2^i \bar{\epsilon} \cdot \sqrt{n}$. Putting these all together, we have

$$|\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H| \leq \|\mathbf{x}_i\|_1 \cdot \|\mathbf{x}_H\|_1 \cdot \|\mathbf{A}\|_\infty \leq \frac{\sqrt{n}}{2^{i-1}} \cdot 2^i \bar{\epsilon} \cdot \sqrt{n} = 2\bar{\epsilon}n. \quad (10)$$

We next bound $|\mathbf{x}_{i,P}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_{H,P}|$. Since $\mathbf{x}_{i,P}$ and $\mathbf{x}_{H,P}$ are both all positive vectors, and since by assumption $\|\mathbf{A}\|_\infty \leq 1$,

$$\begin{aligned} |\mathbf{x}_{i,P}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_{H,P}| &\leq |\mathbf{x}_{i,P}^T \mathbf{S} \mathbf{x}_{H,P}| \\ &\leq |\mathbf{x}_{i,P}^T \mathbf{1} \mathbf{x}_{H,P}| + |\mathbf{x}_{i,P}^T (\mathbf{1} - \mathbf{S}) \mathbf{x}_{H,P}| \\ &\leq \|\mathbf{x}_{i,P}\|_1 \cdot \|\mathbf{x}_{H,P}\|_1 \cdot \|\mathbf{1}\|_\infty + \|\mathbf{x}_{i,P}\|_2 \cdot \|\mathbf{x}_{H,P}\|_2 \cdot \|\mathbf{1} - \mathbf{S}\|_2. \end{aligned} \quad (11)$$

Following the same argument used to show (10), $\|\mathbf{x}_{i,P}\|_1 \cdot \|\mathbf{x}_{H,P}\|_1 \cdot \|\mathbf{1}\|_\infty \leq 2\bar{\epsilon}n$. Further, since by the assumption of the theorem, $\|\mathbf{1} - \mathbf{S}\|_2 \leq \frac{\epsilon^2 n}{c \log^2(1/\epsilon)} \leq \bar{\epsilon}^2 n \leq \bar{\epsilon}n$, and since $\mathbf{x}_{i,P}$ and $\mathbf{x}_{H,P}$ both are at most unit norm, $\|\mathbf{x}_{i,P}\|_2 \cdot \|\mathbf{x}_{H,P}\|_2 \cdot \|\mathbf{1} - \mathbf{S}\|_2 \leq \bar{\epsilon}n$. Thus, plugging back into (11), we have $|\mathbf{x}_{i,P}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_{H,P}| \leq 3\bar{\epsilon}n$. Identical bounds will hold for the remaining three terms of (9), since for each, the two vectors in the quadratic form have entries that either always match or never match on sign. Plugging these bounds and (10) back into (9), we obtain

$$|\mathbf{x}_i^T \mathbf{A} \mathbf{x}_H - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}_H| \leq 2\bar{\epsilon}n + 4 \cdot 3\bar{\epsilon}n \leq 14\bar{\epsilon}n, \quad (12)$$

which completes the required bound in the sparse case.

Concluding the Proof. We finally plug our bounds for the sparse case (12) and the well-spread case (8) into (6) to obtain:

$$|\mathbf{x}_i^T \mathbf{A} \mathbf{x} - \mathbf{x}_i^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| \leq 14\bar{\epsilon}n + \bar{\epsilon} \|\mathbf{A}\|_1 \leq 15\bar{\epsilon} \cdot \max(n, \|\mathbf{A}\|_1).$$

Plugging this bound into (5) gives that $|\mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{x}^T (\mathbf{A} \circ \mathbf{S}) \mathbf{x}| = O(\epsilon \cdot \max(n, \|\mathbf{A}\|_1))$, which completes the theorem after adjusting ϵ by a constant factor. $\square$

4 Spectral Approximation Lower Bounds

We now show that our universal sparsifier upper bounds for both PSD matrices (Theorem 1) and non-PSD matrices (Theorem 4) are nearly tight. We also give $\Omega(n/\epsilon^2)$ query lower bounds against general deterministic (possibly adaptive) spectral approximation algorithms (Theorem 7) and general randomized spectral approximation algorithms (Theorem 11).

4.1 Sparsity Lower Bound for PSD Matrices

We first prove that every matrix which is an ϵn spectral approximation to the all-ones matrix must have $\Omega(\frac{n}{\epsilon^2})$ non-zero entries. This shows that Theorem 1 is optimal up to constant factors, even for algorithms that sparsify just a single bounded entry PSD matrix. The idea of the lower bound is simple: if a matrix $\mathbf{S}$ spectrally approximates the all-ones matrix, its entries must sum to $\Omega(n^2)$. Thus, if $\mathbf{S}$ has just s non-zero entries, it must have Frobenius norm at least $\Omega(n^2/\sqrt{s})$. Unless $s = \Omega(n/\epsilon^2)$, this Frobenius norm is too large for $\mathbf{S}$ to be a ϵn -spectral approximation of the all ones matrix (which has $\|\mathbf{1}\|_F = n$).

Theorem 3 (Sparsity Lower Bound – PSD Matrices). *Let $\mathbf{1} \in \mathbb{R}^{n \times n}$ be the all-ones matrix. Then, for any $\epsilon \in (0, 1/2)$ with $\epsilon \geq c/\sqrt{n}$ for large enough constant c , any $\mathbf{S} \in \mathbb{R}^{n \times n}$ with $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$ must have $\Omega(n/\epsilon^2)$ nonzero entries.*

Proof. If we let $\mathbf{x} \in \mathbb{R}^n$ be the all ones vector, then since $\|\mathbf{x}\|_2^2 = n$, $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$ implies that $|\mathbf{x}^T \mathbf{S} \mathbf{x} - \mathbf{x}^T \mathbf{1} \mathbf{x}| = |\mathbf{x}^T \mathbf{S} \mathbf{x} - n^2| \leq \epsilon n^2$. This in turn implies:

$$\sum_{i,j \in [n]} |\mathbf{S}_{ij}| \geq \sum_{i,j \in [n]} \mathbf{S}_{ij} \geq n^2 - \epsilon n^2 \geq \frac{n^2}{2},$$

where the last inequality follows from the assumption that $\epsilon \leq \frac{1}{2}$. If $\mathbf{S}$ has s non-zero entries, since $\sum_{i,j \in [n]} |\mathbf{S}_{ij}|$ is the ℓ_1 -norm of the entries, and $\|\mathbf{S}\|_F$ is the ℓ_2 -norm of the entries, we conclude by $\ell_1 - \ell_2$ norm equivalence that $\|\mathbf{S}\|_F \geq \frac{1}{\sqrt{s}} \sum_{i,j \in [n]} |\mathbf{S}_{ij}| \geq \frac{n^2}{2\sqrt{s}}$. Therefore, by the triangle inequality,

$$\|\mathbf{1} - \mathbf{S}\|_F \geq \|\mathbf{S}\|_F - \|\mathbf{1}\|_F = \frac{n^2}{2\sqrt{s}} - n \geq \frac{n^2}{4\sqrt{s}},$$

as long as $s \leq \frac{n^2}{16}$ so that $\frac{n^2}{2\sqrt{s}} \geq 2n$. Now, using that $\|\mathbf{A}\|_F^2 = \sum_{i=1}^n \sigma_i^2(\mathbf{A}) \leq n \cdot \|\mathbf{A}\|_2^2$, we have:

$$\|\mathbf{1} - \mathbf{S}\|_F \geq \frac{n^2}{4\sqrt{s}} \implies \|\mathbf{1} - \mathbf{S}\|_F^2 \geq \frac{n^4}{16s} \implies \|\mathbf{1} - \mathbf{S}\|_2^2 \geq \frac{n^3}{16s} \implies \|\mathbf{1} - \mathbf{S}\|_2 \geq \frac{n^{3/2}}{4\sqrt{s}}.$$

By our assumption that $\|\mathbf{1} - \mathbf{S}\|_2 \leq \epsilon n$, this means that $s = \Omega\left(\frac{n}{\epsilon^2}\right)$, concluding the theorem. $\square$

4.2 Lower Bounds for Deterministic Approximation of Non-PSD Matrices

We next show that our universal sparsification bound for non-PSD matrices (Theorem 4) is tight up to a $\log^4(1/\epsilon)$ factor. Our lower bound holds even for the easier problem of top singular value (spectral norm) approximation and against a more general class of algorithms, which non-adaptively and deterministically query entries of the input matrix. We show how to extend the lower bound to possibly adaptive deterministic algorithms in Theorem 7, but with a $1/\epsilon^2$ factor loss.

Theorem 6 (Non-Adaptive Query Lower Bound for Deterministic Spectral Approximation of Non-PSD Matrices). *For any $\epsilon \in (1/n^{1/4}, 1/4)$, any deterministic algorithm that queries entries of a bounded entry matrix $\mathbf{A}$ non-adaptively and outputs $\tilde{\sigma}_1$ satisfying $|\sigma_1(\mathbf{A}) - \tilde{\sigma}_1| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ must read at least $\Omega\left(\frac{n}{\epsilon^4}\right)$ entries. Here $\sigma_1(\mathbf{A}) = \|\mathbf{A}\|_2$ is the largest singular value of $\mathbf{A}$.*

Proof. Assume that we have a deterministic algorithm $\mathcal{A}$ that non-adaptively reads s entries of any bounded symmetric matrix $\mathbf{A}$ and outputs $\tilde{\sigma}_1$ with $|\sigma_1(\mathbf{A}) - \tilde{\sigma}_1| \leq \epsilon \max(n, \|\mathbf{A}\|_1)$. Assume for the sake of contradiction that $s \leq \frac{cn}{\epsilon^4}$ for some sufficiently small constant c .

Let $T \subset [n]$ be a set of $n^{3/2}/s^{1/2}$ rows on which $\mathcal{A}$ reads at most $\sqrt{ns}$ entries. Such a subset must exist since the average number of entries read in any set of $n^{3/2}/s^{1/2}$ rows is $\frac{n^{3/2}}{s^{1/2}} \cdot n \cdot \frac{s}{n^2} = \sqrt{ns}$.

Let $\mathbf{1}_T$ be the matrix which is 1 on all the rows in T and zero everywhere else. Let $\mathbf{S}_T$ be the matrix which matches $\mathbf{1}_T$ on all entries, except is 0 on any entry in the rows of T that is not read by the algorithm $\mathcal{A}$. Observe that $\mathcal{A}$ reads the same entries (all ones) and thus outputs the same approximation $\tilde{\sigma}_1$ for $\mathbf{1}_T$ and $\mathbf{S}_T$. We now bound the allowed error of this approximation. $\mathbf{1}_T$ is rank-1 and so:

$$\|\mathbf{1}_T\|_1 = \sigma_1(\mathbf{1}_T) = \|\mathbf{1}_T\|_F = \frac{n^{5/4}}{s^{1/4}}.$$

We have $\sigma_1(\mathbf{S}_T) \leq \|\mathbf{S}_T\|_F \leq n^{1/4}s^{1/4}$ since $\mathbf{S}_T$ has just $\sqrt{ns}$ entries set to 1 – the entries that $\mathcal{A}$ reads in the rows of T . Using that $\mathbf{S}_T$ is supported only on the $n^{3/2}/s^{1/2}$ rows of T , $\mathbf{S}_T$ has rank at most $n^{3/2}/s^{1/2}$. Thus, we can bound:

$$\|\mathbf{S}_T\|_1 \leq \frac{n^{3/4}}{s^{1/4}} \cdot \|\mathbf{S}_T\|_F \leq n.$$

Now, since $|\sigma_1(\mathbf{1}_T) - \sigma_1(\mathbf{S}_T)| \geq \left| \frac{n^{5/4}}{s^{1/4}} - n^{1/4}s^{1/4} \right|$, our algorithm incurs error on one of the two input instances at least

$$\frac{\left| \frac{n^{5/4}}{s^{1/4}} - n^{1/4}s^{1/4} \right|}{2} \geq \frac{n^{5/4}}{4s^{1/4}},$$

where the inequality follows since, by assumption, $s \leq \frac{cn}{\epsilon^4}$ for some small constant c and $\epsilon \geq \frac{1}{n^{1/4}}$, and thus $n^{1/4}s^{1/4} \leq \frac{n^{5/4}}{2s^{1/4}}$.

The above is a contradiction when $\epsilon < 1/4$ since the above error is at least $1/4 \cdot \|\mathbf{1}_T\|_1$ and further, for $s \leq \frac{cn}{\epsilon^4}$, the error is at least $\frac{\epsilon n}{4c^{1/4}} > \epsilon n \geq \epsilon \cdot \|\mathbf{S}_T\|_1$ if we set c small enough. Thus, on at least one of the two input instances the error exceeds $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$, yielding a contradiction. $\square$

We can prove a variant on Theorem 6 when the algorithm is allowed to make adaptive queries. Here, our lower bound reduces to $\Omega(n/\epsilon^2)$, as we are not able to restrict our hard case to a small set of rows of the input matrix. Closing the gap here – either by giving a stronger lower bound in the adaptive case or giving an adaptive query deterministic algorithm that achieves $\tilde{O}(n/\epsilon^2)$ query complexity is an interesting open question.

Theorem 7 (Adaptive Query Lower Bound for Deterministic Spectral Approximation of Non-PSD Matrices). *For any $\epsilon \in (1/\sqrt{n}, 1/4)$, any deterministic algorithm that queries entries of a bounded entry matrix $\mathbf{A}$ (possibly adaptively) and outputs $\tilde{\sigma}_1$ satisfying $|\sigma_1(\mathbf{A}) - \tilde{\sigma}_1| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ must read at least $\Omega\left(\frac{n}{\epsilon^2}\right)$ entries. Here $\sigma_1(\mathbf{A}) = \|\mathbf{A}\|_2$ is the largest singular value of $\mathbf{A}$.*

Proof. Assume that we have a deterministic algorithm $\mathcal{A}$ that reads at most s entries of any bounded entry matrix $\mathbf{A}$ and outputs $\tilde{\sigma}_1$ with $|\sigma_1(\mathbf{A}) - \tilde{\sigma}_1| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Assume for the sake of contradiction that $s \leq \frac{cn}{\epsilon^2}$ for some sufficiently small constant c .

Let S be the set of entries that $\mathcal{A}$ reads when given the all ones matrix $\mathbf{1}$ as input. Let $\mathbf{S}$ be the matrix which is one on every entry in S and zero elsewhere. Observe that $\mathcal{A}$ reads the same entries and thus outputs the same approximation $\tilde{\sigma}_1$ for $\mathbf{1}$ and $\mathbf{S}$. We now bound the allowed error of this approximation. $\mathbf{1}$ is rank-1 and has $\|\mathbf{1}\|_1 = \sigma_1(\mathbf{1}) = n$. We have $\sigma_1(\mathbf{S}) \leq \|\mathbf{S}\|_F = \sqrt{s}$ and can bound $\|\mathbf{S}\|_1 \leq \sqrt{n} \cdot \|\mathbf{S}\|_F \leq \sqrt{sn} \leq \frac{c^{1/2}n}{\epsilon}$, where the last bound follows from our assumption that $s \leq \frac{cn}{\epsilon^2}$. Now, since $|\sigma_1(\mathbf{1}) - \sigma_1(\mathbf{S})| \geq |n - \sqrt{s}|$, our algorithm incurs an error on one of the two input instances at least

$$\frac{|n - \sqrt{s}|}{2} \geq \frac{n}{4},$$

where the inequality holds since, by assumption, $s \leq \frac{cn}{\epsilon^2}$ for some small constant c and $\epsilon \geq \frac{1}{n^{1/2}}$, and thus $\sqrt{s} \leq \frac{n}{2}$.

The above is a contradiction when $\epsilon < 1/4$ since the error mentioned above is at least $1/2 \cdot \|\mathbf{1}\|_1 = 1/2 \cdot \max(n, \|\mathbf{1}\|_1)$. Furthermore, since we have bounded $\|\mathbf{S}\|_1 \leq \frac{c^{1/2}n}{\epsilon}$, the error is greater than $\epsilon \cdot \max(n, \|\mathbf{S}\|_1)$ when $c < 1$. Thus, on at least one of the two input instances the error exceeds $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$, yielding a contradiction. $\square$

4.3 Lower Bound for Randomized Approximation of Non-PSD Matrices

Finally, we prove Theorem 11, which gives an $\Omega(n/\epsilon^2)$ query lower bound for spectral approximation of bounded entry matrices that holds even for randomized and adaptive algorithms that approximate $\mathbf{A}$ in an arbitrary manner (not necessarily via sparsification). In particular, we show the lower bound against algorithms that produce any data structure, $f(\cdot, \cdot)$, that satisfies $|f(\mathbf{x}, \mathbf{y}) - \mathbf{x}^T \mathbf{A} \mathbf{y}| \leq \epsilon n$ for all unit norm $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$.

To prove this lower bound, we let $\mathbf{A}$ be a random matrix where each row has binary entries that are either i.i.d. fair coin flips, or else are coin flips with bias $+\epsilon$. Each row is unbiased with probability $1/2$ and biased with probability $1/2$. We show that an ϵn -spectral approximation to $\mathbf{A}$ suffices to identify for at least a $9/10$ fraction of the rows, whether or not they are biased – roughly since the approximation must preserve the fact that $\mathbf{x}^T \mathbf{A} \mathbf{x}$ is large when $\mathbf{x}$ is supported just on this biased set. Consider a communication problem in the blackboard model, in which each of n^2 players can access just a single entry of $\mathbf{A}$. It is known that if the n players corresponding to a single row of the matrix want to identify with good probability whether or not it is biased, they must communicate at least $\Omega(1/\epsilon^2)$ bits [SWXZ22]. Further, via a direct sum type argument, we can show that for the n^2 players to identify the bias of a $9/10$ fraction of the rows with good probability, $\Omega(n/\epsilon^2)$ bits must be communicated. I.e., at least $\Omega(n/\epsilon^2)$ of the players must read their input bits, yielding our query complexity lower bound.

Note that there may be other proof approaches here, based on a direct sum of n 2-player Gap-Hamming instances [BJKS04, CR12b, BGPW13], but our argument is simple and already gives optimal bounds.

Section Roadmap. In Section 4.3.1, we formally define the problem of identifying the bias of a large fraction of $\mathbf{A}$'s rows as the (ϵ, n) -distributed detection problem (Definition 4.2). We prove a $\Omega(n/\epsilon^2)$ query lower bound for this problem in Lemma 3.. Then, in Section 4.3.2, we prove Theorem 11 by showing a reduction from the distributed detection problem to spectral approximation.

4.3.1 Distributed Detection Lower Bound

Here, we define additional concepts and notation specific to this section. For random variables X , Y and Z , let $H(X)$ denote the entropy, $H(X|Y)$ denote the conditional entropy, $I(X; Y)$ denote mutual information, and $I(X; Y|Z)$ denote conditional mutual information [Cov99]. We also use some ideas from communication complexity. Namely, we work in the *blackboard model*, where T parties communicate by posting messages to a public blackboard with access to public randomness with unlimited rounds of adaptivity. Let $\Pi \in \{0, 1\}^*$ denote the transcript of a protocol posted to the blackboard, and let $|\Pi|$ denote the length of a transcript. For a fixed protocol with fixed probability of success, we may assume, without loss of generality, that Π has a fixed pre-determined length (see Section 2.2 of [Rou16]).

Our query complexity lower bound for spectral approximation will follow from a reduction from the following testing problem.

Definition 4.1. (ϵ -Distributed detection problem [SWXZ22]). For fixed distributions $\mu_0 = \text{Bernoulli}(1/2)$ and $\mu_1 = \text{Bernoulli}(1/2 + \epsilon)$, with $\epsilon \in [0, \frac{1}{2}]$, let $\mathbf{x} \in \{0, 1\}^n$ be a random vector such that $\mathbf{x}_1, \dots, \mathbf{x}_n$ are sampled i.i.d. from μ_V , for $V \in \{0, 1\}$. The distributed detection problem is the task of determining whether $V = 0$ or $V = 1$, given the values of $\mathbf{x}$.

This decision problem can be naturally interpreted as a communication problem in the blackboard model, where n players each have a single private bit corresponding to whether a unique entry of $\mathbf{x}$ is zero or one. Prior work takes this view to lower bound the mutual information between the

transcript of a protocol which correctly solves the ϵ -Distributed detection problem with constant advantage and the sampled bits in $\mathbf{x}$ in the case $V = 0$.

Theorem 12. (Theorem 6 in [SWXZ22]). Let Π be the transcript of a protocol that solves ϵ -distributed detection problem with probability $1 - p$ for any fixed choice of $p \in [0, 0.5)$. Then

$$I(X; \Pi | V = 0) = \Omega(\epsilon^{-2}).$$

Next, we define the (ϵ, n) -Distributed detection problem, which combines n length- n ϵ -Distributed detection problems into a single joint detection problem.

Definition 4.2. ((ϵ, n) -Distributed detection problem) For $\mathbf{v} \in \{0, 1\}^n$ distributed uniformly on the Hamming cube, generate the matrix $\mathbf{A} \in \{0, 1\}^{n \times n}$ such that if $\mathbf{v}_i = 0$, then all entries in the i -th row of $\mathbf{A}$ are i.i.d. samples from $\text{Bernoulli}(1/2)$. Otherwise, let all entries in the i -th row be sampled from $\text{Bernoulli}(1/2 + \epsilon)$. The (ϵ, n) -Distributed detection problem is the task of recovering a vector $\hat{\mathbf{v}} \in \{0, 1\}^n$ such that $\|\hat{\mathbf{v}} - \mathbf{v}\|_1 \leq \frac{n}{20}$.

Again, this vector recovery problem has a natural interpretation as a communication problem, where n^2 players each hold a single bit of information that corresponds to whether a unique entry of $\mathbf{A}$ is zero or one. Throughout this section, we will view Definition 4.2 as a decision problem or communication problem as needed. Let Π be the transcript of a protocol that solves the (ϵ, n) -Distributed detection problem in the communication model introduced at the beginning of this section. Lower bounding the mutual information between the transcript, Π , and the private information held by the players (the entries of $\mathbf{A}$) lower bounds the length of the transcript via the following argument.

Lemma 1. If Π is the transcript of a protocol that solves the (ϵ, n) -Distributed detection problem and $I(\mathbf{A}; \Pi) \geq b$, then $|\Pi| \geq \frac{b}{\log 2}$.

Proof. For random variables X, Y , $I(X; Y) \leq \min\{H(X), H(Y)\}$ [Cov99]. If a random variable X has finite support, then $H(X) \leq \log |\text{supp}(X)|$. Recall that $|\Pi|$ is the number of bits in the transcript and hence $\log 2^{|\Pi|} = \log 2 \cdot |\Pi| \geq H(\Pi) \geq I(\mathbf{A}; \Pi) \geq b$. Hence, we conclude the statement. $\square$

Next, we lower bound the number of entries in the matrix $\mathbf{A}$ that must be observed to solve a variant of the (ϵ, n) -Distributed detection problem, where we aim to correctly decide each index of $\mathbf{v}$ with constant success probability, rather than constructing $\hat{\mathbf{v}}$ such that $\|\hat{\mathbf{v}} - \mathbf{v}\|_1 \leq \frac{n}{20}$. There are three sources of randomness in this problem: 1) the initial randomness in sampling $\mathbf{v}$, 2) the random variable $\mathbf{A}$ which depends on $\mathbf{v}$, and 3) the random transcript Π which depends on $\mathbf{A}$. When necessary to avoid confusion, we will be explicit regarding which of these variables are considered fixed and which are considered as random in probabilistic statements.

Lemma 2. Let $\mathbf{A}$ and $\mathbf{v}$ be distributed as in the (ϵ, n) -Distributed detection problem. Any protocol with transcript Π that constructs a vector $\hat{\mathbf{v}}$ such that $\mathbb{P}(\mathbf{v}_i = \hat{\mathbf{v}}_i) \geq \frac{9}{10}$, for all $i \in [n]$, (where the randomness is with respect to $\mathbf{A}$, $\mathbf{v}$, and Π), must observe $\Omega(\frac{n}{\epsilon^2})$ entries of $\mathbf{A}$ in the worst case.

Proof. Consider $\mathbf{A}$ as n^2 players each holding a single bit of information corresponding to whether a unique entry of $\mathbf{A}$ is zero or one. Here, let $\Pi \in \{0, 1\}^*$ be the transcript of a protocol that satisfies $\mathbb{P}(\mathbf{v}_i = \hat{\mathbf{v}}_i) \geq \frac{9}{10}$ for every $i \in [n]$. Note that any algorithm which solves this problem by querying k entries of $\mathbf{A}$ implies a protocol for the communication problem which communicates k bits, since the algorithm could be simulated by (possibly adaptively) posting each queried bit to the blackboard.

First, we lower bound the un-conditional mutual information between the private information in the i -th row of $\mathbf{A}$ and the transcript Π . Note that $\mathbf{v}_i$ is one or zero with equal probability, therefore,

$$I(\mathbf{A}_i; \Pi | \mathbf{v}_i) = \frac{1}{2}I(\mathbf{A}_i; \Pi | \mathbf{v}_i = 0) + \frac{1}{2}I(\mathbf{A}_i; \Pi | \mathbf{v}_i = 1) \geq \frac{1}{2}I(\mathbf{A}_i; \Pi | \mathbf{v}_i = 0).$$

Next, we decompose the mutual information by its definition and the chain rule, then, we use the fact that $0 \leq H(\mathbf{v}_i) \leq 1$ and that conditioning can only decrease the entropy of a random variable.

$$\begin{aligned} I(\mathbf{A}_i; \Pi) &= I(\mathbf{A}_i; \Pi | \mathbf{v}_i) + H(\mathbf{v}_i | \mathbf{A}_i, \Pi) - H(\mathbf{v}_i | \mathbf{A}_i) - H(\mathbf{v}_i | \Pi) + H(\mathbf{v}_i) \\ &\Rightarrow I(\mathbf{A}_i; \Pi) \geq I(\mathbf{A}_i; \Pi | \mathbf{v}_i) - 2. \end{aligned}$$

Observe that determining $\hat{\mathbf{v}}_i$ with the guarantee $\mathbb{P}(\mathbf{v}_i = \hat{\mathbf{v}}_i) \geq \frac{9}{10}$ solves the ϵ -Distributed detection problem with probability at least $\frac{9}{10}$. Therefore, by Theorem 12 $I(\mathbf{A}_i; \Pi | \mathbf{v}_i = 0) = \Omega(\epsilon^{-2})$, and so we can use the previous two equations to lower bound for the conditional mutual information:

$$I(\mathbf{A}_i; \Pi) \geq \frac{1}{2}I(\mathbf{A}_i; \Pi | \mathbf{v}_i = 0) - 2 = \Omega\left(\frac{1}{\epsilon^2}\right). \quad (13)$$

Next, we lower bound the total mutual information between $\mathbf{A}$ and Π . Mirroring the argument of Lemma 1 in [SWXZ22], we use that, for independent random variables, entropy is additive and conditional entropy is subadditive to lower bound the mutual information between $\mathbf{A}$ and Π :

$$\begin{aligned} I(\{\mathbf{A}_i\}_{i \in [n]}; \Pi) &= H(\{\mathbf{A}_i\}_{i \in [n]}) - H(\{\mathbf{A}_i\}_{i \in [n]} | \Pi) \\ &\geq \sum_{i=1}^n H(\mathbf{A}_i) - H(\mathbf{A}_i, \Pi) \\ &= \sum_{i=1}^n I(\mathbf{A}_i; \Pi) = \Omega\left(\frac{n}{\epsilon^2}\right), \end{aligned}$$

where the last step follows from (13). By Lemma 1, $|\Pi| = \Omega(I(\{\mathbf{A}_i\}_{i \in [n]}; \Pi)) = \Omega(\frac{n}{\epsilon^2})$. Therefore, every algorithm which samples entries of $\mathbf{A}$ to construct a vector $\hat{\mathbf{v}}$ satisfying $\mathbb{P}(\mathbf{v}_i = \hat{\mathbf{v}}_i) \geq \frac{9}{10}$ must observe at least $\Omega(\frac{n}{\epsilon^2})$ entries of $\mathbf{A}$. $\square$

We now have the necessary results to prove a lower bound on the number of entries of $\mathbf{A}$ that must be observed to solve the (ϵ, n) -Distributed detection problem, which we do by reducing to the problem in Lemma 2.

Lemma 3. *Any adaptive randomized algorithm which solves the (ϵ, n) -Distributed detection problem with probability at least $\frac{19}{20}$ (with respect to the randomness in $\mathbf{v}$, $\mathbf{A}$, and Π) must observe $\Omega(\frac{n}{\epsilon^2})$ entries of $\mathbf{A}$.*

Proof. Any algorithm that can produce a vector $\hat{\mathbf{v}}$ such that $\|\hat{\mathbf{v}} - \mathbf{v}\|_1 \leq \frac{n}{20}$ with probability at least $\frac{19}{20}$ could be used to guarantee that $\mathbb{P}(\hat{\mathbf{v}}_i = \mathbf{v}_i) \geq \frac{9}{10}$ by the following argument.

For any input matrix $\mathbf{A}$, create the matrix $\bar{\mathbf{A}}$ such that $\bar{\mathbf{A}}_i = \mathbf{A}_{\sigma(i)}$, where σ is a permutation sampled uniformly from the symmetric group of size n . Let $\mathbf{v}_\sigma$ be the vector which satisfies $[\mathbf{v}_\sigma]_i = \mathbf{v}_{\sigma(i)}$. Run the considered protocol on $\bar{\mathbf{A}}$ to recover $\mathbf{u}$ such that $\|\mathbf{u} - \mathbf{v}_\sigma\|_1 \leq \frac{n}{20}$ with probability $\frac{19}{20}$. Finally, let $\hat{\mathbf{v}}_i = \mathbf{u}_{\sigma^{-1}(i)}$ for all $i \in [n]$.

Then, $\mathbb{P}(\hat{\mathbf{v}}_i = \mathbf{v}_i) = \mathbb{P}(\|\mathbf{u} - \mathbf{v}_\sigma\|_1 \leq \frac{n}{20}) \cdot \mathbb{P}(\mathbf{u}_{\sigma^{-1}(i)} - \mathbf{v}_{\sigma^{-1}(i)} | \|\mathbf{u} - \mathbf{v}_\sigma\|_1 \leq \frac{n}{20}) = \frac{19}{20} \cdot \frac{19}{20} \geq \frac{9}{10}$. Since $\sigma(i)$ is uniformly distributed over $[n]$, $\mathbb{P}(\mathbf{u}_{\sigma^{-1}(i)} - \mathbf{v}_{\sigma^{-1}(i)} | \|\mathbf{u} - \mathbf{v}_\sigma\|_1 \leq \frac{n}{20})$ is the number of correctly decided entries divided by n .

By Lemma 2, we conclude that solving the (ϵ, n) -Distributed detection with probability at least $\frac{19}{20}$ requires observing $\Omega(\frac{n}{\epsilon^2})$ entries of $\mathbf{A}$ in the worst case. $\square$

4.3.2 Spectral Approximation Query Lower Bound

Next, we show that the (ϵ, n) -Distributed detection problem can be solved using a spectral approximation of $\mathbf{A}$, thereby implying a query complexity lower bound for spectral approximation.

Theorem 11 (Lower Bound for Randomized Spectral Approximation). *Any randomized algorithm that (possibly adaptively) reads entries of a binary matrix $\mathbf{A} \in \{0, 1\}^{n \times n}$ to construct a data structure $f : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ that, for $\epsilon \in (0, 1)$ satisfies with probability at least $99/100$,*

$$|f(\mathbf{x}, \mathbf{y}) - \mathbf{x}^T \mathbf{A} \mathbf{y}| \leq \epsilon n, \text{ for all unit } \mathbf{x}, \mathbf{y} \in \mathbb{R}^n,$$

must read $\Omega\left(\frac{n}{\epsilon^2}\right)$ entries of $\mathbf{A}$ in the worst case, provided $\epsilon = \Omega\left(\frac{\log n}{\sqrt{n}}\right)$.

Proof. We reduce solving the (ϵ, n) -Distributed detection problem to constructing a spectral approximation satisfying the guarantee in the above theorem statement. Throughout this proof, let ϵ_d be the parameter associated with the (ϵ, n) -Distributed detection problem and ϵ_s be the accuracy of the spectral approximation.

Let $\mathbf{A}$ be the matrix associated with the (ϵ_d, n) -Distributed detection problem. Suppose that by reading r entries of $\mathbf{A}$, we can create a data structure $f(\cdot, \cdot)$ satisfying:

$$|f(\mathbf{x}, \mathbf{y}) - \mathbf{x}^T \mathbf{A} \mathbf{y}| \leq \epsilon_s n \|\mathbf{x}\|_2 \|\mathbf{y}\|_2. \text{ for all input } \mathbf{x}, \mathbf{y} \in \mathbb{R}^n.$$

We show that for ϵ_s sufficiently small respect to ϵ_d , such a spectral approximation is sufficient to solve the (ϵ_d, n) -Distributed detection problem with no further queries to $\mathbf{A}$.

First, let B_i denote the number of ones in the i -th row of $\mathbf{A}$. Observe that $B_i \sim \text{Binomial}(n, 1/2)$ if $\mathbf{v}_i = 0$, and, $B_i \sim \text{Binomial}(n, 1/2 + \epsilon_d)$ if $\mathbf{v}_i = 1$. By Hoeffding's inequality [Ver18],

$$P\left(|B_i - E[B_i]| > \sqrt{4n \log n}\right) \leq 2 \exp\left(\frac{-4n \log n}{2n}\right) = \frac{2}{n^2}.$$

Therefore, by the union bound over all $i \in [n]$,

$$P\left(\max_{i \in S} |B_i - E[B_i]| > \sqrt{4n \log n}\right) \leq \frac{2}{n}.$$

Let $S_1, S_2 \subset [n]$, such that $|S_1| = |S_2| = \frac{n}{2}$. Then with probability at least $1 - \frac{2}{n}$,

$$\begin{aligned} \left| \sum_{i \in S_1} B_i - \sum_{i \in S_2} B_i \right| &\leq \left| E \left[\sum_{i \in S_1} B_i - \sum_{i \in S_2} B_i \right] \right| + \frac{2n}{2} \cdot \sqrt{4n \log n} \\ &= \left| E \left[\sum_{i \in S_1} B_i - \sum_{i \in S_2} B_i \right] \right| + O(n^{3/2} \log n). \end{aligned} \quad (14)$$

For large enough n , we have with probability at least $\frac{99}{100}$ that there are at least $\frac{n}{4}$ biased rows, since the number of biased rows is distributed as $\text{Binomial}(n, 1/2)$. Define the set S^* such that $|S^*| = \frac{n}{2}$ and S^* contains a maximal number of biased rows. Define the set $\bar{S}$ as,

$$\bar{S} = \arg \max_{|S| = \frac{n}{2}} f(\mathbf{1}, \mathbf{1}_S).$$

We will show that $\bar{S}$ can contain at most $k \leq \frac{n}{80} + o(n)$ fewer biased rows than S^* . By guarantee of the spectral approximation and (14), (note $\mathbf{1}_{\bar{S}}^T \mathbf{A} \mathbf{1}_S = \sum_{i,j \in S} \mathbf{A}_{ij}$),

$$\begin{aligned} f(\mathbf{1}, \mathbf{1}_{S^*}) - f(\mathbf{1}, \mathbf{1}_{\bar{S}}) &\geq \mathbf{1}^T \mathbf{A} \mathbf{1}_{S^*} - \mathbf{1}^T \mathbf{A} \mathbf{1}_{\bar{S}} - 2\epsilon_s n \|\mathbf{1}\|_2 \|\mathbf{1}_{S^*}\|_2 \\ &\geq E[\mathbf{1}^T \mathbf{A} \mathbf{1}_{S^*} - \mathbf{1}^T \mathbf{A} \mathbf{1}_{\bar{S}}] + O(n^{3/2} \log n) - 2\epsilon_s n^2. \end{aligned}$$

If $\bar{S}$ has k fewer biased rows than S^* , then, $E[\mathbf{1}^T \mathbf{A} \mathbf{1}_{S^*} - \mathbf{1}^T \mathbf{A} \mathbf{1}_{\bar{S}}] = \epsilon_d n k$. By the optimality of $\bar{S}$, $f(\mathbf{1}, \mathbf{1}_{S^*}) - f(\mathbf{1}, \mathbf{1}_{\bar{S}}) \leq 0$. Therefore, if $\epsilon_s = \frac{\epsilon_d}{160}$, then,

$$0 \geq f(\mathbf{1}, \mathbf{1}_{S^*}) - f(\mathbf{1}, \mathbf{1}_{\bar{S}}) \geq \epsilon_d n k - 2\epsilon_s n^2 + O(n^{3/2} \log n) = \epsilon_d n k - \frac{\epsilon_d n^2}{80} + O(n^{3/2} \log n).$$

Solving for k implies, $k \leq \frac{n}{80} + O(\frac{n^{1/2} \log n}{\epsilon_d})$. By assumption of the theorem statement, $\frac{1}{\epsilon_d} = o(\frac{\sqrt{n}}{\log n})$, therefore, $k \leq \frac{n}{80} + o(n)$.

Let b be the number of biased rows in $\mathbf{A}$. First, we consider the case where $b = \frac{n}{2}$ exactly. In this case, S^* contains $\frac{n}{2}$ biased rows, and hence, $\bar{S}$ contains at least $\frac{n}{2} - \frac{n}{80} - o(n) \geq \frac{n}{2} - \frac{n}{40}$ biased rows for large enough n . Therefore, deciding all rows in $\bar{S}$ (i.e., $\hat{\mathbf{v}}_i = 1$ for all $i \in \bar{S}$) are biased will correctly decide $(\frac{n}{2} - \frac{n}{40})/\frac{n}{2} \geq \frac{19}{20}$ of the biased rows. By looking at the complement of S^* and $\bar{S}$, we conclude that $\frac{19}{20}$ of the unbiased rows are decided correctly as well. Hence, we can construct $\hat{\mathbf{v}}$ such that $\|\mathbf{v} - \hat{\mathbf{v}}\|_1 \leq \frac{n}{20}$ by the assignment $\hat{\mathbf{v}}_i = 1$ for $i \in \bar{S}$ and $\hat{\mathbf{v}}_i = 0$ otherwise.

The number of biased rows is distributed as a binomial random variable, i.e., $b \sim \text{Binomial}(n, 1/2)$. Therefore, by Hoeffding's inequality, $\mathbb{P}(|b - \frac{n}{2}| > 10\sqrt{n}) < 0.01$. For large enough n , $10\sqrt{n} \leq 0.01 \cdot n$. Therefore, with probability at least $\frac{99}{100}$, we can reduce to the case $b = \frac{n}{2}$ by assuming that at most an additional $0.01 \cdot n$ rows are misclassified as biased or unbiased. Hence, with probability at least $\frac{99}{100}$, a spectral approximation of $\mathbf{A}$ with $\epsilon_s = \Theta(\epsilon_d)$ accuracy is sufficient to recover at least $\frac{9}{10}$ of the biased rows with probability at least $\frac{99}{100}$.

Correctly classifying $\frac{9}{10}$ of the rows as biased or unbiased requires observing $\Omega(\frac{n}{\epsilon_d^2})$ entries of $\mathbf{A}$ by Lemma 3. Since $\Omega(\frac{n}{\epsilon_d^2}) = \Omega(\frac{n}{\epsilon_s^2})$ entries of $\mathbf{A}$ must be observed to construct the spectral approximation used to solve the (ϵ, n) -Distributed detection problem, we conclude the lower bound of the theorem statement. $\square$

Note that the assumption $\frac{1}{\epsilon} = o(\frac{\sqrt{n}}{\log n})$ in the previous theorem is mild, since if this assumption does not hold, then we must read nearly all entries of the matrix anyways. We also show that our construction provides a lower bound even when the input is restricted to be symmetric.

Corollary 1. *The lower bound in Theorem 11 applies when the input is restricted to **symmetric** binary matrices.*

Proof. While construction used in Theorem 11 is not symmetric, we can modify it to give the same lower bound for **symmetric** input. Let $\mathbf{A}$ be defined as in the proof of Theorem 11, and let $\bar{\mathbf{A}}$ be the Hermitian dilation of $\mathbf{A}$, i.e.,

$$\bar{\mathbf{A}} = \begin{bmatrix} \mathbf{0} & \mathbf{A} \\ \mathbf{A}^T & \mathbf{0} \end{bmatrix}.$$

Any query $\mathbf{x}^T \mathbf{A} \mathbf{y}$ can be simulated by the query $\bar{\mathbf{x}}^T \bar{\mathbf{A}} \bar{\mathbf{y}}$, where $\bar{\mathbf{x}} = [\mathbf{x}, \mathbf{0}]^T$ and $\bar{\mathbf{y}} = [\mathbf{0}, \mathbf{y}]^T$. The size of $\mathbf{A}$ is n and the size of $\bar{\mathbf{A}}$ is $2n$, hence, if $\bar{f}$ is a spectral approximation of $\bar{\mathbf{A}}$ such that,

$$|\bar{f}(\bar{\mathbf{x}}, \bar{\mathbf{y}}) - \bar{\mathbf{x}}^T \bar{\mathbf{A}} \bar{\mathbf{y}}| \leq \left(\frac{\epsilon_s}{2}\right) (2n) \|\bar{\mathbf{x}}\|_2 \|\bar{\mathbf{y}}\|_2, \text{ for all } \bar{\mathbf{x}}, \bar{\mathbf{y}} \in \mathbb{R}^{2n},$$

then we can simulate $f(\cdot, \cdot)$ such that,

$$|f(\mathbf{x}, \mathbf{y}) - \mathbf{x}^T \mathbf{A} \mathbf{y}| \leq \epsilon_s n \|\mathbf{x}\|_2 \|\mathbf{y}\|_2, \text{ for all } \mathbf{x}, \mathbf{y} \in \mathbb{R}^n.$$

Therefore, the proof of Theorem 11 applies to the Hermitian dilation $\bar{\mathbf{A}}$ after adjusting for a constant factor in the accuracy parameter ϵ_s . $\square$

5 Improved Bounds for Binary Magnitude PSD Matrices

We call a PSD matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ a *binary magnitude PSD matrix* if $\forall i, j \in [n], |\mathbf{A}_{ij}| \in \{0, 1\}$, i.e., if $\mathbf{A} \in \{-1, 0, 1\}^{n \times n}$. In this section, we give deterministic spectral approximation algorithms for such matrices with improved ϵ dependencies. Specifically, we show that there exists a deterministic ϵn error spectral approximation algorithm reading just $O\left(\frac{n \log n}{\epsilon}\right)$ entries of the input matrix. This improves on our bound for general bounded entry PSD matrices (Theorem 1) by a $1/\epsilon$ factor, while losing a $\log n$ factor. Further, we show that it is tight up to a logarithmic factor, via a simple application of Turán’s theorem [Tur41].

The key idea of our improved algorithm is that for any PSD $\mathbf{A}$, we can write $\mathbf{A} = \mathbf{B}^T \mathbf{B}$, so that $\mathbf{A}_{ij} = \langle \mathbf{b}_i, \mathbf{b}_j \rangle$, where $\mathbf{b}_i, \mathbf{b}_j$ are the i^{th} and j^{th} columns of $\mathbf{B}$ respectively. Since $\mathbf{A}$ has bounded entries, for any i , $\mathbf{A}_{ii} = \|\mathbf{b}_i\|_2^2 \leq 1$. Thus, we can apply transitivity arguments: if $|\mathbf{A}_{ij}| = |\langle \mathbf{b}_i, \mathbf{b}_j \rangle|$ is large (close to 1) and $|\mathbf{A}_{jk}| = |\langle \mathbf{b}_j, \mathbf{b}_k \rangle|$ is also large, then $|\mathbf{A}_{ik}| = |\langle \mathbf{b}_i, \mathbf{b}_k \rangle|$ must be relatively large as well. Thus, we can hope to infer the contribution of an entry to $\mathbf{x}^T \mathbf{A} \mathbf{x}$ without necessarily reading that entry, allowing for reduced sample complexity.

The logic is particularly clean when $\mathbf{A}$ is binary. In this case, we can restrict to looking at the principal submatrix corresponding to $i \in [n]$ for which $\mathbf{A}_{ii} = 1$. The remainder of the matrix is all zeros. For such i , we have $\|\mathbf{b}_i\|_2 = 1$, and thus if $\mathbf{A}_{ij} = 1$ and $\mathbf{A}_{jk} = 1$, we can conclude that $\mathbf{A}_{ik} = 1$. This implies that, up to a permutation of the row/column indices, $\mathbf{A}$ is a block diagonal matrix, with $\text{rank}(\mathbf{A})$ blocks of all ones, corresponding to groups of indices S_k with $\mathbf{b}_i = \mathbf{b}_j$ for all $i, j \in S_k$. The eigenvalues of this matrix are exactly the sizes of these blocks, and to recover a spectral approximation, it suffices to recover the set S_k corresponding to any block of size $\geq \epsilon n$.

To do so, we employ a weak notion of expanders [Pip87, WZ93], that requires that any two subsets of vertices both of size $\geq \epsilon n$, are connected by at least one edge. Importantly, such expanders can be constructed with $\tilde{O}(n/\epsilon)$ edges – beating Ramanujan graphs, which would require $\Omega(n/\epsilon^2)$ edges, but satisfy much stronger notions of edge discrepancy. Further, if we consider the graph whose edges are in the intersection of those in the expander graph and of the non-zero entries in $\mathbf{A}$, we can show that any block of ones in $\mathbf{A}$ of size $\Omega(\epsilon n)$ will correspond to a large $\Omega(\epsilon n)$ sized connected component in this graph. We can form this graph by querying $\mathbf{A}$ at $\tilde{O}(n/\epsilon)$ positions (the edges in the expander) and then computing these large connected components to recover a spectral approximation. Each sparse connected component becomes a dense block of all ones in our approximation, avoiding the sparsity lower bound of Theorem 3. Noting that we can extend this logic to matrices with entries in $\{-1, 0, 1\}$, yields the main result of this section, Theorem 9.

It is not hard to see that $\tilde{O}(n/\epsilon)$ is optimal, even for binary PSD matrices (Theorem 10). By Turán’s theorem, any sample set of size $o(n/\epsilon)$ does not read any entries in some principal submatrix of size $\epsilon n \times \epsilon n$. By letting $\mathbf{A}$ be the identity plus a block of all ones on this submatrix, we force our algorithm to incur ϵn approximation error, since it cannot find this large block of ones.

Section Roadmap. We formally prove the block diagonal structure of binary magnitude PSD matrices in Section 5.1. We leverage this structure to give our improved spectral approximation algorithm for these matrices in Section 5.2. Finally, in Section 5.3 we prove a nearly matching lower bound via Turán’s theorem.

5.1 Structure of Binary Magnitude PSD Matrices

As discussed, our improved algorithm hinges on the fact that, up to a permutation of the rows and columns, binary magnitude PSD matrices are block diagonal. Further, each block is itself a rank-1 matrix, consisting of two on-diagonal blocks of 1's and two off-diagonal blocks of -1 's. Formally,

Lemma 4 (Binary Magnitude PSD Matrices are Block Matrices). *Let $\mathbf{A} \in \{-1, 0, 1\}^{n \times n}$ be PSD and let λ_i for $i \in [n]$ be the eigenvalues of $\mathbf{A}$, with $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p > 0$ and $\lambda_{p+1} = \dots = \lambda_n = 0$. Then, there exists a permutation matrix $\mathbf{P}$ such that*

$$\mathbf{PAP}^T = \begin{bmatrix} \mathbf{A}^{(1)} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{A}^{(2)} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \dots & \mathbf{A}^{(p)} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} \end{bmatrix}.$$

Further, Each $\mathbf{A}^{(i)} \in \{-1, 1\}^{\lambda_i \times \lambda_i}$ is a $\lambda_i \times \lambda_i$ rank-1 PSD matrix. Letting S_i denote the set of indices corresponding to the principal submatrix $\mathbf{A}^{(i)}$, which we call its support set, S_i can be partitioned in to two subsets S_{i1} and S_{i2} such that, if we let $\mathbf{v}_i(j) = 1$ for $j \in S_{i1}$ and $\mathbf{v}_i(j) = -1$ for $j \in S_{i2}$, then $\mathbf{v}_1, \dots, \mathbf{v}_p$ are orthogonal and $\mathbf{A} = \sum_{i=1}^p \mathbf{v}_i \mathbf{v}_i^T$.

Proof. Since $\mathbf{A}$ is PSD, there exists $\mathbf{B} \in \mathbb{R}^{n \times n}$ with rows $\mathbf{B}_1, \dots, \mathbf{B}_n \in \mathbb{R}^n$ such that $\mathbf{A} = \mathbf{B}\mathbf{B}^T$. Note that either $\mathbf{A}_{ii} = 0$ or $\mathbf{A}_{ii} = 1$ for all i since $\mathbf{A}_{ii} = \|\mathbf{B}_i\|_2^2 \geq 0$. Also for any $i, j \in [n]$, if $\mathbf{A}_{ij} = 1$ or $\mathbf{A}_{ij} = -1$, then $\mathbf{A}_{ii} = \mathbf{A}_{jj} = 1$. Otherwise, if $\mathbf{A}_{ii}$ or $\mathbf{A}_{jj}$ were 0, we would have $\|\mathbf{B}_i\|_2 = 0$ or $\|\mathbf{B}_j\|_2 = 0$ and thus $\mathbf{A}_{ij} = \mathbf{B}_i \mathbf{B}_j^T = 0$.

Thus, if $\mathbf{A}_{ij} = 1$, we have $\mathbf{A}_{ii} = \|\mathbf{B}_i\|_2 = 1$, $\mathbf{A}_{jj} = \|\mathbf{B}_j\|_2 = 1$, and $\mathbf{B}_i^T \mathbf{B}_j = 1$. I.e., we must have $\mathbf{B}_i = \mathbf{B}_j$. Similarly, if $\mathbf{A}_{ij} = -1$, we have $\|\mathbf{B}_i\|_2 = \|\mathbf{B}_j\|_2 = 1$ and $\mathbf{B}_i^T \mathbf{B}_j = -1$ and so $\mathbf{B}_i = -\mathbf{B}_j$. Finally, if $\mathbf{A}_{ij} = 0$, then $\mathbf{B}_i \perp \mathbf{B}_j$. This implies that the indices $[n]$ can be grouped into disjoint subsets $S_1, S_2, \dots, S_p$ ($S_i \subset [n]$) such that for any S_i and any $j, k \in S_i$, either $\mathbf{B}_i = \mathbf{B}_j$ or $\mathbf{B}_i = -\mathbf{B}_j$. Also for any $k \in S_i$ and $\ell \in S_j$ with $i \neq j$, we must have $\mathbf{B}_k \perp \mathbf{B}_\ell$. Label these subsets such that $|S_1| \geq |S_2| \geq \dots \geq |S_p|$. Now observe that any subset S_i can be further divided into two disjoint subsets S_{i1} and S_{i2} such that if $j, k \in S_{i1}$ or $j, k \in S_{i2}$, $\mathbf{B}_j = \mathbf{B}_k$ but if $j \in S_{i1}$ and $k \in S_{i2}$ (or vice-versa), $\mathbf{B}_j = -\mathbf{B}_k$. Thus, the principal submatrix corresponding to S_i , which we label $\mathbf{A}^{(i)}$ is a rank-1 block matrix with two blocks of all ones on the principal submatrices corresponding to S_{i1} and S_{i2} and -1 's on all remaining entries.

For any S_i , let $\mathbf{v}_i \in \{0, 1\}^n$ be a vector such that $\mathbf{v}_i(j) = 1$ if $j \in S_{i1}$, $\mathbf{v}_i(j) = -1$ if $j \in S_{i2}$ and $\mathbf{v}_i(j) = 0$ otherwise. Then, observe that $\mathbf{A}\mathbf{v}_i = |S_i|\mathbf{v}_i$. Thus, each $|S_i|$ is an eigenvalue of $\mathbf{A}$ and each $\mathbf{v}_i$ is an eigenvector. Since $S_1, \dots, S_p$ has disjoint supports, $\mathbf{v}_1, \dots, \mathbf{v}_p$ are orthogonal. Further, $\text{tr}(\mathbf{A}) = \sum_{i=1}^n \lambda_i = \sum_{i=1}^p |S_i|$ and thus, all other eigenvalues of $\mathbf{A}$ are 0. Thus, we can write $\mathbf{A} = \sum_{i=1}^p \mathbf{v}_i \mathbf{v}_i^T$. This concludes the lemma. $\square$

5.2 Improved Spectral Approximation of Binary Magnitude PSD matrices

Given Lemma 4, the key idea to efficiently recovering a spectral approximation to a binary magnitude PSD matrix $\mathbf{A}$ is that we do not need to read very many entries in each block $\mathbf{A}^{(i)}$ to identify the block. We just need enough to identify a large fraction of the indices in the support set S_i .

To ensure that our sampling is able to do so, we will sample according to the edges of a certain type of weak expander graph, which ensures that any two vertex sets of large enough size are connected. I.e., that we read at least one entry in any large enough off-diagonal block of our matrix.

Definition 5.1 (ϵn -expander graphs [WZ93, Pip87]). *For any $\epsilon \in \mathbb{R}$, an ϵn -expander graph is any undirected graph on n vertices such that any two disjoint subsets of vertices containing at least ϵn vertices each are joined by an edge.*

Moreover, it is not hard to show that ϵn -expanders with just $O(n \log n / \epsilon)$ edges exist.

Fact 1. [WZ93] *A random $d = O\left(\frac{\log n}{\epsilon}\right)$ regular graph is an ϵn -expander with high probability.*

We note that while a spectral expander graph, such as a Ramanujan graph, as used to achieve Theorem 1, will also be an ϵn -expander by the expander mixing lemma [CG02], such graphs require $O(n/\epsilon^2)$ edges, as opposed to the $\tilde{O}(n/\epsilon)$ given by Fact 1. Further, while Fact 1 is not constructive, in [WZ93] an explicit polynomial time construction of ϵn -expander graphs is shown, with just an $n^{o(1)}$ loss. This result can be plugged directly into our algorithms to give a polynomial time constructible deterministic spectral approximation algorithm with $O(n^{1+o(1)}/\epsilon)$ sample complexity.

Fact 2 (Constructive ϵn -expanders [WZ93]). *There is a polynomial time algorithm that, given an integer n and $\epsilon \in (0, 1)$, constructs an ϵn -expanding graph on n vertices with maximum degree $\frac{n^{o(1)}}{\epsilon}$.*

We now prove our main technical result, which shows that if we read the entries of a binary magnitude PSD matrix $\mathbf{A}$ according to an ϵn -expander graph, then we can approximately recover all blocks of this matrix by considering the *connected components* of the expander graph restricted to edges where $\mathbf{A}$ is nonzero.

Theorem 13 (Approximation via Sampled Connected Components). *Let $\mathbf{A} \in \{-1, 0, 1\}^{n \times n}$ be a PSD matrix with eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p > 0$ and $\lambda_{p+1} = \dots = \lambda_n = 0$. Let G be an $\epsilon/6 \cdot n$ -expanding graph on n vertices with adjacency matrix $\mathbf{B}_G$. Let $\mathbf{A}_G$ be the binary adjacency matrix with $(\mathbf{A}_G)_{ij} = 1$ whenever $(\mathbf{B}_G)_{ij} = 1$ and $|\mathbf{A}_{ij}| = 1$. Let $\bar{G}$ be the graph whose adjacency matrix is $\mathbf{A}_G$. Let $S_1, \dots, S_p$ denote the support sets of $\mathbf{A}$'s blocks (as defined in Lemma 4) with $|S_i| = \lambda_i$, and let $\bar{G}_{S_i}$ be the induced subgraph of $\bar{G}$ restricted to S_i . Finally, let $C_i \subset [n]$ be the largest connected component of $\bar{G}_{S_i}$. We have:*

$$\lambda_i - \epsilon/2 \cdot n < |C_i| \leq \lambda_i.$$

Proof. First observe that since $|S_i| = \lambda_i$, we trivially have that $|C_i| \leq \lambda_i$. Thus, it remains to show that $|C_i| > \lambda_i - \epsilon/2 \cdot n$. This holds vacuously if $\lambda_i \leq \epsilon/2 \cdot n$. Thus, we focus our attention to $\lambda_i \geq \epsilon/2 \cdot n$. For contradiction assume $|C_i| \leq \lambda_i - \epsilon/2 \cdot n$. Let $M_1, \dots, M_r \subseteq S_i$ denote the connected components of $\bar{G}_{S_i}$ such that $|M_1| \geq \dots \geq |M_r|$. So $|C_i| = |M_1|$. Observe that since all entries in $\mathbf{A}$ in the principal submatrix indexed by S_i (i.e., in $\mathbf{A}^{(i)}$ from Lemma 4) have magnitude 1, $\bar{G}_{S_i} = G_{S_i}$. Consider two cases.

Case 1: $|C_i| \geq \epsilon/6 \cdot n$. From our assumption: $\lambda_i - |C_i| \geq \epsilon/2 \cdot n$ which implies that $\sum_{j=2}^r |M_j| \geq \epsilon/2 \cdot n$. Let $C' = \bigcup_{j=2}^r M_j$ so $|C'| \geq \epsilon/2 \cdot n$. Since $M_1, \dots, M_r$ are disconnected, there are no edges from C' to C_i in $\bar{G}$, and thus no edges in G , since $\bar{G}_{S_i} = G_{S_i}$. However, this contradicts the fact that G is an $\epsilon/6 \cdot n$ expander and thus must have at least one edge between any two set of vertices of size at least $\epsilon/6 \cdot n$ (Definition 1). Thus we have a contradiction and our assumption $\lambda_i - |C_i| \geq \epsilon n$ is incorrect. So we must have $|C_i| > \lambda_i - \epsilon n$ as needed.

Case 2: $|C_i| < \epsilon/6 \cdot n$. Since $|C_i| < \epsilon/6 \cdot n$, we have $\max_{i \in [r]} |M_i| < \epsilon/6 \cdot n$. Consider partitioning these connected components into two sets, T_1 and T_2 . Let V_1, V_2 be the vertex sets corresponding to these two partitions, i.e., $V_1 = \bigcup_{M_i \in T_1} M_i$ and $V_2 = \bigcup_{M_i \in T_2} M_i$. Pick T_1 and T_2 to minimize $||V_1| - |V_2||$. Since $\max_{i \in [r]} |M_i| < \epsilon/6 \cdot n$, for this optimal T_1, T_2 , we must have $||V_1| - |V_2|| < \epsilon/6 \cdot n$.

Since V_1 and V_2 are disconnected in $\bar{G}_{S_i}$, they must be disconnected in G_{S_i} . Thus, we must have either $|V_1| < \epsilon/6 \cdot n$ or $|V_2| < \epsilon/6 \cdot n$ since any two subsets of vertices with size $\geq \epsilon/6 \cdot n$ must have an edge between them in G due to it being an $\epsilon/6 \cdot n$ -expander. Assume w.l.o.g. that $|V_2| < \epsilon/6 \cdot n$. Then since $|V_1| + |V_2| = |S_i| = \lambda_i \geq \epsilon/2 \cdot n$, we must have $|V_1| > \epsilon/3 \cdot n$. However, this means that $||V_1| - |V_2|| > \epsilon/6 \cdot n$, which contradicts the fact that we picked T_1, T_2 such that $||V_1| - |V_2|| < \epsilon/6 \cdot n$. $\square$

We now present our deterministic algorithm for spectral approximation of binary magnitude PSD matrices, based on Theorem 13. The key idea is to compute all connected components of $\bar{G}$ – the graph whose edges lie in the intersection of the edges of G and the non-zero entries of $\mathbf{A}$. Any large enough block in $\mathbf{A}$ (see Lemma 4) will correspond to a large connected component in this graph and can thus be identified.

Algorithm 1 Deterministic spectral approximator using expanders

Input: PSD matrix $\mathbf{A} \in \{-1, 0, 1\}^{n \times n}$ and $\epsilon \in (0, 1)$.

- 1: Construct an $\epsilon/6 \cdot n$ -expanding graph G on n nodes with adjacency matrix $\mathbf{B}_G$.
 - 2: Let $\mathbf{A}_G$ be an $n \times n$ matrix such that $(\mathbf{A}_G)_{ij} = 1$ if $(\mathbf{B}_G)_{ij} = 1$ and $|\mathbf{A}_{ij}| = 1$. Let $\bar{G}$ be the graph corresponding to $\mathbf{A}_G$.
 - 3: Compute all connected components of $\bar{G}$, $C_1, C_2, \dots, C_r \subseteq [n]$.
 - 4: Initialize $\tilde{\mathbf{A}} = \mathbf{0}^{n \times n}$.
 - 5: **for** $i = 1, 2, \dots, r$ **do**
 - 6: **if** $|C_i| > \frac{\epsilon n}{2}$ **then**
 - 7: Pick any $j \in C_i$ and let $\tilde{\mathbf{A}} = \tilde{\mathbf{A}} + \mathbf{A}_j \mathbf{A}_j^T$, where $\mathbf{A}_j$ is the j^{th} column of $\mathbf{A}$.
 - 8: **end if**
 - 9: **end for**
 - 10: **return** $\tilde{\mathbf{A}}$.
-

Query complexity: Observe that Algorithm 1 requires querying $\mathbf{A}$ at the locations where $\mathbf{B}_G$ is non-zero (to perform step 2). By Fact 1, there are $\epsilon/6 \cdot n$ -expander graphs with just $O(\frac{n \log n}{\epsilon})$ edges, so this requires $O(\frac{n \log n}{\epsilon})$ queries. Further, in a graph with n nodes, there can be at most $2/\epsilon$ connected components of size $\geq \frac{\epsilon n}{2}$. Thus, line (7) requires total query complexity $\leq \frac{2n}{\epsilon}$ to read one column of $\mathbf{A}$ corresponding to each large component identified.

Theorem 9 (Deterministic Spectral Approximation of Binary Magnitude PSD Matrices). *Let $\mathbf{A} \in \{-1, 0, 1\}^{n \times n}$ be PSD. Then for any $\epsilon \in (0, 1)$, there exists a deterministic algorithm which reads $O(\frac{n \log n}{\epsilon})$ entries of $\mathbf{A}$ and returns PSD $\tilde{\mathbf{A}} \in \{-1, 0, 1\}^{n \times n}$ such that $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon n$.*

Proof. We will show that Algorithm 1 satisfies the requirements of the theorem. We have already argued above that its query complexity is $O(\frac{n \log n}{\epsilon})$.

By Lemma 4, we know that, up to a permutation, $\mathbf{A}$ is a block matrix with rank-1 blocks $\mathbf{A}^{(1)}, \mathbf{A}^{(2)}, \dots, \mathbf{A}^{(k)}$ with support sets $S_1, \dots, S_k$. Since $\bar{G}$ has edges only where $\mathbf{A}$ is non-zero, each connected component C_i in $\bar{G}$ is entirely contained within one of these support sets. Further, if $|S_i| = \lambda_i \geq \epsilon n$, then by Theorem 13, there is exactly one connected component, call it D_i (the largest connected component of $\bar{G}_{S_i}$) with vertices in S_i and $|D_i| \geq \epsilon/2$. If $|S_i| = \lambda_i < \epsilon n$, then there is at most one such connected component by Theorem 13 (there may be none).

So, Step 7 is triggered for exactly one connected component $D_i \subseteq S_i$ for each S_i with $|S_i| = \lambda_i \geq \epsilon n$ and at most one connected component for all other S_i . If Line 7 is triggered, then we add $\mathbf{A}_j \mathbf{A}_j^T$ to $\tilde{\mathbf{A}}$. Observe that $\mathbf{A}_j$ is a column for $j \in S_i$, and thus by Lemma 4, $\mathbf{A}_j$ is supported only

on $\mathbf{S}_i$. We have that either $\mathbf{A}_j(t) = 1$ on S_{i1} and $\mathbf{A}_j(t) = -1$ on S_{i2} , or $\mathbf{A}_j(t) = -1$ on S_{i1} and $\mathbf{A}_j(t) = 1$. That is, $\mathbf{A}_j$ is equal to either $\mathbf{v}_i$ or $-\mathbf{v}_i$ as defined in Lemma 4.

So overall, letting $Z \subseteq \{S_1, \dots, S_k\}$ be the set of blocks for which one connected component is recovered in Line 7, $\tilde{\mathbf{A}} = \sum_{i \in Z} \mathbf{v}_i \mathbf{v}_i^T$. From Lemma 4, $\mathbf{A} = \sum_{i=1}^k \mathbf{v}_i \mathbf{v}_i^T$ and so $\mathbf{A} - \tilde{\mathbf{A}} = \sum_{i \notin Z} \mathbf{v}_i \mathbf{v}_i^T$. Since the $\mathbf{v}_i$ are orthogonal and since for $i \notin Z$ we must have $\lambda_i = |S_i| = \|\mathbf{v}_i \mathbf{v}_i^T\|_2 \leq \epsilon n$, this gives that $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon n$, completing the theorem. $\square$

5.3 Lower Bounds for Binary PSD Matrix Approximation

We can prove that our $O\left(\frac{n \log n}{\epsilon}\right)$ query complexity for binary magnitude matrices is optimal up to a log n factor via Turán's Theorem, stated below. Our lower bound holds for the easier problem of eigenvalue approximation for PSD $\mathbf{A} \in \{0, 1\}^{n \times n}$.

Fact 3 (Turán's theorem [Tur41, Aig95]). *Let G be a graph on n vertices that does not include a $(k+1)$ -clique as a subgraph. Let $t(n, k)$ be the number of edges in G . Then*

$$t(n, k) \leq \frac{(k-1)n^2}{2k}.$$

Using Turán's theorem we can prove:

Theorem 10 (Deterministic Spectral Approximation of Binary PSD Matrices – Lower Bound). *Let $\mathbf{A} \in \{0, 1\}^{n \times n}$ be PSD. Then for any $\epsilon \in (0, 1)$, any possibly adaptive deterministic algorithm which approximates all eigenvalues of $\mathbf{A}$ up to ϵn additive error must read $\Omega\left(\frac{n}{\epsilon}\right)$ entries of $\mathbf{A}$.*

Proof. Let $\mathcal{A}$ be a deterministic algorithm that approximates the eigenvalues of a binary PSD input matrix $\mathbf{A}$ to ϵn additive error. Let $\mathbf{A}_0 = \mathbf{I}_n$, i.e., the identity matrix. Let $S \subset [n] \times [n]$ be the first $\frac{n}{\epsilon}$ off-diagonal entries read by $\mathcal{A}$ on input $\mathbf{A}$. Since the choices of the algorithm are deterministic, these same entries will be the first $\frac{n}{\epsilon}$ entries queried by $\mathcal{A}$ from any input matrix which has ones on the diagonal and zeroes in all entries of S , even if the algorithm is adaptive.

Without loss of generality, we can assume that the entries in S are above the diagonal, since reading any entry below the diagonal of an input matrix would be redundant due to symmetry. Consider a graph G with vertex set $V = [n]$ and undirected edge set $E = S^c$, that is, the edge set contains all undirected edges (i, j) such that $(i, j), (j, i) \notin S$ for $i, j \in [n]$. Hence, for all $\epsilon \in (0, 1)$,

$$|E| = \binom{n}{2} - \frac{n}{\epsilon} = \frac{n^2}{2} - \left(\frac{n}{2} + \frac{n}{\epsilon}\right) > \frac{n^2}{2} - \frac{2n}{\epsilon} = \frac{(\frac{\epsilon n}{4} - 1)n^2}{2 \cdot \frac{\epsilon n}{4}},$$

where the first inequality is true since $\frac{n}{2} < \frac{n}{\epsilon}$. Therefore, if we set $k = \frac{\epsilon n}{4}$, we conclude by the contrapositive of Fact 3 that there exists a clique of size $\frac{\epsilon n}{4}$ in G . Let $T \subset E$ denote the set of edges which forms this clique in G .

Construct $\mathbf{A}_1$ such that $[\mathbf{A}_1]_{ij} = [\mathbf{A}_0]_{ij}$ for all i, j such that $(i, j) \notin T$ and $[\mathbf{A}_1]_{ij} = [\mathbf{A}_1]_{ji} = 1$ for all i, j such that $(i, j) \in T$. Then, there is a principal submatrix containing all ones in $\mathbf{A}_1$ where the off-diagonal entries correspond to edges in T along with the diagonal entries which are all ones by construction. Hence, $\mathbf{A}_1$ contains a non-zero principal submatrix of size $\frac{\epsilon n}{4}$, and so, it has an eigenvalue that is at least $\frac{\epsilon n}{4}$. Both $\mathbf{A}_0$ and $\mathbf{A}_1$ have identical values on the diagonal and on the first $\frac{n}{\epsilon}$ off-diagonal entries read by $\mathcal{A}$ (i.e., the entries in S). Therefore, the $\frac{n}{\epsilon}$ entries of $\mathbf{A}_0$ and $\mathbf{A}_1$ queried by $\mathcal{A}$ are identical, and hence $\mathcal{A}$ cannot distinguish these two matrices, despite the fact their maximum eigenvalue differs by $\frac{\epsilon n}{4}$.

After adjusting for constant factors, we conclude that any deterministic (possibly adaptive) algorithm that estimates the eigenvalues of a $\{0, 1\}$ -matrix to ϵn additive error must observe $\Omega(\frac{n}{\epsilon})$ entries of the input matrix in the worst case. $\square$

6 Applications to Fast Deterministic Algorithms for Linear Algebra

We finally show how to leverage our universal sparsifiers to design the first $o(n^\omega)$ time deterministic algorithms for several problems related to singular value and vector approximation. As discussed in Section 1.1.1, throughout this section, runtimes are stated assuming that we have already constructed a deterministic sampling matrix $\mathbf{S}$ satisfying the universal sparsification guarantees of Theorem 1 or Theorem 4. When $\mathbf{S}$ is the adjacency matrix of a Ramanujan graph, this can be done in $\tilde{O}(n/\text{poly}(\epsilon))$ time – see Section 1.1.1 and Section 2 for further discussion.

In Section 6.1, we consider approximating the singular values of $\mathbf{A}$ to additive error $\pm\epsilon \max(n, \|\mathbf{A}\|_1)$. Observe that if we deterministically compute $\tilde{\mathbf{A}} = \mathbf{A} \circ \mathbf{S}$ with $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon \max(n, \|\mathbf{A}\|_1)$ via Theorem 4, by Weyl's inequality, all singular values of $\tilde{\mathbf{A}}$ approximate those of $\mathbf{A}$ up to additive error $\pm\epsilon \max(n, \|\mathbf{A}\|_1)$. Thus we can focus on approximating $\tilde{\mathbf{A}}$'s singular values.

To do so efficiently, we will use the fact that $\tilde{\mathbf{A}}$ is sparse and so admits very fast $\tilde{O}(\frac{n}{\epsilon^4})$ time matrix-vector products. Focusing on the top singular value, it is well known that the power method, initialized with a vector that has non-negligible (i.e., $\Omega(1/\text{poly}(n))$) inner product with the top singular vector $\mathbf{v}_1$ of $\tilde{\mathbf{A}}$, converges to a $(1 \pm \epsilon)$ relative error approximation in $O(\log n/\epsilon)$ iterations. I.e., the method outputs $\tilde{\mathbf{v}}_1$ with $(1 - \epsilon)\sigma_1(\tilde{\mathbf{A}}) \leq \|\tilde{\mathbf{A}}\tilde{\mathbf{v}}_1\|_2 \leq \sigma_1(\tilde{\mathbf{A}})$. Each iteration requires a single matrix-vector product with $\tilde{\mathbf{A}}$, yielding total runtime $\tilde{O}(\frac{n}{\epsilon^5})$. Since $\mathbf{v}_1$ is unknown, one typically initializes the power method with a random vector, which has non-negligible inner product with $\mathbf{v}_1$ with high probability. To derandomize this approach, we simply try n orthogonal starting vectors, e.g., the standard basis vectors, of which at least one must have large inner product with $\mathbf{v}_1$. This simple brute force approach yields total runtime $\tilde{O}(\frac{n^2}{\epsilon^5})$, which, for large enough ϵ , is $o(n^\omega)$.

To approximate more than just the top singular value, one would typically use a block power or Krylov method. However, derandomization is difficult here – unlike in the single vector case, it is not clear how to pick a small set of deterministic starting blocks for which at least one gives a good approximation. Thus, we instead use *deflation* [Saa11, AZL16]. We first use our deterministic power method to compute $\tilde{\mathbf{v}}_1$ with $\|\tilde{\mathbf{A}}\tilde{\mathbf{v}}_1\|_2 \approx \sigma_1(\tilde{\mathbf{A}})$. We then run the method again on the deflated matrix $\tilde{\mathbf{A}}(\mathbf{I} - \tilde{\mathbf{v}}_1\tilde{\mathbf{v}}_1^T)$, which we can still multiply by in $O(\frac{n}{\epsilon^4})$ time. This second run outputs $\tilde{\mathbf{v}}_2$ which is orthogonal to $\tilde{\mathbf{v}}_1$ and which has $\|\tilde{\mathbf{A}}\tilde{\mathbf{v}}_2\|_2 \approx \sigma_1(\tilde{\mathbf{A}}(\mathbf{I} - \tilde{\mathbf{v}}_1\tilde{\mathbf{v}}_1^T))$. Since $\tilde{\mathbf{v}}_1$ gives a good approximation to $\sigma_1(\tilde{\mathbf{A}})$, we then argue that $\sigma_1(\tilde{\mathbf{A}}(\mathbf{I} - \tilde{\mathbf{v}}_1\tilde{\mathbf{v}}_1^T)) \approx \sigma_2(\tilde{\mathbf{A}})$ and so $\|\tilde{\mathbf{A}}\tilde{\mathbf{v}}_2\|_2 \approx \sigma_2(\tilde{\mathbf{A}})$. We repeat this process to approximate all large singular values of $\tilde{\mathbf{A}}$. The key challenge is bounding the accumulation of error that occurs at each step.

In Section 6.2 we apply the above derandomized power method approach to the problem of testing whether $\mathbf{A}$ is PSD or has at least one large negative eigenvalue $< -\epsilon \max(n, \|\mathbf{A}\|_1)$. We reduce this problem to approximating the top singular value of a shifted version of $\mathbf{A}$.

Finally, in Section 6.3, we consider approximating $\sigma_1(\mathbf{A})$ to high accuracy – e.g., to relative error $(1 \pm \epsilon)$ with $\epsilon = 1/\text{poly}(n)$, under the assumption that $\sigma_1(\mathbf{A}) \geq \alpha \max(n, \|\mathbf{A}\|_1)$ for some α . Here, one cannot simply approximate $\sigma_1(\tilde{\mathbf{A}})$ in place of $\sigma_1(\mathbf{A})$, since the error due to sparsification is too large. Instead, we use our deflation approach to compute a ‘coarse’ approximate set of top singular vectors for $\tilde{\mathbf{A}}$. We then use these singular vectors to initialize a block Krylov method run on $\mathbf{A}$ itself, which we argue converges in $O(\frac{\log(n/\epsilon)}{\text{poly}(\alpha)})$ iterations, giving total run time $\tilde{O}(\frac{n^2 \log(1/\epsilon)}{\text{poly}(\alpha)})$.

The key idea in showing convergence is to argue that since $\sigma_1(\mathbf{A}) \geq \alpha \max(n, \|\mathbf{A}\|_1)$ by assumption, we have $\sigma_{2/\alpha}(\mathbf{A}) \leq \alpha/2 \|\mathbf{A}\|_1 \leq \sigma_1(\mathbf{A})/2$. So, there must be at least some singular value $\sigma_p(\mathbf{A})$ with $p \leq 2/\alpha$ that is larger than $\sigma_{p+1}(\mathbf{A})$ by $\alpha^2/4 \cdot \max(n, \|\mathbf{A}\|_1)$. The presence of this spectral gap allows us to (1) argue that the first p approximate singular vectors output by our deflation approach applied to $\tilde{\mathbf{A}}$ have non-negligible inner product with the true top k singular vectors of $\mathbf{A}$ and (2) that the block Krylov method initialized with these vectors converges rapidly to a high accuracy approximation to $\sigma_1(\mathbf{A})$, via known gap-dependent convergence bounds [MM15].

6.1 Deterministic Singular Value Approximation

We start by giving a deterministic algorithm to approximate all singular values of a symmetric bounded entry matrix $\mathbf{A}$ up to error $\pm\epsilon \max(n, \|\mathbf{A}\|_1)$ in $\tilde{O}(n^2/\text{poly}(\epsilon))$ time. Our main result appears as Theorem 8.

As discussed, we first deterministically compute a sparsified matrix $\tilde{\mathbf{A}}$ from $\mathbf{A}$ with $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon \max(n, \|\mathbf{A}\|_1)$. We then use a simple brute-force derandomization of the classical power method to approximate to the top singular value of $\tilde{\mathbf{A}}$ (Lemma 6). To approximate the rest of $\tilde{\mathbf{A}}$'s singular values, we repeatedly deflate the matrix and compute the top singular value of the deflated matrix (Lemma 7). By Weyl's inequality, the singular values of $\tilde{\mathbf{A}}$ approximate those of $\mathbf{A}$ up to $\pm\epsilon \max(n, \|\mathbf{A}\|_1)$ error. Thus, we can argue that our approximations will approximate to the singular values of $\mathbf{A}$ itself, yielding Theorem 8. We first present the derandomized power method that will be applied to $\tilde{\mathbf{A}}$ in Algorithm 2.

Algorithm 2 Deterministic Power Method for Largest Singular Value Estimation

Input: $\mathbf{A} \in \mathbb{R}^{n \times n}$, error parameter $\epsilon \in (0, 1)$

- 1: Initialize $\tilde{\sigma} = 0$, $t = \frac{c \log(n/\epsilon)}{\epsilon}$ where c is a sufficiently large constant.
 - 2: **for** $k = 1 \rightarrow n$ **do**
 - 3: $\mathbf{y}_k = \mathbf{e}_k$ where $\mathbf{e}_k$ is the k^{th} standard basis vector.
 - 4: **for** $j = 1 \rightarrow t$ **do**
 - 5: $\mathbf{y}_k \leftarrow (\mathbf{A}\mathbf{A}^T)\mathbf{y}_k$.
 - 6: $\mathbf{y}_k \leftarrow \frac{\mathbf{y}_k}{\|\mathbf{y}_k\|_2}$.
 - 7: **end for**
 - 8: **if** $\|\mathbf{A}^T \mathbf{y}_k\|_2 \geq \tilde{\sigma}$ **then**
 - 9: $\mathbf{z} \leftarrow \mathbf{y}_k$.
 - 10: $\tilde{\sigma} \leftarrow \|\mathbf{A}^T \mathbf{z}\|_2$.
 - 11: **end if**
 - 12: **end for**
 - 13: **return** $\mathbf{z}$
-

Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a matrix with its SVD given by $\mathbf{A} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T$ where $\mathbf{U}, \mathbf{V} \in \mathbb{R}^{n \times n}$ are matrices with orthonormal columns containing the left and right singular vectors respectively and $\mathbf{\Sigma}$ is a diagonal matrix containing the singular values $\sigma_1(\mathbf{A}) \geq \sigma_2(\mathbf{A}) \geq \dots \geq \sigma_n(\mathbf{A}) \geq 0$. Observe that any vector $\mathbf{y} \in \mathbb{R}^n$ can be written as $\mathbf{y} = \sum_{i=1}^n c_i \mathbf{u}_i$ where c_i are scalar coefficients and $\mathbf{u}_i$ are the columns of $\mathbf{U}$. Then, we have the following well-known result which shows that when the starting unit vector $\mathbf{y}$ has a large enough inner product with $\mathbf{u}_1$ i.e., $|c_1| \geq \frac{1}{\sqrt{n}}$, power iterations on $\mathbf{A}\mathbf{A}^T$ converge to a $1 - \epsilon$ approximation to $\sigma_1(\mathbf{A})$ within $\tilde{O}(1/\epsilon)$ iterations. We will leverage this lemma to prove the correctness and runtime of our deterministic power method, Algorithm 2.

Lemma 5 (Power iterations – gap independent bound). *Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a matrix with largest singular value $\sigma_1(\mathbf{A})$ and left singular vectors $\mathbf{u}_1, \dots, \mathbf{u}_n$. Let $\mathbf{y} = \sum_{i=1}^n c_i \mathbf{u}_i$ be a unit vector where $|c_1| \geq \frac{1}{\sqrt{n}}$. Then, for $t = O\left(\frac{\log(n/\epsilon)}{\epsilon}\right)$, if we set $\mathbf{z} = \frac{(\mathbf{A}\mathbf{A}^T)^t \mathbf{y}}{\|(\mathbf{A}\mathbf{A}^T)^t \mathbf{y}\|_2}$,*

$$\|\mathbf{A}^T \mathbf{z}\|_2 \geq (1 - \epsilon) \sigma_1(\mathbf{A}).$$

Proof. We prove Lemma 5 for completeness. Let $\mathbf{r} = (\mathbf{A}\mathbf{A}^T)^t \mathbf{y}$. Then, $\mathbf{z} = \frac{\mathbf{r}}{\|\mathbf{r}\|_2}$. First observe that $\mathbf{A}\mathbf{A}^T = \sum_{i=1}^n \sigma_i^2(\mathbf{A}) \mathbf{u}_i \mathbf{u}_i^T$. Thus, $\mathbf{r} = \sum_{i=1}^n c_i (\sigma_i(\mathbf{A}))^{2t} \mathbf{u}_i$. Let $k \in [n]$ be the smallest k such that $\sigma_k^2(\mathbf{A}) \leq (1 - \epsilon) \sigma_1^2(\mathbf{A})$. Observe that if no such k exists, then for any unit vector $\mathbf{z}$, $\|\mathbf{A}^T \mathbf{z}\|_2 \geq \sigma_n(\mathbf{A}) > (1 - \epsilon) \sigma_1(\mathbf{A})$, and thus the lemma holds trivially.

Let $\mathbf{r} = \mathbf{r}_1 + \mathbf{r}_2$ where $\mathbf{r}_1 = \sum_{i=1}^{k-1} c_i \sigma_i^{2t}(\mathbf{A}) \mathbf{u}_i$ and $\mathbf{r}_2 = \sum_{i=k}^n c_i \sigma_i^{2t}(\mathbf{A}) \mathbf{u}_i$. Then:

$$\frac{\|\mathbf{r}_2\|_2^2}{\|\mathbf{r}_1\|_2^2} = \frac{\sum_{i=k}^n c_i^2 \sigma_i^{4t}(\mathbf{A})}{\sum_{i=1}^{k-1} c_i^2 \sigma_i^{4t}(\mathbf{A})} \leq \frac{\sum_{i=k}^n c_i^2 \sigma_i^{4t}(\mathbf{A})}{c_1^2 \sigma_1^{4t}(\mathbf{A})} \leq \sum_{i=k}^n \left(\frac{c_i}{c_1}\right)^2 \left(\frac{\sigma_i(\mathbf{A})}{\sigma_1(\mathbf{A})}\right)^{4t}.$$

Since $c_i \leq 1$ for all $i \in [n]$, and since by assumption $|c_1| \geq \frac{1}{\sqrt{n}}$, $\left|\frac{c_i}{c_1}\right| \leq \sqrt{n}$ for all $i \in [n]$. Also, for any $i \geq k$, $\left(\frac{\sigma_i(\mathbf{A})}{\sigma_1(\mathbf{A})}\right)^2 \leq (1 - \epsilon)$ and $(1 - \epsilon)^{2t} \leq \delta$ for $t \geq O\left(\frac{\log(1/\delta)}{\epsilon}\right)$. Thus, we have $\frac{\|\mathbf{r}_2\|_2^2}{\|\mathbf{r}_1\|_2^2} \leq n^2 \delta$. Setting $\delta = \frac{\epsilon^2}{n^2}$, we get $\frac{\|\mathbf{r}_2\|_2^2}{\|\mathbf{r}_1\|_2^2} \leq \epsilon^2$.

Next, observe that $\|\mathbf{r}\|_2^2 = \|\mathbf{r}_1\|_2^2 + \|\mathbf{r}_2\|_2^2$ which implies that $\|\mathbf{r}\|_2^2 \leq (1 + \epsilon^2) \|\mathbf{r}_1\|_2^2$ or $\|\mathbf{r}_1\|_2^2 \geq (1 - \epsilon) \|\mathbf{r}\|_2^2$. Thus, using the fact that $\mathbf{r}_1^T \mathbf{A} \mathbf{A}^T \mathbf{r}_2 = 0$ and $\mathbf{r}_2^T \mathbf{A} \mathbf{A}^T \mathbf{r}_2 \geq 0$, and the fact that $\sigma_i^2(\mathbf{A}) \geq (1 - \epsilon) \sigma_1^2(\mathbf{A})$ for $i < k$, we get:

$$\begin{aligned} \mathbf{r}^T \mathbf{A} \mathbf{A}^T \mathbf{r} &= \mathbf{r}_1^T \mathbf{A} \mathbf{A}^T \mathbf{r}_1 + \mathbf{r}_2^T \mathbf{A} \mathbf{A}^T \mathbf{r}_2 \geq \mathbf{r}_1^T \mathbf{A} \mathbf{A}^T \mathbf{r}_1 \\ &\geq \sigma_{k-1}^2(\mathbf{A}) \cdot \|\mathbf{r}_1\|_2^2 \\ &\geq (1 - \epsilon)^2 \cdot \sigma_1^2(\mathbf{A}) \cdot \|\mathbf{r}\|_2^2. \end{aligned}$$

Finally dividing both sides of the above equation by $\|\mathbf{r}\|_2^2$, we get $\mathbf{z}^T \mathbf{A} \mathbf{A}^T \mathbf{z} \geq (1 - \epsilon)^2 \sigma_1^2(\mathbf{A})$. Taking square root on both sides gives us, $\|\mathbf{A}^T \mathbf{z}\|_2 \geq (1 - \epsilon) \sigma_1(\mathbf{A})$. $\square$

Next we prove the correctness of Algorithm 2 using Lemma 5. The idea is to run power iterations using all n basis vectors as starting vectors. At least one of these vectors will have high inner product with $\mathbf{u}_1$, and so so this starting vector will give us a good approximation to $\sigma_1(\mathbf{A})$ by Lemma 5. We also prove that the squared singular values of the deflated matrix $\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}$ are close to the squared singular values of the original matrix up to an additive error of $\epsilon \cdot \sigma_1^2(\mathbf{A})$. This will be used to as part of the correctness proof for our main algorithm for singular value approximation.

Lemma 6. *Let $\mathbf{z}$ be the output of Algorithm 2 for some matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ with singular values $\sigma_1(\mathbf{A}) \geq \sigma_2(\mathbf{A}) \geq \dots \geq \sigma_n(\mathbf{A})$ and error parameter $\epsilon \in (0, 1)$ as input. Then:*

$$(1 - \epsilon) \sigma_1(\mathbf{A}) \leq \|\mathbf{A}^T \mathbf{z}\|_2 \leq \sigma_1(\mathbf{A}).$$

Further, for any $i \in [n]$:

$$\sigma_i^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}) \leq \sigma_{i+1}^2(\mathbf{A}) + \epsilon \sigma_1^2(\mathbf{A}).$$

Proof. Let $\mathbf{u}_1$ be the left singular vector corresponding to the largest singular value of $\mathbf{A}$. Observe that since $\|\mathbf{u}_1\|_2 = 1$, there exists at least one standard basis vector $\mathbf{e}_k$ such that $|\mathbf{e}_k^T \cdot \mathbf{u}_1| \geq \frac{1}{\sqrt{n}}$. After

t iterations of the inner loop of Algorithm 2, we have $\mathbf{y}_k = \frac{(\mathbf{A}\mathbf{A}^T)^t \mathbf{e}_k}{\|(\mathbf{A}\mathbf{A}^T)^t \mathbf{e}_k\|_2}$. Using Lemma 5, we thus have $\|\mathbf{A}^T \mathbf{y}_k\|_2^2 \geq (1 - \epsilon)\sigma_1(\mathbf{A})$. Also since $\mathbf{y}_k$ is a unit vector, we trivially have $\|\mathbf{A}^T \mathbf{y}_k\|_2 \leq \sigma_1(\mathbf{A})$ for all k . Since for $\mathbf{z}$ output by Algorithm 2, $\|\mathbf{A}^T \mathbf{z}\|_2 = \max_k \|\mathbf{A}^T \mathbf{y}_k\|_2$, we thus have:

$$(1 - \epsilon)\sigma_1(\mathbf{A}) \leq \|\mathbf{A}^T \mathbf{z}\|_2 \leq \sigma_1(\mathbf{A}), \quad (15)$$

giving our first error bound.

We now prove the second bound on $\sigma_i^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A})$. Using the Pythagorean theorem, $\|\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}\|_F^2 = \|\mathbf{A}\|_F^2 - \|\mathbf{z}\mathbf{z}^T \mathbf{A}\|_F^2$. We have $\|\mathbf{z}\mathbf{z}^T \mathbf{A}\|_F^2 = \text{tr}(\mathbf{z}\mathbf{z}^T \mathbf{A} \mathbf{A}^T \mathbf{z}\mathbf{z}^T) = \mathbf{z}^T \mathbf{A} \mathbf{A}^T \mathbf{z} = \|\mathbf{A}^T \mathbf{z}\|_2^2$. We thus have $\|\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}\|_F^2 = \|\mathbf{A}\|_F^2 - \|\mathbf{A}^T \mathbf{z}\|_2^2$. Combining this with (15):

$$\begin{aligned} \|\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}\|_F^2 &\leq \|\mathbf{A}\|_F^2 - (1 - \epsilon)^2 \sigma_1^2(\mathbf{A}) \\ &\leq \|\mathbf{A}\|_F^2 - \sigma_1^2(\mathbf{A}) + 2\epsilon \sigma_1^2(\mathbf{A}) \\ &= \|\mathbf{A} - \mathbf{A}_1\|_F^2 + 2\epsilon \sigma_1^2(\mathbf{A}), \end{aligned} \quad (16)$$

where $\mathbf{A}_1 = \mathbf{u}_1 \mathbf{u}_1^T \mathbf{A}$ is the best rank-1 approximation to $\mathbf{A}$ in the Frobenius norm, with $\|\mathbf{A} - \mathbf{A}_1\|_F^2 = \sum_{i=2}^n \sigma_i^2(\mathbf{A})$. So, we have:

$$\sum_{j=1}^n \sigma_j^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}) = \|\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}\|_F^2 \leq \|\mathbf{A} - \mathbf{A}_1\|_F^2 + 2\epsilon \sigma_1^2(\mathbf{A}) = \sum_{j=2}^n \sigma_j^2(\mathbf{A}) + 2\epsilon \sigma_1^2(\mathbf{A}).$$

Simply subtracting $\sigma_n^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A})$ from the lefthand side and pulling $\sigma_i^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A})$ out of the summation for some $i \in [n]$ gives that

$$\sigma_i^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}) + \sum_{j \neq i, j \in [n-1]} \sigma_j^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}) \leq \sum_{j=2}^n \sigma_j^2(\mathbf{A}) + 2\epsilon \sigma_1^2(\mathbf{A}). \quad (17)$$

Using the eigenvalue min-max theorem (Fact 4), we have for all $j \in [n-1]$, $\sigma_j(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}) \geq \sigma_{j+1}(\mathbf{A})$. Using this fact in (17),

$$\sigma_i^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}) + \sum_{j \neq i, j \in [n-1]} \sigma_{j+1}^2(\mathbf{A}) \leq \sum_{j=2}^n \sigma_j^2(\mathbf{A}) + 2\epsilon \sigma_1^2(\mathbf{A}),$$

which implies that $\sigma_i^2(\mathbf{A} - \mathbf{z}\mathbf{z}^T \mathbf{A}) \leq \sigma_{i+1}^2(\mathbf{A}) + 2\epsilon \sigma_1^2(\mathbf{A})$. This completes the proof after adjusting ϵ by a factor of 2. $\square$

We now state as Algorithm 3 our main deterministic algorithm for estimating k singular values of a matrix by leveraging Algorithm 2 as a subroutine. We will then show that by applying Algorithm 3 to a sparse spectral approximation $\tilde{\mathbf{A}}$ of $\mathbf{A}$, we can estimate the singular values of $\mathbf{A}$ up to error $\pm \epsilon \max(n, \|\mathbf{A}\|_1)$ in $\tilde{O}(n^2 / \text{poly}(\epsilon))$ time.

Algorithm 3 Deterministic Singular Value Estimation via Deflation

Input: $\mathbf{A} \in \mathbb{R}^{n \times n}$, error parameter $\epsilon \in (0, 1)$, number of singular values to be estimated $k \leq n$

- 1: $\mathbf{A}^{(1)} = \mathbf{A}$.
 - 2: **for** $i = 1 \rightarrow k$ **do**
 - 3: $\mathbf{z}_i \leftarrow$ output of Algorithm 2 with input $\mathbf{A}^{(i)}$, error parameter ϵ .
 - 4: $\mathbf{A}^{(i+1)} \leftarrow \mathbf{A}^{(i)} - \mathbf{z}_i \mathbf{z}_i^T \mathbf{A}$.
 - 5: **end for**
 - 6: **return** $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_k$
-

Lemma 7. Let $\{\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_k\}$ be the output of Algorithm 3 for some matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ with singular values $\sigma_1(\mathbf{A}) \geq \sigma_2(\mathbf{A}) \geq \dots \geq \sigma_n(\mathbf{A})$ and error parameter $\epsilon \in (0, 1)$ as input. Then, $\{\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_k\}$ are orthogonal unit vectors and for any $i \in [k]$:

$$(1 - \epsilon)\sigma_i(\mathbf{A}) \leq \|\mathbf{A}^T \mathbf{z}_i\|_2 \leq \sigma_i(\mathbf{A}) + \sigma_1(\mathbf{A})\sqrt{i \cdot \epsilon}$$

Proof. Let $\tilde{\sigma}_i = \|(\mathbf{A}^{(i)})^T \mathbf{z}_i\|_2$ for all i . Since each $\mathbf{z}_i$ is the output of Algorithm 2 with $\mathbf{A}^{(i)}$ as the input, using the bound of Lemma 6, for any $i \in [k]$ we get that

$$(1 - \epsilon)\sigma_1(\mathbf{A}^{(i)}) \leq \tilde{\sigma}_i \leq \sigma_1(\mathbf{A}^{(i)}). \quad (18)$$

Note that $\mathbf{A}^{(i)} = \mathbf{A}^{(i-1)} - \mathbf{z}_{i-1}\mathbf{z}_{i-1}^T \mathbf{A} = \mathbf{A} - \sum_{j=1}^{i-1} \mathbf{z}_j \mathbf{z}_j^T \mathbf{A} = \mathbf{A} - \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T \mathbf{A}$ where $\mathbf{Z}_{i-1}$ is a matrix with $i-1$ columns where the j^{th} column is equal to $\mathbf{z}_j$. We first prove that $\mathbf{Z}_{i-1}$ is a matrix with orthonormal columns. First note that each $\mathbf{z}_i$ is a unit vector since Algorithm 2 always outputs a unit vector. We can prove that all $\mathbf{z}_i$ are orthogonal to each other by induction. Suppose that all $\mathbf{z}_k$ with $k \in [j]$ are orthogonal to each other for some $j < i-1$. Now, $\mathbf{A}^{(j+1)} = \mathbf{A} - \mathbf{Z}_j \mathbf{Z}_j^T \mathbf{A} = (\mathbf{I} - \mathbf{Z}_j \mathbf{Z}_j^T) \mathbf{A}$ where $\mathbf{I}$ is the identity matrix. Observe that $\mathbf{z}_{j+1}$ lies in the column span of $\mathbf{A}^{(j+1)}$. Thus, to prove that $\mathbf{z}_{j+1}$ is orthogonal to all $\mathbf{z}_k$ with $k \in [j]$, it is enough to show that each such $\mathbf{z}_k$ is orthogonal to the column space of $\mathbf{A}^{(j+1)}$. This follows since $\mathbf{z}_k = \mathbf{Z}_j \mathbf{Z}_j^T \mathbf{z}_k$ since $\mathbf{Z}_j \mathbf{Z}_j^T$ is a projection onto the span of $\mathbf{z}_1, \dots, \mathbf{z}_j$. Thus, $(\mathbf{A}^{(j+1)})^T \mathbf{z}_k = \mathbf{A}^T (\mathbf{I} - \mathbf{Z}_j \mathbf{Z}_j^T) \mathbf{z}_k = \mathbf{A}^T (\mathbf{z}_k - \mathbf{z}_k) = 0$. Thus, $\mathbf{z}_{j+1}$ is orthogonal to all $\mathbf{z}_k$ with $k \in [j]$. Hence, $\mathbf{Z}_{i-1}$ is a matrix with orthonormal columns. This implies that $\mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T$ is a rank $(i-1)$ projection matrix.

By the above orthogonality claim, $\mathbf{Z}_{i-1}^T \mathbf{z}_i = \mathbf{0}$ and we have $\tilde{\sigma}_i = \|(\mathbf{A}^{(i)})^T \mathbf{z}_i\|_2 = \|(\mathbf{A}^T - \mathbf{A}^T \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T) \mathbf{z}_i\|_2 = \|\mathbf{A}^T \mathbf{z}_i\|_2$. So, to prove the lemma, it is enough to bound $\tilde{\sigma}_i$. We first prove the lower bound on $\tilde{\sigma}_i$ using the min-max principle of eigenvalues, which we state below.

Fact 4 (Eigenvalue Minimax Principle – see e.g., [Bha13]). Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a symmetric matrix with eigenvalues $\lambda_1(\mathbf{A}) \geq \lambda_2(\mathbf{A}) \geq \dots \geq \lambda_n(\mathbf{A})$, then for all $i \in [n]$,

$$\lambda_i(\mathbf{A}) = \max_{S: \dim(S)=i} \min_{\mathbf{x} \in S, \|\mathbf{x}\|_2=1} \mathbf{x}^T \mathbf{A} \mathbf{x} = \min_{S: \dim(S)=n-i+1} \max_{\mathbf{x} \in S, \|\mathbf{x}\|_2=1} \mathbf{x}^T \mathbf{A} \mathbf{x}.$$

By Fact 4, $\sigma_1(\mathbf{A}^{(i)})^2 = \sigma_1(\mathbf{A} - \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T \mathbf{A})^2 = \lambda_1((\mathbf{I} - \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T) \mathbf{A} \mathbf{A}^T (\mathbf{I} - \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T)) \geq \lambda_i(\mathbf{A} \mathbf{A}^T) = \sigma_i(\mathbf{A})^2$ since $\mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T$ is a rank $i-1$ projection matrix. Thus, from (18), we get $\tilde{\sigma}_i \geq (1 - \epsilon)\sigma_1(\mathbf{A}^{(i)}) \geq (1 - \epsilon)\sigma_i(\mathbf{A})$.

We now prove the upper bound on $\tilde{\sigma}_i$. For any $i \in [k]$ and any $j \in [n]$ from Lemma 6 we have:

$$\sigma_j^2(\mathbf{A}^{(i)} - \mathbf{z}_i \mathbf{z}_i^T \mathbf{A}^{(i)}) \leq \sigma_{j+1}^2(\mathbf{A}^{(i)}) + \epsilon \sigma_1^2(\mathbf{A}^{(i)}).$$

Now using the fact that $\mathbf{A}^{(i)} - \mathbf{z}_i \mathbf{z}_i^T \mathbf{A}^{(i)} = \mathbf{A} - \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T \mathbf{A} - \mathbf{z}_i \mathbf{z}_i^T \mathbf{A} = \mathbf{A} - \mathbf{Z}_i \mathbf{Z}_i^T \mathbf{A}$ and $\sigma_1(\mathbf{A}^{(i)}) \leq \sigma_1(\mathbf{A})$, we get:

$$\sigma_j^2(\mathbf{A} - \mathbf{Z}_i \mathbf{Z}_i^T \mathbf{A}) \leq \sigma_{j+1}^2(\mathbf{A} - \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T \mathbf{A}) + \epsilon \sigma_1^2(\mathbf{A}). \quad (19)$$

From (18) we have $\tilde{\sigma}_i \leq \sigma_1(\mathbf{A} - \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T \mathbf{A})$. Squaring both sides and applying the bound from (19) $i-1$ times repeatedly on the RHS, we get

$$\begin{aligned} \tilde{\sigma}_i^2 &\leq \sigma_1^2(\mathbf{A} - \mathbf{Z}_{i-1} \mathbf{Z}_{i-1}^T \mathbf{A}) \leq \sigma_2^2(\mathbf{A} - \mathbf{Z}_{i-2} \mathbf{Z}_{i-2}^T \mathbf{A}) + \epsilon \sigma_1^2(\mathbf{A}) \\ &\leq \sigma_3^2(\mathbf{A} - \mathbf{Z}_{i-3} \mathbf{Z}_{i-3}^T \mathbf{A}) + 2\epsilon \sigma_1^2(\mathbf{A}) \leq \dots \leq \sigma_i^2(\mathbf{A}) + i \cdot \epsilon \cdot \sigma_1^2(\mathbf{A}). \end{aligned}$$

Finally, taking the square root on both sides and using that for any non-negative a, b $\sqrt{a} + \sqrt{b} \geq \sqrt{a+b}$ gives us the upper bound. $\square$

Lemma 7 in place, we are now ready to state and prove our main result on deterministic singular value estimation for any bounded entry matrix.

Theorem 8 (Deterministic Singular Value Approximation). *Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a bounded entry symmetric matrix with singular values $\sigma_1(\mathbf{A}) \geq \sigma_2(\mathbf{A}) \geq \dots \geq \sigma_n(\mathbf{A})$. Then there exists a deterministic algorithm that, given $\epsilon \in (0, 1)$, reads $\tilde{O}(\frac{n}{\epsilon^4})$ entries of $\mathbf{A}$, runs in $\tilde{O}(\frac{n^2}{\epsilon^8})$ time, and returns singular value approximations $\tilde{\sigma}_1(\mathbf{A}), \dots, \tilde{\sigma}_n(\mathbf{A})$ satisfying for all i ,*

$$|\sigma_i(\mathbf{A}) - \tilde{\sigma}_i(\mathbf{A})| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1).$$

Further, for all $i \leq 1/\epsilon$, the algorithm returns a unit vector $\mathbf{z}_i$ such that $|\|\mathbf{A}\mathbf{z}_i\|_2 - \sigma_i(\mathbf{A})| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ and the returned vectors are all mutually orthogonal.

Proof. Note that $\mathbf{A}$ can have at most $\frac{1}{\epsilon}$ singular values greater than $\epsilon\|\mathbf{A}\|_1$. Thus, it is enough to approximate only the top $\frac{1}{\epsilon}$ singular values of $\mathbf{A}$ up to error $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. The rest of the singular values can be approximated by 0 to still have the required approximation error of $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Let $\tilde{\mathbf{A}} = \mathbf{A} \circ \mathbf{S}$ be the sparsification of $\mathbf{A}$ as described in Theorem 4, which satisfies $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Using Weyl's inequality we have for all $i \in [n]$:

$$|\sigma_i(\mathbf{A}) - \sigma_i(\tilde{\mathbf{A}})| \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1).$$

Thus it is enough to approximate the top $\frac{1}{\epsilon}$ singular values of $\tilde{\mathbf{A}}$ up to $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ error and then use triangle inequality to get the final approximation to the top $\frac{1}{\epsilon}$ singular values of $\mathbf{A}$. To do so, we run Algorithm 3 with $\tilde{\mathbf{A}}$ as the input matrix, ϵ^3 as the error parameter, and $k = \lceil 1/\epsilon \rceil$ as the number of singular values to be estimated. Using Lemma 7 we get orthonormal vectors $\mathbf{z}_i$ for $i \in \lceil \frac{1}{\epsilon} \rceil$ such that:

$$(1 - \epsilon^3)\sigma_i(\tilde{\mathbf{A}}) \leq \|\tilde{\mathbf{A}}\mathbf{z}_i\|_2 \leq \sigma_i(\tilde{\mathbf{A}}) + \sigma_1(\tilde{\mathbf{A}})\sqrt{i \cdot \epsilon^3}. \quad (20)$$

Let $\tilde{\sigma}_i(\mathbf{A}) = \|\tilde{\mathbf{A}}\mathbf{z}_i\|_2$ for $i \in \lceil \frac{1}{\epsilon} \rceil$. Now since $\sigma_1(\mathbf{A}) \leq n$, using Weyl's inequality, we have $\sigma_1(\tilde{\mathbf{A}}) \leq \sigma_1(\mathbf{A}) + \|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \sigma_1(\mathbf{A}) + \epsilon \max(n, \|\mathbf{A}\|_1) \leq 2 \max(n, \|\mathbf{A}\|_1)$. Thus, for any $i \in \lceil \frac{1}{\epsilon} \rceil$, $\sigma_1(\tilde{\mathbf{A}})\sqrt{i \cdot \epsilon^3} \leq 2\epsilon \max(n, \|\mathbf{A}\|_1)$. So, from Equation (20), we get:

$$(1 - \epsilon)\sigma_i(\tilde{\mathbf{A}}) \leq \tilde{\sigma}_i(\mathbf{A}) \leq \sigma_i(\tilde{\mathbf{A}}) + 2\epsilon \cdot \max(n, \|\mathbf{A}\|_1).$$

This gives us $|\tilde{\sigma}_i(\mathbf{A}) - \sigma_i(\tilde{\mathbf{A}})| \leq 2\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Finally, we get:

$$|\tilde{\sigma}_i(\mathbf{A}) - \sigma_i(\mathbf{A})| \leq |\tilde{\sigma}_i(\mathbf{A}) - \sigma_i(\tilde{\mathbf{A}})| + |\sigma_i(\tilde{\mathbf{A}}) - \sigma_i(\mathbf{A})| \leq 3\epsilon \max(n, \|\mathbf{A}\|_1),$$

where the second inequality follows from the bound on $|\tilde{\sigma}_i(\mathbf{A}) - \sigma_i(\tilde{\mathbf{A}})|$ and Weyl's inequality which gives us $|\sigma_i(\tilde{\mathbf{A}}) - \sigma_i(\mathbf{A})| \leq \epsilon \max(n, \|\mathbf{A}\|_1)$. This gives us the required additive approximation error after adjusting ϵ by constants. Next, observe that for $i \in \lceil \frac{1}{\epsilon} \rceil$, using triangle inequality:

$$|\|\mathbf{A}\mathbf{z}_i\|_2 - \sigma_i(\mathbf{A})| \leq |\|\mathbf{A}\mathbf{z}_i\|_2 - \|\tilde{\mathbf{A}}\mathbf{z}_i\|_2| + |\|\tilde{\mathbf{A}}\mathbf{z}_i\|_2 - \sigma_i(\mathbf{A})| \leq 4\epsilon \max(n, \|\mathbf{A}\|_1),$$

where in the second step, the first term is bounded as $|\|\mathbf{A}\mathbf{z}_i\|_2 - \|\tilde{\mathbf{A}}\mathbf{z}_i\|_2| \leq \|(\mathbf{A} - \tilde{\mathbf{A}})\mathbf{z}_i\|_2 \leq \|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \epsilon \max(n, \|\mathbf{A}\|_1)$. This gives us the second bound (after adjusting ϵ by constants).

Sample Complexity and Runtime Analysis. By Theorem 4, the number of queries to $\mathbf{A}$'s entries need to construct $\tilde{\mathbf{A}} = \mathbf{A} \circ \mathbf{S}$ is $\tilde{O}(\frac{n}{\epsilon^4})$. The loop estimating the singular values in Algorithm 3 runs $\frac{1}{\epsilon}$ times and Algorithm 2 is called in Line 3 inside the loop each time with error parameter ϵ^3 . At the i^{th} iteration of the loop in Algorithm 3, in Line 4 we have $\tilde{\mathbf{A}}^{(i+1)} = \tilde{\mathbf{A}}^{(i)} - \mathbf{z}_i^T \mathbf{z}_i^T \tilde{\mathbf{A}} = \tilde{\mathbf{A}} - \mathbf{z}_i^T \mathbf{z}_i^T \tilde{\mathbf{A}}$.

Note that since the matrix $\tilde{\mathbf{A}}^{(i+1)}$ could be dense we don't explicitly compute $\tilde{\mathbf{A}}^{(i+1)}$ to do power iterations in Algorithm 2 with this matrix. Instead, in each step of power iteration in Line 5 of Algorithm 2, we can first calculate $\tilde{\mathbf{A}}\mathbf{y}_j$ in time $\tilde{O}(\frac{n}{\epsilon^4})$ and then multiply this vector first with $\mathbf{Z}_i^T$ and then with $\mathbf{Z}_i$ in $O(\frac{n}{\epsilon})$ time to get $\mathbf{Z}_i\mathbf{Z}_i^T\tilde{\mathbf{A}}$. Thus, each power iteration (Lines 5 and 6 of Algorithm 2) takes $\tilde{O}(\frac{n}{\epsilon^4})$ time. Since we call Algorithm 2 with error parameter ϵ^3 , the number of power iterations with each starting vector is $O(\frac{\log(n/\epsilon)}{\epsilon^3})$. There are also n different starting vectors. Thus, each call to Algorithm 2 in Line 3 of Algorithm 3 takes $\tilde{O}(\frac{n^2}{\epsilon^7})$ time. Thus, the total running time of the algorithm is $\tilde{O}(\frac{n^2}{\epsilon^8})$ (as the loop in Algorithm 3 runs $\frac{1}{\epsilon}$ times). $\square$

6.2 Deterministic PSD Testing

We next leverage our deterministic power method approach (Algorithm 2) to give an $o(n^\omega)$ time deterministic algorithm for testing if a bounded entry matrix is either PSD or has at least one large negative eigenvalue $\leq -\epsilon \max(n, \|\mathbf{A}\|_1)$. An optimal randomized algorithm for this problem with detection threshold $-\epsilon n$ was presented in [BCJ20]. The idea of our approach is to approximate the maximum singular value of $\mathbf{I} - \frac{\tilde{\mathbf{A}}}{\|\tilde{\mathbf{A}}\|_2}$ which will be at smaller than 1 if $\mathbf{A}$ is PSD and greater than 1 if $\mathbf{A}$ is ϵn far from being PSD.

Theorem 14 (Deterministic PSD testing with ℓ_∞ -gap). *Given $\epsilon \in (0, 1)$, let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a symmetric bounded entry matrix with eigenvalues $\lambda_1(\mathbf{A}) \geq \lambda_2(\mathbf{A}) \geq \dots \geq \lambda_n(\mathbf{A})$ such that either $\mathbf{A}$ is PSD i.e. $\lambda_n(\mathbf{A}) \geq 0$ or $\mathbf{A}$ is at least $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ -far from being PSD, i.e., $\lambda_n(\mathbf{A}) \leq -\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. There exists a deterministic algorithm that distinguishes between these two cases by reading $\tilde{O}(\frac{n}{\epsilon^4})$ entries of $\mathbf{A}$ and which runs in time $\tilde{O}(\frac{n^2}{\epsilon^5})$.*

Proof. Let $\tilde{\mathbf{A}} = \mathbf{A} \circ \mathbf{S}$ be a sparsification of $\mathbf{A}$ as described in Theorem 4 with approximation parameter $\frac{\epsilon}{10}$ that satisfies $\|\mathbf{A} - \tilde{\mathbf{A}}\|_2 \leq \frac{\epsilon}{10} \max(n, \|\mathbf{A}\|_1)$. Using Weyl's inequality we have for all $i \in [n]$, $|\lambda_i(\mathbf{A}) - \lambda_i(\tilde{\mathbf{A}})| \leq \frac{\epsilon}{10} \max(n, \|\mathbf{A}\|_1)$.

Let $\tilde{\sigma}_1(\tilde{\mathbf{A}})$ be the estimate of $\|\tilde{\mathbf{A}}\|_2 = \sigma_1(\tilde{\mathbf{A}})$ output by Algorithm 2 with input $\tilde{\mathbf{A}}$ and error parameter $\frac{\epsilon}{100}$. Then by Lemma 6, we have:

$$\left(1 - \frac{\epsilon}{100}\right) \sigma_1(\tilde{\mathbf{A}}) \leq \tilde{\sigma}_1(\tilde{\mathbf{A}}) \leq \sigma_1(\tilde{\mathbf{A}}).$$

Let $\mathbf{B} = \mathbf{I} - \frac{1}{\tilde{\sigma}_1(\tilde{\mathbf{A}})} \tilde{\mathbf{A}}$. Let $\tilde{\sigma}_1(\mathbf{B})$ be the estimate of $\sigma_1(\mathbf{B})$ output by Algorithm 2 with input $\mathbf{B}$ and error parameter $\frac{\epsilon}{1000}$. Then by Lemma 6, we again have:

$$\left(1 - \frac{\epsilon}{1000}\right) \sigma_1(\mathbf{B}) \leq \tilde{\sigma}_1(\mathbf{B}) \leq \sigma_1(\mathbf{B}).$$

We now distinguish between the two cases as follows:

Case 1: When $\mathbf{A}$ is PSD, we have $\lambda_n(\tilde{\mathbf{A}}) \geq \lambda_n(\mathbf{A}) - \frac{\epsilon n}{10} \geq -\frac{\epsilon n}{10}$. Now, $\tilde{\sigma}_1(\mathbf{B}) \leq \sigma_1(\mathbf{B}) = \max(|\lambda_1(\mathbf{B})|, |\lambda_n(\mathbf{B})|) = \max\left(\left|1 - \frac{\lambda_n(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\right|, \left|1 - \frac{\lambda_1(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\right|\right)$.

First assume that $\left|1 - \frac{\lambda_n(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\right| \geq \left|1 - \frac{\lambda_1(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\right|$. Now, $1 - \frac{\lambda_n(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})} \leq 1 + \frac{\epsilon n}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}$. Thus, $\tilde{\sigma}_1(\mathbf{B}) \leq 1 + \frac{\epsilon n}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}$. Now, assume $\left|1 - \frac{\lambda_n(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\right| < \left|1 - \frac{\lambda_1(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\right|$. Then, we must have $\lambda_1(\tilde{\mathbf{A}}) \geq 0$ as otherwise, we will have $1 - \frac{\lambda_n(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})} \geq 1 - \frac{\lambda_1(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})} > 0$. Thus, we get that $\tilde{\sigma}_1(\mathbf{B}) \leq \left|1 - \frac{\lambda_1(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\right| = \left|1 - \frac{\sigma_1(\tilde{\mathbf{A}})}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\right| = \left|1 - \frac{1}{1 - \frac{\epsilon}{100}}\right| = \frac{\epsilon/100}{1 - \epsilon/100} < \frac{\epsilon}{50} < 1 + \frac{\epsilon n}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}$.

Case 2: When $\mathbf{A}$ is $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ far from being PSD, we have $\lambda_n(\tilde{\mathbf{A}}) \leq \lambda_n(\mathbf{A}) + \frac{\epsilon \cdot \max(n, \|\mathbf{A}\|_1)}{10} \leq -\frac{9}{10}\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Observe that $\sigma_1(\mathbf{B}) \geq 1 - \frac{1}{\tilde{\sigma}_1(\tilde{\mathbf{A}})}\lambda_n(\tilde{\mathbf{A}}) \geq 1 + \frac{9}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. Then, we have $\tilde{\sigma}_1(\mathbf{B}) \geq (1 - \frac{\epsilon}{1000})\sigma_1(\mathbf{B}) \geq (1 - \frac{\epsilon}{1000})\left(1 + \frac{9}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}\epsilon \cdot \max(n, \|\mathbf{A}\|_1)\right) > 1 + \frac{2}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}\epsilon \cdot \max(n, \|\mathbf{A}\|_1) > 1 + \frac{2\epsilon n}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}$.

Thus, in **Case 1**, we have $\tilde{\sigma}_1(\mathbf{B}) \leq 1 + \frac{\epsilon n}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}$ while for **Case 2**, we have $\tilde{\sigma}_1(\mathbf{B}) > 1 + \frac{2\epsilon n}{10\tilde{\sigma}_1(\tilde{\mathbf{A}})}$. So we can distinguish between the above two cases.

Sample Complexity and Runtime Analysis. The total runtime is the time taken to estimate $\sigma_1(\tilde{\mathbf{A}})$ and $\sigma_1(\mathbf{B})$ using Algorithm 2 with error parameter $\Theta(\epsilon)$. This is $\tilde{O}\left(\frac{n^2}{\epsilon^5}\right)$ since we run $O\left(\frac{\log(n/\epsilon)}{\epsilon}\right)$ iterations of the power method with n different starting vectors, where each iteration takes time $\tilde{O}\left(\frac{n}{\epsilon^4}\right)$ since $\tilde{\mathbf{A}}$ and $\mathbf{B}$ both have $\tilde{O}(n/\epsilon^4)$ non-zero entries. The sample complexity required for form $\tilde{\mathbf{A}}$ and $\mathbf{B}$ is $\tilde{O}(n/\epsilon^4)$. $\square$

6.3 High Accuracy Spectral Norm Approximation

Finally, we show that, under the assumption that $\sigma_1(\mathbf{A}) \geq \alpha \max(n, \|\mathbf{A}\|_1)$ for some $\alpha \in (0, 1)$, we can deterministically compute $\tilde{\sigma}_1(\mathbf{A})$ such that $\tilde{\sigma}_1(\mathbf{A}) \geq (1 - \epsilon)\sigma_1(\mathbf{A})$ in $\tilde{O}\left(\frac{n^2 \log(1/\epsilon)}{\text{poly}(\alpha)}\right)$ time. This allows us to set ϵ to $\frac{1}{n^c}$ to get a high accuracy approximation to $\sigma_1(\mathbf{A})$ in roughly linear time in the input matrix size. This is the first $o(n^\omega)$ time deterministic algorithm for this problem, and to the best of our knowledge matches the best known randomized methods for high accuracy top singular value computation, up to a $\text{poly} \log(n, 1/\alpha)$ factor. Our approach is described in Algorithm 4 below.

Algorithm 4 High Accuracy Spectral Norm Approximator

Input: Symmetric $\mathbf{A} \in \mathbb{R}^{n \times n}$, parameter α , and error parameter $\epsilon \in (0, 1)$

- 1: Let $\tilde{\mathbf{A}} = \mathbf{A} \circ \mathbf{S}$ where $\mathbf{S}$ is as specified in Theorem 4 with error parameter $\epsilon = c\alpha^4$ for a sufficiently small constant c .
 - 2: $\mathbf{Z} \leftarrow$ output of Algorithm 3 with inputs $\tilde{\mathbf{A}}$, error parameter $\tilde{\epsilon} = c\alpha^4$, and $k = 2/\alpha$.
 - 3: $\tilde{\mathbf{Z}} \leftarrow$ Output of Block Krylov method (Algorithm 2 of [MM15]) with input matrix $\mathbf{A}$, starting block $\mathbf{Z}$, and number of iterations $q = \frac{c' \log(n/\epsilon)}{\alpha^{3/2}}$ for some enough large constant c' .
 - 4: **return** $\tilde{\sigma}_1(\mathbf{A}) = \sqrt{\tilde{\mathbf{z}}_1^T \mathbf{A} \tilde{\mathbf{z}}_1}$, where $\tilde{\mathbf{z}}_1$ is the first column of $\tilde{\mathbf{Z}}$.
-

Algorithm 4 first computes (Step 2) a coarse approximate subspace of top singular vectors for $\mathbf{A}$ using our deflation based Algorithm 3 applied to the sparsified matrix $\tilde{\mathbf{A}}$. This subspace is then used in Step 3 to initialize a Block Krylov method applied to $\mathbf{A}$ to compute a much more accurate approximation to $\sigma_1(\mathbf{A})$. Our proof will leverage the analysis of Block Krylov methods given in [MM15]. Specifically, we will apply Theorem 13 of [MM15] which bounds the number of iterations required to find a $(1 \pm \epsilon)$ accurate approximation to the top k singular values in terms of the singular value gap $\frac{\sigma_k(\mathbf{A})}{\sigma_{p+1}(\mathbf{A})}$ for some block size $p \geq k$.

The proof of Theorem 13 in [MM15] uses a degree q amplifying polynomial $f(\mathbf{A})$ (for large enough q) in the Krylov subspace to significantly separate the top k singular values from the rest. It is shown that if the starting block $\mathbf{Z}$ is such that $f(\mathbf{A})$ projected onto the column span of $f(\mathbf{A})\mathbf{Z}$ is a good rank p approximation to $f(\mathbf{A})$, then the approximate singular values and vectors output by the Block Krylov method will be accurate. To prove this condition for our starting block, we

will apply the following well-known result, which bounds the error of any low rank approximation to a matrix in terms of the error of the best possible low rank approximation.

Lemma 8 (Lemma 48 of [Woo14], proved in [BDMI14]). *Let $\mathbf{A} = \mathbf{A}\mathbf{S}\mathbf{S}^T + \mathbf{E}$ be a low-rank matrix factorization of $\mathbf{A}$, with $\mathbf{S} \in \mathbb{R}^{n \times k}$, and $\mathbf{S}^T \mathbf{S} = \mathbf{I}_k$. Let $\mathbf{Z} \in \mathbb{R}^{n \times c}$ with $c \geq k$ be any matrix such that $\text{rank}(\mathbf{S}^T \mathbf{Z}) = \text{rank}(\mathbf{S}) = k$. Let $\mathbf{Q}$ be an orthonormal basis for the column span of $\mathbf{A}\mathbf{Z}$. Then,*

$$\|\mathbf{A} - \mathbf{Q}\mathbf{Q}^T \mathbf{A}\|_F^2 \leq \|\mathbf{E}\|_F^2 + \|\mathbf{E}\mathbf{Z}(\mathbf{S}^T \mathbf{Z})^+\|_F^2.$$

The above result lets us bound the difference between the Frobenius norm error of projecting $f(\mathbf{A})$ onto the column span of $f(\mathbf{A})\mathbf{Z}$, and the error of the best rank p approximation of $f(\mathbf{A})$. The proof of Theorem 13 of [MM15] uses Lemma 4 of [MM15], (which in turn is proven using Lemma 8 above) to show that $\|f(\mathbf{A}) - \mathbf{Q}\mathbf{Q}^T f(\mathbf{A})\|_F^2 \leq n^c \cdot \|f(\mathbf{A}) - (f(\mathbf{A}))_p\|_F^2$ where $\mathbf{Q}$ is an orthonormal basis for $\mathbf{A}\mathbf{Z}$, where $\mathbf{Z}$ is a random starting block, c is some constant, and $(f(\mathbf{A}))_p$ is the best rank p approximation to $f(\mathbf{A})$. In order to apply this result in our setting, we need to show that the same bound applies for our starting block $\mathbf{Z}$, which is computed deterministically by Algorithm 3.

In particular, by Lemma 8, it suffices to bound $\|(\mathbf{V}_p^T \mathbf{Z})^+\|_2$, where $\mathbf{V}_p$ contains the top p singular vectors of $\mathbf{A}$ (which are identical to those of $f(\mathbf{A})$). That is, we need to ensure that the columns of $\mathbf{Z}$ and the top p singular vectors of $\mathbf{A}$ have large inner product with each other. We next prove that this is the case when there is a large gap between the top p singular values and the rest. We will later show that under the assumption that $\sigma_1(\mathbf{A}) \geq \alpha \cdot \max(n, \|\mathbf{A}\|_1)$, there is always some p that satisfies this gap assumption.

Lemma 9 (Starting Block Condition). *Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a bounded entry symmetric matrix with singular values $\sigma_1(\mathbf{A}) \geq \sigma_2(\mathbf{A}) \geq \dots \geq \sigma_n(\mathbf{A})$. Let $\tilde{\mathbf{A}} = \mathbf{A} \circ \mathbf{S}$ where the sampling matrix $\mathbf{S}$ is as specified in Theorem 4 with error parameter $\epsilon = c\alpha^4$ for a sufficiently small constant c and some $\alpha \in (0, 1)$. For some $k \in [n]$, let $\mathbf{Z} \in \mathbb{R}^{n \times k}$ be the output of Algorithm 3 with input $\tilde{\mathbf{A}}$ and parameters $k = \frac{2}{\alpha}$ and $\epsilon = c\alpha^4$. For some $p \leq \frac{2}{\alpha}$, assume that for some constant c_1 ,*

$$\sigma_p^2(\mathbf{A}) - \sigma_{p+1}^2(\mathbf{A}) \geq c_1 \alpha^3 \cdot (\max(n, \|\mathbf{A}\|_1))^2.$$

Also let $\mathbf{V}_p \in \mathbb{R}^{n \times p}$ be the matrix with columns equal to the p singular vectors of $\mathbf{A}$ corresponding to its largest p singular values. Let $\mathbf{Z}_p$ contain the first p columns of $\mathbf{Z}$. Then for some constant c_2 ,

$$\|(\mathbf{V}_p^T \mathbf{Z}_p)^+\|_2 \leq c_2 n.$$

Proof. Before proving the main result, we will bound the error between $\sigma_j^2(\mathbf{A})$ and $\|\mathbf{A}\mathbf{z}_j\|_2^2$ where $\mathbf{z}_j$ is a column of $\mathbf{Z}$ for any $j \in [\frac{2}{\alpha}]$. From Lemma 7, we have for any $j \in [\frac{2}{\alpha}]$:

$$\|\tilde{\mathbf{A}}\mathbf{z}_j\|_2 \geq (1 - c\alpha^4)\sigma_j(\tilde{\mathbf{A}}),$$

since we run Algorithm 3 with error parameter $\epsilon = c\alpha^4$. Also since $\|\tilde{\mathbf{A}} - \mathbf{A}\|_2 \leq c\alpha^4 \max(n, \|\mathbf{A}\|_1)$, using Weyl's inequality and the fact that $\sigma_j(\mathbf{A}) \leq \max(n, \|\mathbf{A}\|_1)$ we have (for some constant b)

$$\sigma_j(\mathbf{A}) - c\alpha^4 \max(n, \|\mathbf{A}\|_1) \leq \sigma_j(\tilde{\mathbf{A}}) \leq \sigma_j(\mathbf{A}) + c\alpha^4 \max(n, \|\mathbf{A}\|_1) \leq b \max(n, \|\mathbf{A}\|_1).$$

Now, using triangle inequality,

$$\|\mathbf{A}\mathbf{z}_j\|_2 \geq \|\tilde{\mathbf{A}}\mathbf{z}_j\|_2 - \|(\tilde{\mathbf{A}} - \mathbf{A})\mathbf{z}_j\|_2 \geq (1 - c\alpha^4)\sigma_j(\tilde{\mathbf{A}}) - c\alpha^4 \max(n, \|\mathbf{A}\|_1),$$

where the second inequality uses that $\|(\tilde{\mathbf{A}} - \mathbf{A})\mathbf{z}_j\|_2 \leq \|\tilde{\mathbf{A}} - \mathbf{A}\|_2 \|\mathbf{z}_j\|_2 \leq c\alpha^2 \max(n, \|\mathbf{A}\|_1)$. Thus, using the bounds on $\sigma_j(\tilde{\mathbf{A}})$, we get:

$$\begin{aligned} \|\mathbf{A}\mathbf{z}_j\|_2 &\geq \sigma_j(\mathbf{A}) - c\alpha^4 \max(n, \|\mathbf{A}\|_1) - c\alpha^4 \sigma_j(\tilde{\mathbf{A}}) - c\alpha^4 \max(n, \|\mathbf{A}\|_1) \\ &\geq \sigma_j(\mathbf{A}) - c''\alpha^4 \max(n, \|\mathbf{A}\|_1), \end{aligned}$$

for some constant c'' . Finally, squaring both sides, we get for any $j \in [\frac{2}{\alpha}]$:

$$\begin{aligned} \|\mathbf{A}\mathbf{z}_j\|_2^2 &\geq \sigma_j^2(\mathbf{A}) - 2c''\alpha^4 \sigma_j(\mathbf{A}) \max(n, \|\mathbf{A}\|_1) + (c'')^2 \alpha^8 (\max(n, \|\mathbf{A}\|_1))^2 \\ &\geq \sigma_j^2(\mathbf{A}) - 2c''\alpha^4 (\max(n, \|\mathbf{A}\|_1))^2. \end{aligned} \quad (21)$$

We are now ready to prove the main result of the lemma. Note that showing $\|(\mathbf{V}_p^T \mathbf{Z}_p)^+\|_2 \leq c_2 n$ is equivalent to showing that $\sigma_{\min}(\mathbf{V}_p^T \mathbf{Z}_p) \geq \frac{1}{c_2 n}$. Assume for contradiction that $\sigma_{\min}(\mathbf{V}_p^T \mathbf{Z}_p) \leq \frac{1}{c_2 n}$. Then there is some vector $\mathbf{y}$ with $\|\mathbf{y}\|_2 = 1$ such that $\|\mathbf{Z}_p^T \mathbf{V}_p \mathbf{y}\|_2 \leq \frac{1}{c_2 n}$. Let $\mathbf{w} = \mathbf{V}_p \mathbf{y}$. Observe that $\|\mathbf{w}\|_2 = 1$ since $\mathbf{V}_p$ has orthonormal columns. Then, for any column $\mathbf{z}_i$ of $\mathbf{Z}_p$,

$$\|\mathbf{w}\mathbf{w}^T \mathbf{z}_i\|_2 = |\mathbf{w}^T \mathbf{z}_i| \leq \|\mathbf{w}^T \mathbf{Z}_p\|_2 \leq \frac{1}{c_2 n}, \quad (22)$$

where the last step follows by our assumption. Then, using triangle inequality and the above bound, we have $\|\mathbf{A}(\mathbf{I} - \mathbf{w}\mathbf{w}^T)\mathbf{z}_i\|_2 \geq \|\mathbf{A}\mathbf{z}_i\|_2 - \|\mathbf{A}\mathbf{w}\mathbf{w}^T \mathbf{z}_i\|_2 \geq \|\mathbf{A}\mathbf{z}_i\|_2 - \frac{\|\mathbf{A}\|_2}{c_2 n} \geq \|\mathbf{A}\mathbf{z}_i\|_2 - \frac{1}{c_2}$, where the last equality uses the fact that $\|\mathbf{A}\|_2 \leq n$ since it has bounded entries. Thus, for any $i \in [p]$, squaring both sides of this bound, we get:

$$\begin{aligned} \|\mathbf{A}(\mathbf{I} - \mathbf{w}\mathbf{w}^T)\mathbf{z}_i\|_2^2 &\geq \|\mathbf{A}\mathbf{z}_i\|_2^2 - \frac{2\|\mathbf{A}\mathbf{z}_i\|_2}{c_2} \\ &\geq \|\mathbf{A}\mathbf{z}_i\|_2^2 - \frac{2}{c_2} \max(n, \|\mathbf{A}\|_1) \\ &\geq \sigma_i^2(\mathbf{A}) - c_3 \alpha^4 \max(n, \|\mathbf{A}\|_1)^2, \end{aligned} \quad (23)$$

for some constant c_3 , where the second inequality uses $\|\mathbf{A}\mathbf{z}_i\|_2 \leq n$, and the last inequality uses (21). Moreover since $\mathbf{w}$ is spanned by $\mathbf{V}_p$, we have:

$$\|\mathbf{A}\mathbf{w}\|_2^2 \geq \min_{\mathbf{x}: \|\mathbf{x}\|_2=1} \|\mathbf{A}\mathbf{V}_p \mathbf{x}\|_2^2 \geq \sigma_p^2(\mathbf{A}) > \sigma_{p+1}^2(\mathbf{A}) + c_1 \alpha^3 \cdot (\max(n, \|\mathbf{A}\|_1))^2, \quad (24)$$

where the final inequality is by the gap assumption in the lemma statement.

Consider the matrix $\mathbf{W} = [\mathbf{w}, (\mathbf{I} - \mathbf{w}\mathbf{w}^T)\mathbf{Z}_p]$. Next, we show that under our assumption that $\sigma_{\min}(\mathbf{V}_p^T \mathbf{Z}_p) \leq \frac{1}{c_2 n}$, $\mathbf{W}$ will be very close to orthonormal, and further that by (23) and (24), $\|\mathbf{A}\mathbf{W}\|_F^2 > \sum_{i=1}^{p+1} \sigma_i^2(\mathbf{A})$. Since $\mathbf{W}$ has $p+1$ columns, this will lead to a contradiction since we must have $\|\mathbf{A}\mathbf{W}\|_F^2 \leq \|\mathbf{A}\mathbf{V}_{p+1}\|_F^2 = \sum_{i=1}^{p+1} \sigma_i^2(\mathbf{A})$. In particular, using (23) and (24),

$$\begin{aligned} \|\mathbf{A}\mathbf{W}\|_F^2 &= \|\mathbf{A}\mathbf{w}\|_2^2 + \sum_{i=1}^p \|\mathbf{A}(\mathbf{I} - \mathbf{w}\mathbf{w}^T)\mathbf{z}_i\|_2^2 \\ &\geq \sigma_{p+1}^2(\mathbf{A}) + c_1 \alpha^3 \cdot (\max(n, \|\mathbf{A}\|_1))^2 + \sum_{i=1}^p \sigma_i^2(\mathbf{A}) - c_3 p \alpha^4 \max(n, \|\mathbf{A}\|_1)^2 \\ &\geq \sum_{i=1}^{p+1} \sigma_i^2(\mathbf{A}) + c_4 \alpha^3 \max(n, \|\mathbf{A}\|_1)^2, \end{aligned} \quad (25)$$

for some constant c_4 which will be positive as long as c_1 is large enough compared to c_3 (which can be made arbitrary small by our setting of c_2, c in the lemma statement).

We will now show that all eigenvalues of $\mathbf{W}^T \mathbf{W}$ are close to 1, which implies that $\mathbf{W}$ is approximately orthogonal. This allows us to translate the bound on $\|\mathbf{A}\mathbf{W}\|_F^2$ given in (25) to an orthonormal basis $\mathbf{Q}$ for the column span of $\mathbf{W}$.

Observe that $(\mathbf{W}^T \mathbf{W})_{i,i} = \|\mathbf{w}\|_2^2 = 1$ for $i = 1$ and using (22), $|(\mathbf{W}^T \mathbf{W})_{i,i} - \mathbf{z}_{i-1}^T \mathbf{z}_{i-1}| = |\mathbf{z}_{i-1} \mathbf{w} \mathbf{w}^T \mathbf{z}_{i-1}| \leq \frac{1}{c_2^2 n^2}$ otherwise. Thus, $|(\mathbf{W}^T \mathbf{W})_{i,i} - 1| \leq \frac{1}{c_2^2 n^2}$ for all i . Further, $(\mathbf{W}^T \mathbf{W})_{i,j} = \mathbf{w}^T (\mathbf{I} - \mathbf{w} \mathbf{w}^T) \mathbf{z}_{j-1} = 0$ when $i \neq j$ and $i = 1$ as $\|\mathbf{w}\|_2 = 1$ (and similarly when $i \neq j$ and $j = 1$). Also, when $i \neq j$ and $i, j > 1$, $|(\mathbf{W}^T \mathbf{W})_{i,j}| = |\mathbf{z}_i^T \mathbf{z}_j - \mathbf{z}_i^T \mathbf{w} \mathbf{w}^T \mathbf{z}_j| = |\mathbf{z}_i^T \mathbf{w} \mathbf{w}^T \mathbf{z}_j| \leq \frac{1}{c_2^2 n^2}$ otherwise. We can then apply Gershgorin's circle theorem to bound the eigenvalues of $\mathbf{W}^T \mathbf{W}$.

Fact 5 (Gershgorin's circle theorem [Ger31]). *Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ with entries $\mathbf{A}_{ij}$. For $i \in [n]$, let $\mathbf{R}_i$ be the sum of absolute values of non-diagonal entries in the i^{th} row. Let $D(\mathbf{A}_{ii}, \mathbf{R}_i)$ be the closed disc centered at $\mathbf{A}_{ii}$ with radius $\mathbf{R}_i$. Then every eigenvalue of $\mathbf{A}$ lies within one of the discs $D(\mathbf{A}_{ii}, \mathbf{R}_i)$.*

From Fact 5, all eigenvalues of $\mathbf{W}^T \mathbf{W}$ lie in the range $\left[1 - \frac{1}{c_2^2 n}, 1 + \frac{1}{c_2^2 n}\right]$. Thus, $\mathbf{W}^T \mathbf{W}$ is a full rank matrix and so, $\mathbf{W}$ is a rank $p+1$ matrix. Let $\mathbf{Q}\mathbf{S}\mathbf{R}^T$ be the SVD of $\mathbf{W}$, where $\mathbf{Q} \in \mathbb{R}^{n \times p+1}$ and $\mathbf{R} \in \mathbb{R}^{p+1 \times p+1}$ are orthonormal matrices and $\mathbf{S} \in \mathbb{R}^{p+1 \times p+1}$ is a diagonal matrix of singular values of $\mathbf{W}$. Then, all eigenvalues of $(\mathbf{W}^T \mathbf{W})^{-1/2}$ or all diagonal entries of $\mathbf{S}^{-1}$ lie in $\left[1 - \frac{c_5}{n}, 1 + \frac{c_5}{n}\right]$ for some constant c_5 .

Now $\|\mathbf{A}\mathbf{Q}\|_F^2 = \|\mathbf{A}\mathbf{W}\mathbf{R}\mathbf{S}^{-1}\|_F^2 \geq (1 - \frac{c_6}{n}) \|\mathbf{A}\mathbf{W}\mathbf{R}\|_F^2 \geq (1 - \frac{c_6}{n}) \|\mathbf{A}\mathbf{W}\|_F^2$ (for some constant c_6) where the last step follows from the fact that $\|\mathbf{A}\mathbf{W}\mathbf{R}\|_F^2 = \|\mathbf{A}\mathbf{W}\|_F^2$ since $\mathbf{R}$ is an orthonormal square matrix. Then we have

$$\|\mathbf{A}\mathbf{Q}\|_F^2 \geq (1 - \frac{c_6}{n}) \|\mathbf{A}\mathbf{W}\|_F^2 \geq \sum_{i=1}^{p+1} \sigma_i^2(\mathbf{A}) + c_7 \alpha^3 \max(n, \|\mathbf{A}\|_1)^2, \quad (26)$$

for some constant c_7 , where the last inequality uses (25) and the fact that $(1 - \frac{c_6}{n}) c_4 \alpha^3 \max(n, \|\mathbf{A}\|_1)^2 - \frac{c_5 \sum_{i=1}^{p+1} \sigma_i^2(\mathbf{A})}{n} = \Omega(\alpha^3 \max(n, \|\mathbf{A}\|_1)^2)$. However, by the optimality of the SVD for Frobenius norm low-rank approximation (based on the eigenvalue min-max theorem of Fact 4), for any matrix $\mathbf{X} \in \mathbb{R}^{n \times (p+1)}$ with orthonormal columns, $\|\mathbf{A}\mathbf{X}\|_F^2 \leq \|\mathbf{A}\mathbf{V}_{p+1}\|_F^2 = \sum_{i=1}^{p+1} \sigma_i^2(\mathbf{A})$. This contradicts (26) which was derived using the assumption $\sigma_{\min}(\mathbf{V}_p^T \mathbf{Z}_p) \leq \frac{1}{c_2 n}$. $\square$

We now prove the main result of this section. We will first show that there will be a large (in terms of $O((\alpha \cdot \max(n, \|\mathbf{A}\|_1))^c)$) singular value gap for some $p \leq O(\frac{1}{\alpha})$ and then use Lemmas 8 and 9 to prove a gap dependent bound similar to Theorem 13 of [MM15].

Theorem 15 (High accuracy spectral norm approximation). *Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a bounded entry symmetric matrix with largest singular value $\sigma_1(\mathbf{A})$, such that $\sigma_1(\mathbf{A}) \geq \alpha \cdot \max(n, \|\mathbf{A}\|_1)$ for some $\alpha \in (0, 1)$. Then there exists an $\tilde{O}\left(\frac{n^2 \log(n/\epsilon)}{\alpha^{29}}\right)$ time deterministic algorithm that computes $\tilde{\mathbf{z}} \in \mathbb{R}^n$, such that $\|\mathbf{A}\tilde{\mathbf{z}}\|_2 \geq (1 - \epsilon)\sigma_1(\mathbf{A})$.*

Proof. We will first show that there exists some $p \in \frac{2}{\alpha}$ such that there is a large enough gap between the squared singular values $\sigma_p^2(\mathbf{A})$ and $\sigma_{p+1}^2(\mathbf{A})$. For $k = \frac{2}{\alpha}$ we have $\sigma_k(\mathbf{A}) \leq \frac{\alpha}{2} \cdot (\max(n, \|\mathbf{A}\|_1))$ and squaring both sides we have $\sigma_k^2(\mathbf{A}) \leq \frac{\alpha^2}{4} \cdot (\max(n, \|\mathbf{A}\|_1))^2$. Also since $\sigma_1(\mathbf{A}) \geq \alpha^2 \cdot \max(n, \|\mathbf{A}\|_1)^2$, we have $\sigma_1^2(\mathbf{A}) - \sigma_k^2(\mathbf{A}) \geq \frac{3}{4} \alpha^2 \cdot (\max(n, \|\mathbf{A}\|_1))^2$. Thus there exists $p \in [\frac{2}{\alpha}]$ such that

$$\sigma_p^2(\mathbf{A}) - \sigma_{p+1}^2(\mathbf{A}) \geq \frac{\frac{3}{4} \alpha^2 \cdot (\max(n, \|\mathbf{A}\|_1))^2}{\frac{2}{\alpha}} = \frac{3}{8} \alpha^3 \cdot (\max(n, \|\mathbf{A}\|_1))^2. \quad (27)$$

Let $\mathbf{Z}$ be the output of step 2 of Algorithm 4. Let $\mathbf{Z}_p \in \mathbb{R}^{n \times p}$ be the matrix with the first p columns of $\mathbf{Z}$. Then performing Block Krylov iterations (i.e. step 3 of Algorithm 4) with $\mathbf{Z}$ as a starting block can only yield a larger approximation for $\sigma_1(\mathbf{A})$ as compared to doing Block Krylov iterations with $\mathbf{Z}_p$ as a starting block. Thus it suffices to show that Block Krylov iterations (step 3 of Algorithm 4) with starting block $\mathbf{Z}_p$ produces some matrix $\tilde{\mathbf{Z}}_p \in \mathbb{R}^{n \times p}$ with orthonormal columns such that $\sqrt{\tilde{\mathbf{z}}_1^T \mathbf{A} \mathbf{A}^T \tilde{\mathbf{z}}_1} \geq (1 - \epsilon)\sigma_1(\mathbf{A})$ where $\tilde{\mathbf{z}}_1$ is the first columns of $\tilde{\mathbf{Z}}_p$.

To show this, we will be using Theorem 13 of [MM15] which bounds the number of iterations of randomized Block Krylov iterations in terms of the singular value gap. To apply this in our setting, we need to show that $\mathbf{Z}_p$ is a good enough starting block. Specifically, let $f(\mathbf{A})$ be the amplifying polynomial used in the proof of Theorem 13 of [MM15]. Let $\mathbf{Y}$ be an orthonormal basis for the span of $f(\mathbf{A})\mathbf{Z}_p$. Then, following the proof in [MM15], to prove the gap-dependent convergence bound, we just need to show that

$$\|f(\mathbf{A}) - \mathbf{Y}\mathbf{Y}^T f(\mathbf{A})\|_F^2 \leq cn^2 \|f(\mathbf{A}) - (f(\mathbf{A}))_p\|_F^2$$

where $(f(\mathbf{A}))_p$ is the best rank p approximation to $f(\mathbf{A})$ and c is a constant. To prove this, we will be proceeding in a similar manner as to the proof of Lemma 4 of [MM15]. Using Lemma 8, we have:

$$\|f(\mathbf{A}) - \mathbf{Y}\mathbf{Y}^T f(\mathbf{A})\|_F^2 \leq \|f(\mathbf{A}) - (f(\mathbf{A}))_p\|_F^2 + \|(f(\mathbf{A}) - (f(\mathbf{A}))_p)\mathbf{Z}_p(\mathbf{U}_p^T \mathbf{Z}_p)^+\|_F^2$$

where $\mathbf{U}_p \in \mathbb{R}^{n \times p}$ contains the top p singular vectors of $\mathbf{A}$ as its columns (note that $\mathbf{A}$ is symmetric so it has same left and right singular vectors). Thus, it suffices to show that

$$\|(f(\mathbf{A}) - (f(\mathbf{A}))_p)\mathbf{Z}_p(\mathbf{U}_p^T \mathbf{Z}_p)^+\|_F^2 \leq cn^2 \|f(\mathbf{A}) - (f(\mathbf{A}))_p\|_F^2$$

for some constant c . Using the spectral submultiplicativity property, we have

$$\begin{aligned} \|(f(\mathbf{A}) - (f(\mathbf{A}))_p)\mathbf{Z}_p(\mathbf{U}_p^T \mathbf{Z}_p)^+\|_F^2 &\leq \|(f(\mathbf{A}) - (f(\mathbf{A}))_p)\|_F^2 \|\mathbf{Z}_p(\mathbf{U}_p^T \mathbf{Z}_p)^+\|_2^2 \\ &\leq \|f(\mathbf{A}) - (f(\mathbf{A}))_p\|_F^2 \|(\mathbf{U}_p^T \mathbf{Z}_p)^+\|_2^2, \end{aligned}$$

where the second step uses the fact that $\|\mathbf{Z}_p\|_2^2 = 1$. From Lemma 9, we have $\|(\mathbf{U}_p^T \mathbf{Z}_p)^+\|_2 \leq c_1 n$ for some constant c_1 . Thus, we have shown that doing Block Krylov iterations with starting block $\mathbf{Z}_p$ gives us the same guarantees as Theorem 13 of [MM15] for gap-dependent convergence of randomized Block Krylov.

Specifically, we can apply the per-vector guarantee (equation 3 of [MM15]) i.e. if $\tilde{\mathbf{z}}_1$ is the first column of the matrix $\tilde{\mathbf{Z}}_p$ after doing Block Krylov iterations, and $\mathbf{v}_1$ is the singular vector of $\mathbf{A}$ corresponding to its largest singular value, we have $|\mathbf{v}_1^T \mathbf{A} \mathbf{A}^T \mathbf{v}_1 - \tilde{\mathbf{z}}_1^T \mathbf{A} \mathbf{A}^T \tilde{\mathbf{z}}_1| \leq \epsilon \sigma_{p+1}^2(\mathbf{A}) \leq \epsilon \sigma_1^2(\mathbf{A})$. This implies that $\sqrt{\tilde{\mathbf{z}}_1^T \mathbf{A} \mathbf{A}^T \tilde{\mathbf{z}}_1} \geq (1 - \epsilon)\sigma_1(\mathbf{A})$ (as $\mathbf{v}_1^T \mathbf{A} \mathbf{A}^T \mathbf{v}_1 = \sigma_1^2(\mathbf{A})$).

Runtime Analysis. The time to compute $\mathbf{Z}$ from $\tilde{\mathbf{A}}$ in step 2 of Algorithm 4 with $k = \frac{2}{\alpha}$ and error parameter $\tilde{\epsilon} = c\alpha^4$ (for some small $c < 1$) is $\tilde{O}\left(\frac{n^2}{\alpha^{29}}\right)$, since we only need to estimate the top $\frac{2}{\alpha}$ singular values, and estimating each singular value using power iterations in Algorithm 3 takes time $\tilde{O}\left(\frac{kn^2 \log(n/\tilde{\epsilon})}{\tilde{\epsilon}^T}\right)$ as described in the analysis of Theorem 8.

The number of iterations of the Block Krylov with starting block $\mathbf{Z}_p$ is given by $O\left(\log(n/\epsilon)/\sqrt{\min(1, \frac{\sigma_1(\mathbf{A})}{\sigma_{p+1}(\mathbf{A})} - 1)}\right)$ as specified in Theorem 13 of [MM15]. Thus, we first need to bound $\frac{\sigma_1(\mathbf{A})}{\sigma_{p+1}(\mathbf{A})}$. We have:

$$\frac{\sigma_1^2(\mathbf{A})}{\sigma_{p+1}^2(\mathbf{A})} \geq \frac{\sigma_p^2(\mathbf{A})}{\sigma_{p+1}^2(\mathbf{A})} \geq \frac{\sigma_p^2(\mathbf{A})}{\sigma_p^2(\mathbf{A}) - \frac{3}{8}\alpha^3 \cdot (\max(n, \|\mathbf{A}\|_1))^2} \geq \frac{1}{1 - \frac{3}{8}\alpha^3} \geq 1 + \frac{3}{8}\alpha^3,$$

where the second inequality uses (27) and for the third inequality we use the fact that $\sigma_p(\mathbf{A}) \leq \max(n, \|\mathbf{A}\|_1)$. Taking square root we have $\frac{\sigma_1(\mathbf{A})}{\sigma_{p+1}(\mathbf{A})} \geq 1 + c\alpha^3$, where c is a large enough constant. Thus, the number of iterations of Block Krylov in step 3 of Algorithm 4 is $O\left(\frac{\log(n/\epsilon)}{\alpha^{3/2}}\right)$. From Theorem 7 of [MM15], for q iterations, Block Krylov takes time $O(n^2 k q + n k^2 q^2 + k^3 q^3)$ where $k = \frac{2}{\alpha}$. Thus, the total time taken by step 3 of Algorithm 4 is $\tilde{O}\left(\frac{n^2 \log(n/\epsilon)}{\alpha^{5/2}} + \frac{n \log^2(n/\epsilon)}{\alpha^5} + \frac{\log^3(n/\epsilon)}{\alpha^{15/2}}\right)$. Since $\tilde{O}\left(\frac{n^2 \log(n/\epsilon)}{\alpha^{29}}\right)$ asymptotically dominates each term of $\tilde{O}\left(\frac{n^2 \log(n/\epsilon)}{\alpha^{5/2}} + \frac{n \log^2(n/\epsilon)}{\alpha^5} + \frac{\log^3(n/\epsilon)}{\alpha^{15/2}}\right)$, the total time taken by the Algorithm 4 is $\tilde{O}\left(\frac{n^2 \log(n/\epsilon)}{\alpha^{29}}\right)$. $\square$

Remark. Since we can apply the per-vector guarantee (Equation 3 of [MM15]) for $\tilde{\mathbf{Z}}_p$, i.e., the output of Block Krylov in Algorithm 4, then for all $p \in \lceil \frac{2}{\alpha} \rceil$ such that the condition in (27) holds and any $i \leq p$, $|\mathbf{v}_i^T \mathbf{A} \mathbf{A}^T \mathbf{v}_i - \tilde{\mathbf{z}}_i^T \mathbf{A} \mathbf{A}^T \tilde{\mathbf{z}}_i| \leq \epsilon \sigma_{p+1}^2(\mathbf{A}) \leq \epsilon \sigma_i^2(\mathbf{A})$, where $\tilde{\mathbf{z}}_i$ is the i^{th} column of $\tilde{\mathbf{Z}}_p$. This implies for all $i \leq p$ we have $\sqrt{\tilde{\mathbf{z}}_i^T \mathbf{A} \mathbf{A}^T \tilde{\mathbf{z}}_i} \geq (1 - \epsilon) \sigma_i(\mathbf{A})$ (as $\mathbf{v}_i^T \mathbf{A} \mathbf{A}^T \mathbf{v}_i = \sigma_i^2(\mathbf{A})$). Thus we are able to approximate the top p singular values and corresponding singular vectors of $\mathbf{A}$ with high accuracy using Algorithm 4, where the p^{th} singular value of $\mathbf{A}$ satisfies the condition of (27).

7 Conclusion

Our work has shown that it is possible to deterministically construct an entrywise sampling matrix $\mathbf{S}$ (a *universal sparsifier*) with just $\tilde{O}(n/\epsilon^c)$ non-zero entries such that for any bounded entry $\mathbf{A}$ with $\|\mathbf{A}\|_\infty \leq 1$, $\|\mathbf{A} - \mathbf{A} \circ \mathbf{S}\|_2 \leq \epsilon \cdot \max(n, \|\mathbf{A}\|_1)$. We show how to achieve sparsity $O(n/\epsilon^2)$ when $\mathbf{A}$ is restricted to be PSD (Theorem 1) and $O(n/\epsilon^4)$ when $\mathbf{A}$ is general (Theorem 4), and prove that both these bounds are tight up to logarithmic factors (Theorems 3 and 6). Further, our proofs are based on simple reductions, which show that any $\mathbf{S}$ that spectrally approximates the all ones matrix to sufficient accuracy (i.e., is a sufficiently good spectral expander) yields a universal sparsifier.

We apply our universal sparsification bounds to give the first $o(n^\omega)$ time deterministic algorithms for several core linear algebraic problems, including singular value/vector approximation and positive semidefiniteness testing. When $\mathbf{A}$ is restricted to be PSD and have entries in $\{-1, 0, 1\}$, we show how to give achieve improved deterministic query complexity of $\tilde{O}(n/\epsilon)$ to construct a general spectral approximation $\tilde{\mathbf{A}}$, which may not be sparse (Theorem 9). We again show that this bound is tight up to a logarithmic factor (Theorem 10).

Our work leaves several open questions:

1. An interesting question is if $\tilde{O}(n/\epsilon)$ sample complexity can be achieved for deterministic spectral approximation of any bounded entry PSD matrix if the sparsity of the output is not restricted, thereby generalizing the upper bound for PSD $\{-1, 0, 1\}$ -matrices proven in Theorem 9. This query complexity is known for randomized algorithms based on column sampling [MM17], however it is not currently known how to derandomize such results.
2. It would also be interesting to close the $\tilde{O}(1/\epsilon^2)$ factor gap between our universal sparsification upper bound of $\tilde{O}(n/\epsilon^4)$ queries for achieving $\epsilon \cdot \max(n, \|\mathbf{A}\|_1)$ spectral approximation error for non-PSD matrices (Theorem 4) and our $\Omega(n/\epsilon^2)$ query lower bound for general deterministic algorithms that make possibly adaptive queries to $\mathbf{A}$ (Theorem 7). By Theorem 6, our universal sparsification bound is tight up to log factors for algorithms that make *non-adaptive* deterministic queries to $\mathbf{A}$. It is unknown if adaptive queries can be used to give improved algorithms.

3. Finally, it would be interesting to understand if our deterministic algorithms for spectrum approximation can be improved. For example, can one compute an ϵn additive error or a $(1 + \epsilon)$ relative error approximation to the top singular value $\|\mathbf{A}\|_2$ for bounded entry $\mathbf{A}$ and constant ϵ in $o(n^\omega)$ time deterministically? Are there fundamental barriers that make doing so difficult?

Acknowledgements

We thank Christopher Musco for helpful conversations about this work. RB, CM and AR was partially supported by an Adobe Research grant, along with NSF Grants 2046235, 1763618, and 1934846. AR was also partially supported by the Manning College of Information and Computer Sciences Dissertation Writing Fellowship. GD was partially supported by NSF AF 1814041, NSF FRG 1760353, and DOE-SC0022085. SS was supported by an NSERC Discovery Grant RGPIN-2018-06398 and a Sloan Research Fellowship. DW was supported by a Simons Investigator Award.

References

- [ACK⁺16] Alexandr Andoni, Jiecao Chen, Robert Krauthgamer, Bo Qin, David P Woodruff, and Qin Zhang. On sketching quadratic forms. In *Proceedings of the 7th Conference on Innovations in Theoretical Computer Science (ITCS)*, 2016.
- [AHK05] Sanjeev Arora, Elad Hazan, and Satyen Kale. Fast algorithms for approximate semidefinite programming using the multiplicative weights update method. In *Proceedings of the 46th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 2005.
- [AHK06] Sanjeev Arora, Elad Hazan, and Satyen Kale. A fast random sampling algorithm for sparsifying matrices. In *Proceedings of the 10th International Workshop on Randomization and Computation (RANDOM)*, 2006.
- [Aig95] Martin Aigner. Turán’s graph theorem. *The American Mathematical Monthly*, 1995.
- [AKL13] Dimitris Achlioptas, Zohar Shay Karnin, and Edo Liberty. Near-optimal entrywise sampling for data matrices. In *Advances in Neural Information Processing Systems 26 (NeurIPS)*, 2013.
- [AKM⁺17] Haim Avron, Michael Kapralov, Cameron Musco, Christopher Musco, Ameya Velingker, and Amir Zandieh. Random Fourier features for kernel ridge regression: Approximation bounds and statistical guarantees. In *Proceedings of the 34th International Conference on Machine Learning (ICML)*, 2017.
- [Alo86] Noga Alon. Eigenvalues and expanders. *Combinatorica*, 1986.
- [Alo21] Noga Alon. Explicit expanders of every degree and size. *Combinatorica*, 2021.
- [AM07] Dimitris Achlioptas and Frank McSherry. Fast computation of low rank matrix approximations. In *Proceedings of the 39th Annual ACM Symposium on Theory of Computing (STOC)*, 2007.
- [AW21] Josh Alman and Virginia Vassilevska Williams. A refined laser method and faster matrix multiplication. In *Proceedings of the 32nd Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, 2021.

- [AZL16] Zeyuan Allen-Zhu and Yuanzhi Li. LazySVD: Even faster SVD decomposition yet without agonizing pain. *Advances in Neural Information Processing Systems 29 (NeurIPS)*, 2016.
- [BCJ20] Ainesh Bakshi, Nadiia Chepurko, and Rajesh Jayaram. Testing positive semi-definiteness via random submatrices. In *Proceedings of the 61st Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 2020.
- [BDD⁺23] Rajarshi Bhattacharjee, Gregory Dexter, Petros Drineas, Cameron Musco, and Archan Ray. Sublinear time eigenvalue approximation via random sampling. *Proceedings of the 50th International Colloquium on Automata, Languages and Programming (ICALP)*, 2023.
- [BDMI14] Christos Boutsidis, Petros Drineas, and Malik Magdon-Ismael. Near-optimal column-based matrix reconstruction. *SIAM Journal on Computing*, 2014.
- [BGPW13] Mark Braverman, Ankit Garg, Denis Pankratov, and Omri Weinstein. Information lower bounds via self-reducibility. In *Proceedings of the 8th International Computer Science Symposium in Russia (CSR)*, 2013.
- [Bha13] Rajendra Bhatia. *Matrix analysis*. Springer Science & Business Media, 2013.
- [BJKS04] Ziv Bar-Yossef, T. S. Jayram, Ravi Kumar, and D. Sivakumar. An information statistics approach to data stream and communication complexity. *Journal of Computer and System Sciences*, 2004.
- [BKKS21] Vladimir Braverman, Robert Krauthgamer, Aditya R Krishnan, and Shay Sapir. Near-optimal entrywise sampling of numerically sparse matrices. In *Proceedings of the 34th Annual Conference on Computational Learning Theory (COLT)*, 2021.
- [BKM22] Vladimir Braverman, Aditya Krishnan, and Christopher Musco. Sublinear time spectral density estimation. In *Proceedings of the 54th Annual ACM Symposium on Theory of Computing (STOC)*, 2022.
- [BSST13] Joshua Batson, Daniel A Spielman, Nikhil Srivastava, and Shang-Hua Teng. Spectral sparsification of graphs: theory and algorithms. *Communications of the ACM*, 2013.
- [CG02] Fan Chung and Ronald Graham. Sparse quasi-random graphs. *Combinatorica*, 2002.
- [CGL⁺20] Julia Chuzhoy, Yu Gao, Jason Li, Danupon Nanongkai, Richard Peng, and Thatchaphol Saranurak. A deterministic algorithm for balanced cut with applications to dynamic connectivity, flows, and beyond. In *Proceedings of the 61st Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 2020.
- [Cha11] Françoise Chatelin. *Spectral approximation of linear operators*. SIAM, 2011.
- [Coh16] Michael B Cohen. Ramanujan graphs in polynomial time. In *Proceedings of the 57th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 2016.
- [Cov99] Thomas M Cover. *Elements of information theory*. John Wiley & Sons, 1999.
- [CR12a] Emmanuel J. Candès and Benjamin Recht. Exact matrix completion via convex optimization. *Communications of the ACM*, 2012.

- [CR12b] Amit Chakrabarti and Oded Regev. An optimal lower bound on the communication complexity of gap-Hamming-distance. *SIAM Journal on Computing*, 2012.
- [CT10] Emmanuel J. Candès and Terence Tao. The power of convex relaxation: near-optimal matrix completion. *IEEE Transactions on Information Theory*, 2010.
- [CW17] Kenneth L Clarkson and David P Woodruff. Low-rank PSD approximation in input-sparsity time. In *Proceedings of the 28th Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, 2017.
- [d’A11] Alexandre d’Aspremont. Subsampling algorithms for semidefinite programming. *Stochastic Systems*, 2011.
- [DZ11] Petros Drineas and Anastasios Zouzias. A note on element-wise matrix sparsification via a matrix-valued Bernstein inequality. *Information Processing Letters*, 2011.
- [Ger31] Semyon Aranovich Gershgorin. Über die abgrenzung der eigenwerte einer matrix. *Izvestiya Rossiyskoy akademii nauk. Seriya matematicheskaya*, 1931.
- [Kun15] Abhisek Kundu. *Element-wise matrix sparsification and reconstruction*. PhD thesis, Rensselaer Polytechnic Institute, USA, 2015.
- [KW92] Jacek Kuczyński and Henryk Woźniakowski. Estimating the largest eigenvalue by the power and Lanczos algorithms with a random start. *SIAM Journal on Matrix Analysis and Applications*, 1992.
- [LPS88] Alexander Lubotzky, Ralph Phillips, and Peter Sarnak. Ramanujan graphs. *Combinatorica*, 1988.
- [LSY16] Lin Lin, Yousef Saad, and Chao Yang. Approximating spectral densities of large matrices. *SIAM Review*, 2016.
- [Mar73] Grigorii Aleksandrovich Margulis. Explicit constructions of concentrators. *Problemy Peredachi Informatsii*, 1973.
- [MM15] Cameron Musco and Christopher Musco. Randomized block Krylov methods for stronger and faster approximate singular value decomposition. *Advances in Neural Information Processing Systems 28 (NeurIPS)*, 2015.
- [MM17] Cameron Musco and Christopher Musco. Recursive sampling for the Nyström method. *Advances in Neural Information Processing Systems 30 (NeurIPS)*, 2017.
- [MMMWW21] Raphael A Meyer, Cameron Musco, Christopher Musco, and David P Woodruff. Hutch++: Optimal stochastic trace estimation. In *Symposium on Simplicity in Algorithms (SOSA)*. SIAM, 2021.
- [Mor94] Moshe Morgenstern. Existence and explicit constructions of $q + 1$ regular Ramanujan graphs for every prime power q . *Journal of Combinatorial Theory, Series B*, 1994.
- [MW17] Cameron Musco and David P. Woodruff. Sublinear time low-rank approximation of positive semidefinite matrices. In *Proceedings of the 58th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 2017.

- [NSW22] Deanna Needell, William Swartworth, and David P. Woodruff. Testing positive semidefiniteness using linear measurements. *Proceedings of the 63rd Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 2022.
- [Pip87] Nicholas Pippenger. Sorting and selecting in rounds. *SIAM Journal on Computing*, 1987.
- [Rou16] Tim Roughgarden. Communication complexity (for algorithm designers). *Foundations and Trends in Theoretical Computer Science*, 2016.
- [Saa11] Yousef Saad. *Numerical methods for large eigenvalue problems: Revised edition*. SIAM, 2011.
- [SKO21] Florian Schäfer, Matthias Katzfuss, and Houman Owhadi. Sparse Cholesky factorization by Kullback–Leibler minimization. *SIAM Journal on Scientific Computing*, 2021.
- [SS08] Daniel A Spielman and Nikhil Srivastava. Graph sparsification by effective resistances. In *Proceedings of the 40th Annual ACM Symposium on Theory of Computing (STOC)*, 2008.
- [ST04] Daniel A. Spielman and Shang-Hua Teng. Nearly-linear time algorithms for graph partitioning, graph sparsification, and solving linear systems. In *Proceedings of the 36th Annual ACM Symposium on Theory of Computing (STOC)*, 2004.
- [SWXZ22] Vaidehi Srinivas, David P Woodruff, Ziyu Xu, and Samson Zhou. Memory bounds for the experts problem. *Proceedings of the 54th Annual ACM Symposium on Theory of Computing (STOC)*, 2022.
- [Tur41] Pál Turán. On an extremal problem in graph theory. *Mat. Fiz. Lapok*, 1941.
- [TYUC17] JA Tropp, A Yurtsever, M Udell, and V Cevher. Randomized single-view algorithms for low-rank matrix approximation, 2017.
- [Ver18] Roman Vershynin. *High-dimensional probability: An introduction with applications in data science*. Cambridge University Press, 2018.
- [Wey12] Hermann Weyl. The asymptotic distribution law of the eigenvalues of linear partial differential equations (with an application to the theory of cavity radiation). *Mathematical Annals*, 1912.
- [WLZ16] Shusen Wang, Luo Luo, and Zhihua Zhang. SPSD matrix approximation via column selection: Theories, algorithms, and extensions. *The Journal of Machine Learning Research*, 2016.
- [Woo14] David P. Woodruff. Sketching as a tool for numerical linear algebra. *Foundations and Trends in Theoretical Computer Science*, 2014.
- [WS00] Christopher Williams and Matthias Seeger. Using the Nyström method to speed up kernel machines. *Advances in Neural Information Processing Systems 13 (NeurIPS)*, 2000.

- [WS23] David Woodruff and William Swartworth. Optimal eigenvalue approximation via sketching. In *Proceedings of the 55th Annual ACM Symposium on Theory of Computing (STOC)*, 2023.
- [WWAF06] Alexander Weisse, Gerhard Wellein, Andreas Alvermann, and Holger Fehske. The kernel polynomial method. *Reviews of modern physics*, 2006.
- [WZ93] Avi Wigderson and David Zuckerman. Expanders that beat the eigenvalue bound: Explicit construction and applications. In *Proceedings of the 25th Annual ACM Symposium on Theory of Computing (STOC)*, 1993.
- [XG10] Jianlin Xia and Ming Gu. Robust approximate Cholesky factorization of rank-structured symmetric positive definite matrices. *SIAM Journal on Matrix Analysis and Applications*, 2010.