

The seal of the California Institute of Technology is a circular emblem. It features a central torch held by a hand, with a flame at the top. The words "CALIFORNIA INSTITUTE OF TECHNOLOGY" are written in a circular path around the torch. The year "1891" is inscribed below the torch.

**ERROR BOUNDS FOR RANDOM
MATRIX APPROXIMATION SCHEMES**

ALEX GITTENS AND JOEL A. TROPP

Technical Report No. 2014-01
August 2014

ERROR BOUNDS FOR RANDOM MATRIX APPROXIMATION SCHEMES

A. GITTENS AND J. A. TROPP

ABSTRACT. Randomized matrix sparsification has proven to be a fruitful technique for producing faster algorithms in applications ranging from graph partitioning to semidefinite programming. In the decade or so of research into this technique, the focus has been—with few exceptions—on ensuring the quality of approximation in the spectral and Frobenius norms. For certain graph algorithms, however, the $\infty \rightarrow 1$ norm may be a more natural measure of performance.

This paper addresses the problem of approximating a real matrix $\mathbf{A}$ by a sparse random matrix $\mathbf{X}$ with respect to several norms. It provides the first results on approximation error in the $\infty \rightarrow 1$ and $\infty \rightarrow 2$ norms, and it uses a result of Latała to study approximation error in the spectral norm. These bounds hold for a reasonable family of random sparsification schemes, those which ensure that the entries of $\mathbf{X}$ are independent and average to the corresponding entries of $\mathbf{A}$. Optimality of the $\infty \rightarrow 1$ and $\infty \rightarrow 2$ error estimates is established. Concentration results for the three norms hold when the entries of $\mathbf{X}$ are uniformly bounded. The spectral error bound is used to predict the performance of several sparsification and quantization schemes that have appeared in the literature; the results are competitive with the performance guarantees given by earlier scheme-specific analyses.

1. INTRODUCTION

Massive datasets are ubiquitous in modern data processing. Classical dense matrix algorithms are poorly suited to such problems because their running times scale superlinearly with the size of the matrix. When the dataset is sparse, one prefers to use sparse matrix algorithms, whose running times depend more on the sparsity of the matrix than on the size of the matrix. Of course, in many applications the matrix is *not* sparse. Accordingly, one may wonder whether it is possible to approximate a computation on a large dense matrix with a related computation on a sparse approximant to the matrix.

Let $\|\cdot\|$ be a norm on matrices. We may phrase the following question: Given a matrix $\mathbf{A}$, how can one efficiently generate a sparse matrix $\mathbf{X}$ for which the approximation error $\|\mathbf{A} - \mathbf{X}\|$ is small?

In the seminal papers [AM07, AM01], Achlioptas and McSherry demonstrate that one can bound, *a priori*, the spectral and Frobenius norm errors incurred when using a particular random approximation scheme. They use this scheme as the basis of an efficient algorithm for calculating near optimal low-rank approximations to large matrices. In a related work [AHK06], Arora, Hazan, and Kale present another randomized scheme for computing sparse approximants with controlled Frobenius and spectral norm errors.

The literature has concentrated on the behavior of the approximation error in the spectral and Frobenius norms; however, these norms are not always the most natural choice. For instance, the problem of graph sparsification is naturally posed as a question of preserving the cut-norm of a graph Laplacian. The strong equivalency of the cut-norm and the $\infty \rightarrow 1$ norm suggests that, for graph-theoretic applications, it may be fruitful to consider the behavior of the $\infty \rightarrow 1$ norm under sparsification. In other applications, e.g., the column subset selection algorithm in [Tro09], the $\infty \rightarrow 2$ norm is the norm of interest.

This paper investigates the errors incurred by approximating a fixed real matrix with a random matrix. Our results apply to any scheme in which the entries of the approximating matrix are independent and average to the corresponding entries of the fixed matrix. Our main contribution is a bound on the expected $\infty \rightarrow p$ norm error, which we specialize to the case of the $\infty \rightarrow 1$ and $\infty \rightarrow 2$

norms. We also use a result of Latała [Lat05] to find an optimal bound on the expected spectral approximation error, and we establish the subgaussianity of the spectral approximation error.

1.1. Summary of results. Consider the matrix $\mathbf{A}$ as an operator from ℓ_∞^n to ℓ_p^m , where $1 \leq p \leq \infty$. The operator norm of $\mathbf{A}$ is defined as

$$\|\mathbf{A}\|_{\infty \rightarrow p} = \max_{\mathbf{u} \neq \mathbf{0}} \frac{\|\mathbf{A}\mathbf{u}\|_p}{\|\mathbf{u}\|_\infty}.$$

The major novel contributions of this paper are estimates of the $\infty \rightarrow 1$ and $\infty \rightarrow 2$ norm approximation errors induced by random matrix sparsification schemes. Let $\mathbf{A}$ be a target matrix. Suppose that $\mathbf{X}$ is a random matrix with independent entries that satisfies $\mathbb{E}\mathbf{X} = \mathbf{A}$. Then

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1} \leq 2 \left[\sum_k \left(\sum_j \text{Var}(X_{jk}) \right)^{1/2} + \sum_j \left(\sum_k \text{Var}(X_{jk}) \right)^{1/2} \right]$$

and

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 2} \leq 2 \left(\sum_{jk} \text{Var}(X_{jk}) \right)^{1/2} + 2\sqrt{m} \min_{\mathbf{D}} \max_j \left(\sum_k \frac{\text{Var}(X_{jk})}{d_k^2} \right)^{1/2}$$

where $\mathbf{D}$ is a positive diagonal matrix with $\text{trace}(\mathbf{D}^2) = 1$. We complement these estimates with tail bounds for the approximation error in each norm.

We also note that

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\| \leq C \left[\max_j \left(\sum_k \text{Var}(X_{jk}) \right)^{1/2} + \max_k \left(\sum_j \text{Var}(X_{jk}) \right)^{1/2} + \left(\sum_{jk} \mathbb{E}(X_{jk} - a_{jk})^4 \right)^{1/4} \right]$$

where C is a universal constant. This estimate follows from Latała's work [Lat05]. We use our spectral norm results to analyze the sparsification and quantization schemes in [AHK06] and [AM07], and we show that our analysis yields error estimates competitive with those established in the respective works.

While on the path to proving the $\infty \rightarrow 1$ and $\infty \rightarrow 2$ norm approximation error estimates, we derive analogous relations of independent interest that bound the expected norm of a random matrix $\mathbf{Z}$ with independent, zero-mean entries:

$$\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow 1} \leq 2\mathbb{E} (\|\mathbf{Z}\|_{\text{col}} + \|\mathbf{Z}^T\|_{\text{col}}),$$

where $\|\mathbf{A}\|_{\text{col}}$ is the sum of the ℓ_2 norms of the columns of $\mathbf{A}$, and

$$\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow 2} \leq 2\mathbb{E} \|\mathbf{Z}\|_F + 2 \min_{\mathbf{D}} \mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty},$$

where $\mathbf{D}$ is a diagonal positive matrix with $\text{trace}(\mathbf{D}^2) = 1$. More generally,

$$\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow p} \leq 2\mathbb{E} \left\| \sum_k \varepsilon_k \mathbf{z}_k \right\|_p + 2 \max_{\|\mathbf{u}\|_q=1} \mathbb{E} \sum_k \left| \sum_j \varepsilon_j Z_{jk} u_j \right|,$$

where q is the conjugate exponent to p and $\mathbf{z}_k$ is the k th column of $\mathbf{Z}$. Here, $\{\varepsilon_k\}$ is a sequence of independent random variables uniform on $\{\pm 1\}$.

1.2. Outline. Section 2 offers an overview of the related strands of research, and Section 3 introduces the reader to our notations. In Section 4, we establish the foundations for our $\infty \rightarrow 1$ and $\infty \rightarrow 2$ error estimates: an estimate of the expected $\infty \rightarrow p$ norm of a random matrix with independent, zero-mean entries and a tail bound for this quantity. In Sections 5 and 6, these estimates are specialized to the cases $p = 1$ and $p = 2$, respectively, and we establish the optimality of the resulting expressions. The bounds are provided in both their most generic forms and ones more suitable for applications. In Section 7, we use a result due to Latała [Lat05] to find an optimal estimate of the spectral approximation error. A deviation bound is provided for the spectral norm which captures the correct (subgaussian) tail behavior. We show that our estimates applied to the

sparsification schemes in [AHK06] and [AM07] recover performance guarantees comparable with those obtained from the original analyses, under slightly weaker hypotheses.

2. RELATED WORK

In recent years, much attention has been paid to the problem of using sampling methods to approximate linear algebra computations efficiently. Such algorithms find applications in areas like data mining, computational biology, and other areas where the relevant matrices may be too large to fit in RAM, or large enough that the computational requirements of standard algorithms become prohibitive. Here we review the streams of literature that have motivated and influenced this work.

2.1. Randomized sparsification. The seminal research of Frieze, Kannan, and Vempala [FKV98, FKV04] on using random sampling to decrease the running time of linear algebra algorithms focuses on approximations to $\mathbf{A}$ constructed from random subsets of its rows. The subsequent influential works of Drineas, Kannan, and Mahoney [DKM06a, DKM06b, DKM06c] also analyze the performance of Monte Carlo algorithms which sample from the columns or rows of $\mathbf{A}$.

In [AM01], Achlioptas and McSherry advance a different approach: instead of using low-rank approximants, they sample the entries of $\mathbf{A}$ to produce a sparse matrix $\mathbf{X}$ that has a spectral decomposition close to that of $\mathbf{A}$. This transition from row/column sampling to independent sampling of the entries allows them to bring to bear powerful techniques from random matrix theory. In [AM01], their main tool is a result on the concentration of the spectral norm of a random matrix with independent, zero-mean, bounded entries. In a follow-up paper, [AM07], they obtain better estimates by using a sharper concentration result derived from Talagrand's inequality. They point out the importance of using schemes which sparsify a given entry in $\mathbf{A}$ with a probability proportional to its magnitude. Such adaptive sparsification schemes keep the variance of the individual entries small, which tends to keep the error in approximating $\mathbf{A}$ with $\mathbf{X}$ small. Their scheme requires either two passes through the matrix or prior knowledge of an upper bound for the largest magnitude in $\mathbf{A}$.

In [AHK06], Arora, Hazan, and Kale describe a random sparsification algorithm which partially quantizes its inputs and requires only one pass through the matrix. They use an epsilon-net argument and Chernoff bounds to establish that with high probability the resulting approximant has small error and high sparsity.

2.2. Probability in Banach spaces. In [RV07], Rudelson and Vershynin take a different approach to the Monte Carlo methodology for low-rank approximation. They consider $\mathbf{A}$ as a linear operator between finite-dimensional Banach spaces and apply techniques of probability in Banach spaces: decoupling, symmetrization, Slepian's lemma for Rademacher random variables, and a law of large numbers for operator-valued random variables. They show that, if $\mathbf{A}$ can be approximated by any rank- r matrix, then it is possible to obtain an accurate rank- k approximation to $\mathbf{A}$ by sampling $O(r \log r)$ rows of $\mathbf{A}$. Additionally, they quantify the behavior of the $\infty \rightarrow 1$ and $\infty \rightarrow 2$ norms of random submatrices.

Our methods are similar to those of Rudelson and Vershynin in [RV07] in that we consider $\mathbf{A}$ as a linear operator between finite-dimensional Banach spaces and use some of the same tools of probability in Banach spaces. Whereas Rudelson and Vershynin consider the behavior of the norms of random submatrices of $\mathbf{A}$, we consider the behavior of the norms of matrices formed by randomly sparsifying (or quantizing) the entries of $\mathbf{A}$. This yields error bounds applicable to schemes that sparsify or quantize matrices entrywise. Since some graph algorithms depend more on the number of edges in the graph than the number of vertices, such schemes may be useful in developing algorithms for handling large graphs.

2.3. Random sampling of graphs. The bulk of the literature has focused on the behavior of the spectral and Frobenius norms under randomized sparsification, but the $\infty \rightarrow 1$ norm occurs naturally in connection with graph theory. Let us review the relevant ideas.

Consider a weighted simple graph $G = (V, E, \omega)$ with adjacency matrix $\mathbf{A}$ given by

$$a_{jk} = \begin{cases} \omega_{jk}, & (j, k) \in E \\ 0, & \text{otherwise.} \end{cases}$$

A *cut* is a partition of the vertices into two blocks: $V = S \cup \bar{S}$. The *cost* of a cut is the sum of the weights of all edges in E which have one vertex in S and one vertex in $\bar{S}$. Several problems relating to cuts are of considerable practical interest. In particular, the MAXCUT problem, to determine the cut of maximum cost in a graph, is common in computer science applications. The cuts of maximum cost are exactly those which realize the *cut-norm* of the adjacency matrix, which is defined as

$$\|\mathbf{A}\|_C = \max_{S \subseteq V} \left| \sum_{(j,k) \in E} \omega_{jk} (\mathbb{1}_S)_j (\mathbb{1}_{\bar{S}})_k \right|,$$

where $\mathbb{1}_S$ is the indicator vector for S . Finding the cut-norm of a general matrix is NP-hard, but in [AN04], the authors offer a randomized polynomial-time algorithm which finds a submatrix $\tilde{\mathbf{A}}$ of $\mathbf{A}$ such that $|\sum_{jk} \tilde{a}_{jk}| \geq 0.56 \|\mathbf{A}\|_C$. One crucial point in the derivation of the algorithm is the fact that the $\infty \rightarrow 1$ norm is strongly equivalent with the cut-norm:

$$\|\mathbf{A}\|_C \leq \|\mathbf{A}\|_{\infty \rightarrow 1} \leq 4 \|\mathbf{A}\|_C.$$

In his thesis [Kar95] and the sequence of papers [Kar94a, Kar94b, Kar96], Karger introduces the idea of random sampling to increase the efficiency of calculations with graphs, with a focus on cuts. In [Kar96], he shows that by picking each edge of the graph with a probability inversely proportional to the density of edges in a neighborhood of that edge, one can construct a *sparsifier*, i.e., a graph with the same vertex set and significantly fewer edges that preserves the value of each cut to within a factor of $(1 \pm \epsilon)$.

In [SS08], Spielman and Srivastava improve upon this sampling scheme, instead keeping an edge with probability proportional to its *effective resistance*—a measure of how likely it is to appear in a random spanning tree of the graph. They provide an algorithm which produces a sparsifier with $O((n \log n)/\epsilon^2)$ edges, where n is the number of vertices in the graph. They obtain this result by reducing the problem to the behavior of projection matrices Π_G and $\Pi_{G'}$ associated with the original graph and the sparsifier, and appealing to a spectral norm concentration result.

The $\log n$ factor in [SS08] seems to be an unavoidable consequence of using spectral norm concentration. In [BSS09], Batson et. al. prove that the $\log n$ factor is not intrinsic: they establish that every graph has a sparsifier that has $O(n)$ edges. The proof is constructive and provides a deterministic algorithm for constructing such optimal sparsifiers in $O(n^3 m)$ time, where m is the number of edges in the original graph.

The algorithm of [BSS09] is clearly not suitable for sparsifying graphs with a large number of vertices. Part of our motivation for investigating the $\infty \rightarrow 1$ approximation error is the belief that the equivalence of the cut-norm with the $\infty \rightarrow 1$ norm means that matrix sparsification in the $\infty \rightarrow 1$ norm might be useful for efficiently constructing optimal sparsifiers for such graphs.

3. BACKGROUND

We establish the notation used in the sequel and introduce our key technical tools.

All quantities are real. The k th column of the matrix $\mathbf{A}$ is denoted by $\mathbf{a}_k$ and the entries are denoted a_{jk} .

For $1 \leq p \leq \infty$, the ℓ_p norm of $\mathbf{x}$ is written as $\|\mathbf{x}\|_p$. We treat $\mathbf{A}$ as an operator from ℓ_p^n to ℓ_q^m , and the $p \rightarrow q$ operator norm of $\mathbf{A}$ is written as $\|\mathbf{A}\|_{p \rightarrow q}$. The $p \rightarrow q$ and $q' \rightarrow p'$ norms are dual

in the sense that

$$\|\mathbf{A}\|_{p \rightarrow q} = \|\mathbf{A}^T\|_{q' \rightarrow p'}.$$

This paper is concerned primarily with the spectral norm and the $\infty \rightarrow 1$ and $\infty \rightarrow 2$ norms. The spectral norm $\|\mathbf{A}\|$ is the largest singular value of $\mathbf{A}$. The $\infty \rightarrow 1$ and $\infty \rightarrow 2$ norms do not have such nice interpretations, and they are NP-hard to compute for general matrices [Roh00]. We remark that $\|\mathbf{A}\|_{\infty \rightarrow 1} = \|\mathbf{A}\mathbf{x}\|_1$ and $\|\mathbf{A}\|_{\infty \rightarrow 2} = \|\mathbf{A}\mathbf{y}\|_2$ for certain vectors $\mathbf{x}$ and $\mathbf{y}$ whose components take values ± 1 .

An additional operator norm, the $2 \rightarrow \infty$ norm, is also of interest: it is the largest ℓ_2 norm achieved by a row of $\mathbf{A}$. We encounter two norms in the sequel that are not operator norms: the Frobenius norm, denoted by $\|\mathbf{A}\|_F$, and the column norm

$$\|\mathbf{A}\|_{\text{col}} = \sum_k \|\mathbf{a}_k\|_2.$$

The expectation of a random variable X is written $\mathbb{E}X$ and its variance is written $\text{Var}(X) = \mathbb{E}(X - \mathbb{E}X)^2$. The expectation taken with respect to one variable X , with all others fixed, is written $\mathbb{E}_X$. The L_q norm of X is denoted by $\mathbb{E}^q(X) = (\mathbb{E}|X|^q)^{1/q}$.

The expression $X \sim Y$ indicates the random variables X and Y are identically distributed. Given a random variable X , the symbol X' denotes a random variable independent of X such that $X' \sim X$. The indicator variable of the event $X > Y$ is written $\mathbb{1}_{X>Y}$.

The Bernoulli distribution with expectation p is written $\text{Bern}(p)$ and the Binomial distribution of n independent trials each with success probability p is written $\text{Bin}(n, p)$. We write $X \sim \text{Bern}(p)$ to indicate X is Bernoulli.

A Rademacher random variable takes on the values ± 1 with equal probability. A vector whose components are independent Rademacher variables is called a Rademacher vector. A real sum whose terms are weighted by independent Rademacher variables is called a Rademacher sum. The Khintchine inequality [Sza76] gives information on the moments of a Rademacher sum; in particular, it tells us the expected value of the sum is equivalent with the ℓ_2 norm of the vector $\mathbf{x}$:

Proposition 1 (Khintchine inequality). *Let $\mathbf{x}$ be a real vector, and let ε be a Rademacher vector. Then*

$$\frac{1}{\sqrt{2}} \|\mathbf{x}\|_2 \leq \mathbb{E} \left| \sum_k \varepsilon_k x_k \right| \leq \|\mathbf{x}\|_2.$$

4. THE $\infty \rightarrow p$ NORM OF A RANDOM MATRIX

We are interested in schemes that approximate a given matrix $\mathbf{A}$ by means of a random matrix $\mathbf{X}$ in such a way that the entries of $\mathbf{X}$ are independent and $\mathbb{E}\mathbf{X} = \mathbf{A}$. It follows that the error matrix $\mathbf{Z} = \mathbf{A} - \mathbf{X}$ has independent, zero-mean entries. Our intellectual concern is the class of sparse random matrices, but this property does not play a role at this stage of the analysis.

In this section, we derive a bound on the expected value of the $\infty \rightarrow p$ norm of a random matrix with independent, zero-mean entries. We also study the tails of this error. In the next two sections, we use the results of this section to reach more detailed conclusions on the $\infty \rightarrow 1$ and $\infty \rightarrow 2$ norms of $\mathbf{Z}$.

4.1. Expected $\infty \rightarrow p$ norm. The main tool used to derive the bound on the expected norm of $\mathbf{Z}$ is the following symmetrization argument [vW96, Lemma 2.3.1 et seq.].

Proposition 2. *Let $Z_1, \dots, Z_n, Z'_1, \dots, Z'_n$ be independent random variables satisfying $Z_i \sim Z'_i$, and let ε be a Rademacher vector. Let $\mathcal{F}$ be a family of functions such that*

$$\sup_{f \in \mathcal{F}} \sum_{k=1}^n (f(Z_k) - f(Z'_k))$$

is measurable. Then

$$\mathbb{E} \sup_{f \in \mathcal{F}} \sum_{k=1}^n (f(Z_k) - f(Z'_k)) = \mathbb{E} \sup_{f \in \mathcal{F}} \sum_{k=1}^n \varepsilon_k (f(Z_k) - f(Z'_k)).$$

Since we work with finite-dimensional probability models and linear functions, measurability concerns can be ignored.

The other critical tool is a version of Talagrand's Rademacher comparison theorem [LT91, Theorem 4.12 et seq.].

Proposition 3. Fix finite-dimensional vectors $\mathbf{z}_1, \dots, \mathbf{z}_n$ and let ε be a Rademacher vector. Then

$$\mathbb{E} \max_{\|\mathbf{u}\|_q=1} \sum_{k=1}^n \varepsilon_k |\langle \mathbf{z}_k, \mathbf{u} \rangle| \leq \mathbb{E} \max_{\|\mathbf{u}\|_q=1} \sum_{k=1}^n \varepsilon_k \langle \mathbf{z}_k, \mathbf{u} \rangle.$$

Now we state and prove the bound on the expected norm of $\mathbf{Z}$.

Theorem 1. Let $\mathbf{Z}$ be a random matrix with independent, zero-mean entries and let ε be a Rademacher vector independent of $\mathbf{Z}$. Then

$$\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow p} \leq 2\mathbb{E} \left\| \sum_k \varepsilon_k \mathbf{z}_k \right\|_p + 2 \max_{\|\mathbf{u}\|_q=1} \mathbb{E} \sum_k \left| \sum_j \varepsilon_j Z_{jk} u_j \right|$$

where q is the conjugate exponent of p .

Proof of Theorem 1. By duality,

$$\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow p} = \mathbb{E} \|\mathbf{Z}^T\|_{q \rightarrow 1} = \mathbb{E} \max_{\|\mathbf{u}\|_q=1} \sum_k |\langle \mathbf{z}_k, \mathbf{u} \rangle|.$$

Center the terms in the sum and apply subadditivity of the maximum to get

$$\begin{aligned} \mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow p} &\leq \mathbb{E} \max_{\|\mathbf{u}\|_q=1} \sum_k (|\langle \mathbf{z}_k, \mathbf{u} \rangle| - \mathbb{E}' |\langle \mathbf{z}'_k, \mathbf{u} \rangle|) + \max_{\|\mathbf{u}\|_q=1} \mathbb{E} \sum_k |\langle \mathbf{z}_k, \mathbf{u} \rangle| \\ &=: F + S. \end{aligned} \tag{1}$$

Begin with the first term in (1). Use Jensen's inequality to draw the expectation outside of the maximum:

$$F \leq \mathbb{E} \max_{\|\mathbf{u}\|_q=1} \sum_k (|\langle \mathbf{z}_k, \mathbf{u} \rangle| - |\langle \mathbf{z}'_k, \mathbf{u} \rangle|).$$

Now apply Proposition 2 to symmetrize the random variable:

$$F \leq \mathbb{E} \max_{\|\mathbf{u}\|_q=1} \sum_k \varepsilon_k (|\langle \mathbf{z}_k, \mathbf{u} \rangle| - |\langle \mathbf{z}'_k, \mathbf{u} \rangle|).$$

By the subadditivity of the maximum,

$$F \leq \mathbb{E} \left(\max_{\|\mathbf{u}\|_q=1} \sum_k \varepsilon_k |\langle \mathbf{z}_k, \mathbf{u} \rangle| + \max_{\|\mathbf{u}\|_q=1} \sum_k -\varepsilon_k |\langle \mathbf{z}_k, \mathbf{u} \rangle| \right) = 2\mathbb{E} \max_{\|\mathbf{u}\|_q=1} \sum_k \varepsilon_k |\langle \mathbf{z}_k, \mathbf{u} \rangle|,$$

where we have invoked the fact that $-\varepsilon_k$ has the Rademacher distribution. Apply Proposition 3 to get the final estimate of F :

$$F \leq 2\mathbb{E} \max_{\|\mathbf{u}\|_q=1} \sum_k \varepsilon_k \langle \mathbf{z}_k, \mathbf{u} \rangle = 2\mathbb{E} \max_{\|\mathbf{u}\|_q=1} \left\langle \sum_k \varepsilon_k \mathbf{z}_k, \mathbf{u} \right\rangle = 2\mathbb{E} \left\| \sum_k \varepsilon_k \mathbf{z}_k \right\|_p.$$

Now consider the last term in (1). Use Jensen's inequality to prepare for symmetrization:

$$\begin{aligned} S &= \max_{\|\mathbf{u}\|_q=1} \mathbb{E} \sum_k \left| \sum_j Z_{jk} u_j \right| = \max_{\|\mathbf{u}\|_q=1} \mathbb{E} \sum_k \left| \sum_j (Z_{jk} - \mathbb{E}' Z'_{jk}) u_j \right| \\ &\leq \max_{\|\mathbf{u}\|_q=1} \sum_k \mathbb{E} \left| \sum_j (Z_{jk} - Z'_{jk}) u_j \right|. \end{aligned}$$

Apply Proposition 2 to the expectation of the inner sum to see

$$S \leq \max_{\|\mathbf{u}\|_q=1} \sum_k \mathbb{E} \left| \sum_j \varepsilon_j (Z_{jk} - Z'_{jk}) u_j \right|.$$

The triangle inequality gives us the final expression:

$$S \leq \max_{\|\mathbf{u}\|_q=1} 2\mathbb{E} \sum_k \left| \sum_j \varepsilon_j Z_{jk} u_j \right|.$$

Introduce the bounds for F and S into (1) to complete the proof. $\square$

4.2. Tail bound for $\infty \rightarrow p$ norm. We now develop a deviation bound for the $\infty \rightarrow p$ approximation error. The argument is based on a bounded differences inequality.

First we establish some notation. Let $g : \mathbb{R}^n \rightarrow \mathbb{R}$ be a measurable function of n random variables. Let $X_1, \dots, X_n$ be independent random variables, and write $W = g(X_1, \dots, X_n)$. Let W_i denote the random variable obtained by replacing the i th argument of g with an independent copy: $W_i = g(X_1, \dots, X'_i, \dots, X_n)$.

The following bounded differences inequality states that if g is insensitive to changes of a single argument, then W does not deviate much from its mean.

Proposition 4 ([BLM03]). *Let W and $\{W_i\}$ be random variables defined as above. Assume that there exists a positive number C such that, almost surely,*

$$\sum_{i=1}^n (W - W_i)^2 \mathbb{1}_{W > W_i} \leq C.$$

Then, for all $t > 0$,

$$\mathbb{P}(W > \mathbb{E}W + t) \leq e^{-t^2/(4C)}.$$

To apply Proposition 4, we let $\mathbf{Z} = \mathbf{A} - \mathbf{X}$ be our error matrix, $W = \|\mathbf{Z}\|_{\infty \rightarrow p}$, and $W^{jk} = \|\mathbf{Z}^{jk}\|_{\infty \rightarrow p}$, where $\mathbf{Z}^{jk}$ is a matrix obtained by replacing $a_{jk} - X_{jk}$ with an identically distributed variable $a_{jk} - X'_{jk}$ while keeping all other variables fixed. The $\infty \rightarrow p$ norms are sufficiently insensitive to each entry of the matrix that Proposition 4 gives us a useful deviation bound.

Theorem 2. *Fix an $m \times n$ matrix $\mathbf{A}$, and let $\mathbf{X}$ be a random matrix with independent entries for which $\mathbb{E}X = \mathbf{A}$. Assume $|X_{jk}| \leq \frac{D}{2}$ almost surely for all j, k . Then, for all $t > 0$,*

$$\mathbb{P}\left(\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow p} > \mathbb{E}\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow p} + t\right) \leq e^{-t^2/(4D^2nm^s)}$$

where $s = \max\{0, 1 - 2/q\}$ and q is the conjugate exponent to p .

Proof. Let q be the conjugate exponent of p , and choose $\mathbf{u}, \mathbf{v}$ such that $W = \mathbf{u}^T \mathbf{Z} \mathbf{v}$ and $\|\mathbf{u}\|_q = 1$ and $\|\mathbf{v}\|_\infty = 1$. Then

$$(W - W^{jk}) \mathbb{1}_{W > W^{jk}} \leq \mathbf{u}^T (\mathbf{Z} - \mathbf{Z}^{jk}) \mathbf{v} \mathbb{1}_{W > W^{jk}} = (X'_{jk} - X_{jk}) u_j v_k \mathbb{1}_{W > W^{jk}} \leq D |u_j v_k|.$$

This implies

$$\sum_{j,k} (W - W^{jk})^2 \mathbb{1}_{W > W^{jk}} \leq D^2 \sum_{j,k} |u_j v_k|^2 \leq nD^2 \|\mathbf{u}\|_2^2,$$

so we can apply Proposition 4 if we have an estimate for $\|\mathbf{u}\|_2^2$. We have the bounds $\|\mathbf{u}\|_2 \leq \|\mathbf{u}\|_q$ for $q \in [1, 2]$ and $\|\mathbf{u}\|_2 \leq m^{1/2-1/q} \|\mathbf{u}\|_q$ for $q \in [2, \infty]$. Therefore,

$$\sum_{j,k} (W - W^{jk})^2 \mathbb{1}_{W > W^{jk}} \leq D^2 \begin{cases} nm^{1-2/q}, & q \in [2, \infty] \\ n, & q \in [1, 2]. \end{cases}$$

It follows from Proposition 4 that

$$\mathbb{P}\left(\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow p} > \mathbb{E}\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow p} + t\right) = \mathbb{P}(W > \mathbb{E}W + t) \leq e^{-t^2/(4D^2nm^s)}$$

where $s = \max\{0, 1 - 2/q\}$. $\square$

It is often convenient to measure deviations on the scale of the mean. Taking $t = \delta \mathbb{E}\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow p}$ in Theorem 2 gives the following result.

Corollary 1. *Under the conditions of Theorem 2, for all $\delta > 0$,*

$$\mathbb{P}\left(\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow p} > (1 + \delta)\mathbb{E}\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow p}\right) \leq e^{-\delta^2(\mathbb{E}\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow p})^2/(4D^2nm^s)}.$$

5. APPROXIMATION IN THE $\infty \rightarrow 1$ NORM

In this section, we develop the $\infty \rightarrow 1$ error bound as a consequence of Theorem 1. We then prove that one form of the error bound is optimal, and we describe an example of its application to matrix sparsification.

5.1. Expected $\infty \rightarrow 1$ norm. To derive the $\infty \rightarrow 1$ error bound, we first apply Theorem 1 with $p = 1$.

Theorem 3. *Suppose that $\mathbf{Z}$ is a random matrix with independent, zero-mean entries. Then*

$$\mathbb{E}\|\mathbf{Z}\|_{\infty \rightarrow 1} \leq 2\mathbb{E}(\|\mathbf{Z}\|_{\text{col}} + \|\mathbf{Z}^T\|_{\text{col}}).$$

Proof. Apply Theorem 1 to get

$$\begin{aligned} \mathbb{E}\|\mathbf{Z}\|_{\infty \rightarrow 1} &\leq 2\mathbb{E}\left\|\sum_k \varepsilon_k \mathbf{z}_k\right\|_1 + 2 \max_{\|\mathbf{u}\|_\infty=1} \mathbb{E} \sum_k \left|\sum_j \varepsilon_j Z_{jk} u_j\right| \\ &=: F + S. \end{aligned} \tag{2}$$

Use Hölder's inequality to bound the first term in (2) with a sum of squares:

$$\begin{aligned} F &= 2\mathbb{E} \sum_j \left|\sum_k \varepsilon_k Z_{jk}\right| = 2\mathbb{E}_{\mathbf{Z}} \sum_j \mathbb{E}_{\boldsymbol{\varepsilon}} \left|\sum_k \varepsilon_k Z_{jk}\right| \\ &\leq 2\mathbb{E}_{\mathbf{Z}} \sum_j \left(\mathbb{E}_{\boldsymbol{\varepsilon}} \left|\sum_k \varepsilon_k Z_{jk}\right|^2\right)^{1/2}. \end{aligned}$$

The inner expectation can be computed exactly by expanding the square and using the independence of the Rademacher variables:

$$F \leq 2\mathbb{E} \sum_j \left(\sum_k Z_{jk}^2\right)^{1/2} = 2\mathbb{E}\|\mathbf{Z}^T\|_{\text{col}}.$$

We treat the second term in the same manner. Use Hölder's inequality to replace the sum with a sum of squares and invoke the independence of the Rademacher variables to eliminate cross terms:

$$S \leq 2 \max_{\|\mathbf{u}\|_\infty=1} \mathbb{E}_{\mathbf{Z}} \sum_k \left(\mathbb{E}_{\boldsymbol{\varepsilon}} \left|\sum_j \varepsilon_j Z_{jk} u_j\right|^2\right)^{1/2} = 2 \max_{\|\mathbf{u}\|_\infty=1} \mathbb{E} \sum_k \left(\sum_j Z_{jk}^2 u_j^2\right)^{1/2}.$$

Since $\|\mathbf{u}\|_\infty = 1$, it follows that $u_j^2 \leq 1$ for all j , and

$$S \leq 2\mathbb{E} \sum_k \left(\sum_j Z_{jk}^2\right)^{1/2} = 2\mathbb{E}\|\mathbf{Z}\|_{\text{col}}.$$

Introduce these estimates for F and S into (2) to complete the proof. $\square$

Taking $\mathbf{Z} = \mathbf{A} - \mathbf{X}$ in Theorem 3, we find

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1} \leq 2\mathbb{E} \left[\sum_k \left(\sum_j (a_{jk} - X_{jk})^2 \right)^{1/2} + \sum_j \left(\sum_k (a_{jk} - X_{jk})^2 \right)^{1/2} \right].$$

A simple application of Jensen's inequality gives an error bound in terms of the variances of the entries of $\mathbf{X}$.

Corollary 2. *Fix the matrix $\mathbf{A}$, and let $\mathbf{X}$ be a random matrix with independent entries for which $\mathbb{E}X_{jk} = a_{jk}$. Then*

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1} \leq 2 \left[\sum_k \left(\sum_j \text{Var}(X_{jk}) \right)^{1/2} + \sum_j \left(\sum_k \text{Var}(X_{jk}) \right)^{1/2} \right].$$

5.2. Optimality. The bound in Corollary 2 is optimal in the sense that there are families of matrices $\mathbf{A}$ and random approximants $\mathbf{X}$ for which $\mathbb{E} \|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1}$ grows like one of the terms in the bound and dominates the other term in the bound. To show this, we construct specific examples.

Let $\mathbf{A}$ be a tall $m \times \sqrt{m}$ matrix of ones and choose the approximant $X_{jk} \sim 2 \text{ Bern}(\frac{1}{2})$. With this choice, $\text{Var}(X_{jk}) = 1$, so the first term in the bound is m and the second term is $m^{5/4}$. The following argument from [RV07, Sec. 4.2] establishes that $\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1}$ grows like $m^{5/4}$.

Observe that the matrix $\mathbf{A} - \mathbf{X} = [\varepsilon_{jk}]$, where ε_{jk} are i.i.d. Rademacher variables. Its $\infty \rightarrow 1$ norm is

$$\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1} = \max_{\substack{\|\mathbf{y}\|_{\infty} = 1 \\ \|\mathbf{x}\|_{\infty} = 1}} \sum_{j,k} \varepsilon_{jk} x_j y_k.$$

Let δ_k be a sequence of i.i.d. Rademacher variables. By the scalar Khintchine inequality,

$$\mathbb{E}_{\delta} \sum_j \left| \sum_k \varepsilon_{jk} \delta_k \right| \geq \sum_j \frac{1}{\sqrt{2}} \|\varepsilon_{j\cdot}\|_2 = \frac{1}{\sqrt{2}} m^{5/4}.$$

The probabilistic method shows that there is a sign vector $\mathbf{x}$ for which

$$\sum_j \left| \sum_k \varepsilon_{jk} x_k \right| \geq \frac{1}{\sqrt{2}} m^{5/4}.$$

Choose the vector $\mathbf{y}$ with components $y_j = \text{sgn}(\sum_k \varepsilon_{jk} x_k)$. Then

$$\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1} \geq \sum_j \sum_k \varepsilon_{jk} y_j x_k = \sum_j \left| \sum_k \varepsilon_{jk} x_k \right| \geq \frac{1}{\sqrt{2}} m^{5/4}.$$

This shows $\|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1}$ grows like the second term in the error bound, so this term of the bound cannot be ignored.

These arguments, applied to a fat $\sqrt{n} \times n$ matrix of ones, also establish the necessity of the first term.

5.3. Example application. In this section we provide an example illustrating the application of Corollary 2 to matrix sparsification.

From Corollary 2 we infer that a good scheme for sparsifying a matrix $\mathbf{A}$ while minimizing the expected relative $\infty \rightarrow 1$ error is one which drastically increases the sparsity of $\mathbf{X}$ while keeping the relative error

$$\frac{\sum_k \left(\sum_j \text{Var}(X_{jk}) \right)^{1/2} + \sum_j \left(\sum_k \text{Var}(X_{jk}) \right)^{1/2}}{\|\mathbf{A}\|_{\infty \rightarrow 1}}$$

small. Once a sparsification scheme is chosen, the hardest part of estimating this quantity is probably estimating the $\infty \rightarrow 1$ norm of $\mathbf{A}$. The example shows, for a simple family of approximation schemes, what kind of sparsification results can be obtained using Corollary 2 when we have a very good handle on this quantity.

Consider the case where $\mathbf{A}$ is an $n \times n$ matrix whose entries all lie within an interval bounded away from zero; for definiteness, take them to be positive. Let γ be a desired bound on the expected relative $\infty \rightarrow 1$ norm error. We choose the randomization strategy $X_{jk} \sim \frac{a_{jk}}{p} \text{Bern}(p)$ and ask how small can p be without violating our bound on the expected error.

In this case,

$$\|\mathbf{A}\|_{\infty \rightarrow 1} = \sum_{j,k} a_{jk} = O(n^2),$$

and $\text{Var}(X_{jk}) = \frac{a_{jk}^2}{p} - a_{jk}^2$. Consequently, the first term in Corollary 2 satisfies

$$\begin{aligned} \sum_k \left(\sum_j \text{Var}(X_{jk}) \right)^{1/2} &= \sum_k \left(\frac{1}{p} \|\mathbf{a}_k\|_2^2 - \|\mathbf{a}_k\|_2^2 \right)^{1/2} = \left(\frac{1-p}{p} \right)^{1/2} \|\mathbf{A}\|_{\text{col}} \\ &= O\left(\left(\frac{1-p}{p} \right)^{1/2} n\sqrt{n} \right) \end{aligned}$$

and likewise the second term satisfies

$$\sum_j \left(\sum_k \text{Var}(X_{jk}) \right)^{1/2} = O\left(\left(\frac{1-p}{p} \right)^{1/2} n\sqrt{n} \right).$$

Therefore the relative $\infty \rightarrow 1$ norm error satisfies

$$\frac{\sum_k \left(\sum_j \text{Var}(X_{jk}) \right)^{1/2} + \sum_j \left(\sum_k \text{Var}(X_{jk}) \right)^{1/2}}{\|\mathbf{A}\|_{\infty \rightarrow 1}} = O\left(\left(\frac{1-p}{pn} \right)^{1/2} \right).$$

It follows that $\mathbb{E} \|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 1} < \gamma$ for p on the order of $(1 + n\gamma^2)^{-1}$ or larger. The expected number of nonzero entries in $\mathbf{X}$ is pn^2 , so for matrices with this structure, we can sparsify with a relative $\infty \rightarrow 1$ norm error smaller than γ while reducing the number of expected nonzero entries to as few as $O(\frac{n^2}{1+n\gamma^2}) = O(\frac{n}{\gamma^2})$. Intuitively, this sparsification result is optimal in the dimension: it seems we must keep on average at least one entry per row and column if we are to faithfully approximate $\mathbf{A}$.

6. APPROXIMATION IN THE $\infty \rightarrow 2$ NORM

In this section, we develop the $\infty \rightarrow 2$ error bound stated in the introduction, establish the optimality of a related bound, and provide examples of its application to matrix sparsification. To derive the error bound, we first specialize Theorem 1 to the case of $p = 2$.

Theorem 4. *Suppose that $\mathbf{Z}$ is a random matrix with independent, zero-mean entries. Then*

$$\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow 2} \leq 2\mathbb{E} \|\mathbf{Z}\|_{\text{F}} + 2 \min_{\mathbf{D}} \mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty}$$

where $\mathbf{D}$ is a positive diagonal matrix that satisfies $\text{trace}(\mathbf{D}^2) = 1$.

Proof. Apply Theorem 1 to get

$$\begin{aligned} \mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow 2} &\leq 2\mathbb{E} \left\| \sum_k \varepsilon_k \mathbf{z}_k \right\|_2 + 2 \max_{\|\mathbf{u}\|_2=1} \mathbb{E} \sum_k \left| \sum_j \varepsilon_j Z_{jk} u_j \right| \\ &=: F + S. \end{aligned} \tag{3}$$

Expand the first term, and use Jensen's inequality to move the expectation with respect to the Rademacher variables inside the square root:

$$F = 2\mathbb{E} \left(\sum_j \left| \sum_k \varepsilon_k Z_{jk} \right|^2 \right)^{1/2} \leq 2\mathbb{E}_{\mathbf{Z}} \left(\sum_j \mathbb{E}_{\varepsilon} \left| \sum_k \varepsilon_k Z_{jk} \right|^2 \right)^{1/2}.$$

The independence of the Rademacher variables implies that the cross terms cancel, so

$$F \leq 2\mathbb{E} \left(\sum_j \sum_k Z_{jk}^2 \right)^{1/2} = 2\mathbb{E} \|\mathbf{Z}\|_F.$$

We use the Cauchy–Schwarz inequality to replace the ℓ_1 norm with an ℓ_2 norm in the second term of (3). A direct application would introduce a possibly suboptimal factor of $\sqrt{n}$ (where n is the number of columns in $\mathbf{Z}$), so instead we choose $d_k > 0$ such that $\sum_k d_k^2 = 1$ and use the corresponding weighted ℓ_2 norm:

$$S = 2 \max_{\|\mathbf{u}\|_2=1} \mathbb{E} \sum_k \frac{\left| \sum_j \varepsilon_j Z_{jk} u_j \right|}{d_k} d_k \leq 2 \max_{\|\mathbf{u}\|_2=1} \mathbb{E} \left(\sum_k \frac{\left| \sum_j \varepsilon_j Z_{jk} u_j \right|^2}{d_k^2} \right)^{1/2}.$$

Move the expectation with respect to the Rademacher variables inside the square root and observe that the cross terms cancel:

$$S \leq 2 \max_{\|\mathbf{u}\|_2=1} \mathbb{E}_{\mathbf{Z}} \left(\sum_k \frac{\mathbb{E}_{\varepsilon} \left| \sum_j \varepsilon_j Z_{jk} u_j \right|^2}{d_k^2} \right)^{1/2} = 2 \max_{\|\mathbf{u}\|_2=1} \mathbb{E} \left(\sum_{j,k} \frac{Z_{jk}^2 u_j^2}{d_k^2} \right)^{1/2}.$$

Use Jensen's inequality to pass the maximum through the expectation, and note that if $\|\mathbf{u}\|_2 = 1$ then the vector formed by elementwise squaring $\mathbf{u}$ lies on the ℓ_1 unit ball, thus

$$S \leq 2\mathbb{E} \left(\max_{\|\mathbf{u}\|_1=1} \sum_{j,k} \left(\frac{Z_{jk}}{d_k} \right)^2 u_j \right)^{1/2}.$$

Clearly this maximum is achieved when $\mathbf{u}$ is chosen so $u_j = 1$ at an index j for which $\left(\sum_k \left(\frac{Z_{jk}}{d_k} \right)^2 \right)^{1/2}$ is maximal and $u_j = 0$ otherwise. Consequently, the maximum is the largest of the ℓ_2 norms of the rows of $\mathbf{Z}\mathbf{D}^{-1}$, where $\mathbf{D} = \text{diag}(d_1, \dots, d_n)$. Recall that this quantity is, by definition, $\|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty}$. Therefore $S \leq 2\mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty}$. The theorem follows by optimizing our choice of $\mathbf{D}$ and introducing our estimates for F and S into (3). $\square$

Taking $\mathbf{Z} = \mathbf{A} - \mathbf{X}$ in Theorem 4, we have

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 2} \leq 2\mathbb{E} \left(\sum_{j,k} (X_{jk} - a_{jk})^2 \right)^{1/2} + 2 \min_{\mathbf{D}} \mathbb{E} \max_j \left(\sum_k \frac{(X_{jk} - a_{jk})^2}{d_k^2} \right)^{1/2}. \quad (4)$$

We now derive a bound which depends only on the variances of the X_{jk} .

Corollary 3. *Fix the $m \times n$ matrix $\mathbf{A}$ and let $\mathbf{X}$ be a random matrix with independent entries so that $\mathbb{E}\mathbf{X} = \mathbf{A}$. Then*

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\|_{\infty \rightarrow 2} \leq 2 \left(\sum_{j,k} \text{Var}(X_{jk}) \right)^{1/2} + 2\sqrt{m} \min_{\mathbf{D}} \max_j \left(\sum_k \frac{\text{Var}(X_{jk})}{d_k^2} \right)^{1/2}$$

where $\mathbf{D}$ is a positive diagonal matrix with $\text{trace}(\mathbf{D}^2) = 1$.

Proof. Let F and S denote, respectively, the first and second term of (4). An application of Jensen's inequality shows that $F \leq 2 \left(\sum_{j,k} \text{Var}(X_{jk}) \right)^{1/2}$. A second application shows that

$$S \leq 2 \min_{\mathbf{D}} \left(\mathbb{E} \max_j \sum_k \frac{(X_{jk} - a_{jk})^2}{d_k^2} \right)^{1/2}.$$

Bound the maximum with a sum:

$$S \leq 2 \min_{\mathbf{D}} \left(\sum_j \mathbb{E} \sum_k \frac{(X_{jk} - a_{jk})^2}{d_k^2} \right)^{1/2}.$$

The sum is controlled by a multiple of its largest term, so

$$S \leq 2\sqrt{m} \min_{\mathbf{D}} \left(\max_j \sum_k \frac{\text{Var}(X_{jk})}{d_k^2} \right)^{1/2},$$

where m is the number of rows of $\mathbf{A}$. □

6.1. Optimality. We now show that Theorem 4 gives an optimal bound, in the sense that each of its terms is necessary. In the following, we reserve the letter $\mathbf{D}$ for a positive diagonal matrix with $\text{trace}(\mathbf{D}^2) = 1$.

First, we establish the necessity of the Frobenius term by identifying a class of random matrices whose $\infty \rightarrow 2$ norms are larger than their weighted $2 \rightarrow \infty$ norms but comparable to their Frobenius norms. Let $\mathbf{Z}$ be a random $m \times \sqrt{m}$ matrix such that the entries in the first column of $\mathbf{Z}$ are equally likely to be positive or negative ones, and all other entries are zero. With this choice, $\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow 2} = \mathbb{E} \|\mathbf{Z}\|_F = \sqrt{m}$. Meanwhile, $\mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty} = \frac{1}{d_{11}}$, so $\min_{\mathbf{D}} \mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty} = 1$, which is much smaller than $\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow 2}$. Clearly, the Frobenius term is necessary.

Similarly, to establish the necessity of the weighted $2 \rightarrow \infty$ norm term, we consider a class of matrices whose $\infty \rightarrow 2$ norms are larger than their Frobenius norms but comparable to their weighted $2 \rightarrow \infty$ norms. Consider a $\sqrt{n} \times n$ matrix $\mathbf{Z}$ whose entries are all equally likely to be positive or negative ones. It is a simple task to confirm that $\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow 2} \geq n$ and $\mathbb{E} \|\mathbf{Z}\|_F = n^{3/4}$; it follows that the weighted $2 \rightarrow \infty$ norm term is necessary. In fact,

$$\min_{\mathbf{D}} \mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty} = \min_{\mathbf{D}} \mathbb{E} \max_{j=1, \dots, \sqrt{n}} \left(\sum_{k=1}^n \frac{Z_{jk}^2}{d_{kk}^2} \right)^{1/2} = \min_{\mathbf{D}} \left(\sum_{k=1}^n \frac{1}{d_{kk}^2} \right)^{1/2} = n,$$

so we see that $\mathbb{E} \|\mathbf{Z}\|_{\infty \rightarrow 2}$ and the weighted $2 \rightarrow \infty$ norm term are comparable.

6.2. Example application. From Theorem 4 we infer that a good scheme for sparsifying a matrix $\mathbf{A}$ while minimizing the expected relative $\infty \rightarrow 2$ norm error is one which drastically increases the sparsity of $\mathbf{X}$ while keeping the relative error

$$\frac{\mathbb{E} \|\mathbf{Z}\|_F + \min_{\mathbf{D}} \mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty}}{\|\mathbf{A}\|_{\infty \rightarrow 2}}$$

small, where $\mathbf{Z} = \mathbf{A} - \mathbf{X}$.

As before, consider the case where $\mathbf{A}$ is an $n \times n$ matrix all of whose entries are positive and in an interval bounded away from zero. Let γ be a desired bound on the expected relative $\infty \rightarrow 2$ norm error. We choose the randomization strategy $X_{jk} \sim \frac{a_{jk}}{p} \text{Bern}(p)$ and ask how much can we sparsify while respecting our bound on the relative error. That is, how small can p be? We appeal to Theorem 4. In this case,

$$\|\mathbf{A}\|_{\infty \rightarrow 2} = \left(\sum_j \sum_k a_{jk}^2 + 2 \sum_j \sum_{\ell < m} a_{j\ell} a_{jm} \right)^{\frac{1}{2}} = O\left((n^2 + n^2(n-1))^{\frac{1}{2}}\right).$$

By Jensen's inequality,

$$\mathbb{E} \|\mathbf{Z}\|_F \leq \mathbb{E} \|\mathbf{A}\|_F + \mathbb{E} \|\mathbf{X}\|_F \leq \left(1 + \frac{1}{\sqrt{p}}\right) \|\mathbf{A}\|_F = O\left(n \left(1 + \frac{1}{\sqrt{p}}\right)\right).$$

We bound the other term in the numerator, also using Jensen's inequality:

$$\min_{\mathbf{D}} \mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty} \leq \sqrt{n} \mathbb{E} \|\mathbf{Z}\|_{2 \rightarrow \infty} \leq \sqrt{n} \left(1 + \frac{1}{\sqrt{p}}\right) \|\mathbf{A}\|_{2 \rightarrow \infty} = O\left(n \left(1 + \frac{1}{\sqrt{p}}\right)\right)$$

to get

$$\frac{\mathbb{E} \|\mathbf{Z}\|_{\text{F}} + \min_{\mathbf{D}} \mathbb{E} \|\mathbf{Z}\mathbf{D}^{-1}\|_{2 \rightarrow \infty}}{\|\mathbf{A}\|_{\infty \rightarrow 2}} = O\left(\frac{1}{\sqrt{n}} + \frac{1}{\sqrt{pn}}\right) = O\left(\frac{1}{\sqrt{pn}}\right)$$

We conclude that, for this class of matrices and this family of sparsification schemes, we can reduce the number of expected nonzero terms to $O\left(\frac{n}{\gamma^2}\right)$ while maintaining an expected $\infty \rightarrow 2$ norm relative error of γ .

7. SPECTRAL ERROR BOUND

In this section we establish a bound on $\mathbb{E} \|\mathbf{A} - \mathbf{X}\|$ as an immediate consequence of Latała's result [Lat05]. We then derive a deviation inequality for the spectral approximation error using a log-Sobolev inequality from [BLM03], and use it to compare our results to those of Achlioptas and McSherry [AM07] and Arora, Hazan, and Kale [AHK06].

Theorem 5. *Suppose $\mathbf{A}$ is a fixed matrix, and let $\mathbf{X}$ be a random matrix with independent entries for which $\mathbb{E}\mathbf{X} = \mathbf{A}$. Then*

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\| \leq C \left[\max_j \left(\sum_k \text{Var}(X_{jk}) \right)^{1/2} + \max_k \left(\sum_j \text{Var}(X_{jk}) \right)^{1/2} + \left(\sum_{jk} \mathbb{E}(X_{jk} - a_{jk})^4 \right)^{1/4} \right]$$

where C is a universal constant.

In [Lat05], Latała considered the spectral norm of random matrices with independent, zero-mean entries, and he showed that, for any such matrix $\mathbf{Z}$,

$$\mathbb{E} \|\mathbf{Z}\| \leq C \left[\max_j \left(\sum_k \mathbb{E} Z_{jk}^2 \right)^{1/2} + \max_k \left(\sum_j \mathbb{E} Z_{jk}^2 \right)^{1/2} + \left(\sum_{jk} \mathbb{E} Z_{jk}^4 \right)^{1/4} \right],$$

where C is some universal constant. Unfortunately, no estimate for C is available. Theorem 5 follows from Latała's result, by taking $\mathbf{Z} = \mathbf{A} - \mathbf{X}$.

The bounded differences argument from Section 4 establishes the correct (subgaussian) tail behavior of $\mathbb{E} \|\mathbf{A} - \mathbf{X}\|$.

Theorem 6. *Fix the matrix $\mathbf{A}$, and let $\mathbf{X}$ be a random matrix with independent entries for which $\mathbb{E}\mathbf{X} = \mathbf{A}$. Assume $|X_{jk}| \leq D/2$ almost surely for all j, k . Then, for all $t > 0$,*

$$\mathbb{P}(\|\mathbf{A} - \mathbf{X}\| > \mathbb{E} \|\mathbf{A} - \mathbf{X}\| + t) \leq e^{-t^2/(4D^2)}.$$

Proof. The proof is exactly that of Theorem 2, except now $\mathbf{u}$ and $\mathbf{v}$ are both in the ℓ_2 unit sphere. $\square$

We find it convenient to measure deviations on the scale of the mean.

Corollary 4. *Under the conditions of Theorem 6, for all $\delta > 0$,*

$$\mathbb{P}(\|\mathbf{A} - \mathbf{X}\| > (1 + \delta)\mathbb{E} \|\mathbf{A} - \mathbf{X}\|) \leq e^{-\delta^2(\mathbb{E} \|\mathbf{A} - \mathbf{X}\|)^2/(4D^2)}.$$

7.1. Comparison with previous results. To demonstrate the applicability of our bound on the spectral norm error, we consider the sparsification and quantization schemes used by Achlioptas and McSherry [AM07], and the quantization scheme proposed by Arora, Hazan, and Kale [AHK06]. We show that our spectral norm error bound and the associated concentration result give results of the same order, with less effort. Throughout these comparisons, we take $\mathbf{A}$ to be a $m \times n$ matrix, with $m < n$, and we define $b = \max_{jk} |a_{jk}|$.

7.1.1. *A matrix quantization scheme.* First we consider the scheme proposed by Achlioptas and McSherry for quantization of the matrix entries:

$$X_{jk} = \begin{cases} b & \text{with probability } \frac{1}{2} + \frac{a_{jk}}{2b} \\ -b & \text{with probability } \frac{1}{2} - \frac{a_{jk}}{2b} \end{cases}.$$

With this choice $\text{Var}(X_{jk}) = b^2 - a_{jk}^2 \leq b^2$, and $\mathbb{E}(X_{jk} - a_{jk})^4 = b^2 - 3a^4 + 2a^2b^2 \leq 3b^4$, so the expected spectral error satisfies

$$\mathbb{E}\|\mathbf{A} - \mathbf{X}\| \leq C(\sqrt{nb} + \sqrt{mb} + b^4\sqrt{3mn}) \leq 4Cb\sqrt{n}.$$

Applying Corollary 4, we find that the error satisfies

$$\mathbb{P}(\|\mathbf{A} - \mathbf{X}\| > 4Cb\sqrt{n}(1 + \delta)) \leq e^{-\delta^2 C^2 n}.$$

In particular, with probability at least $1 - \exp(-C^2 n)$,

$$\|\mathbf{A} - \mathbf{X}\| \leq 8Cb\sqrt{n}.$$

Achlioptas and McSherry proved that for $n \geq n_0$, where n_0 is on the order of 10^9 , with probability at least $1 - \exp(-19(\log n)^4)$,

$$\|\mathbf{A} - \mathbf{X}\| < 4b\sqrt{n}.$$

Thus, Theorem 6 provides a bound of the same order in n which holds with higher probability and over a larger range of n .

7.1.2. *A nonuniform sparsification scheme.* Next we consider an analog to the nonuniform sparsification scheme proposed in the same paper. Fix a number p in the range $(0, 1)$ and sparsify entries with probabilities proportional to their magnitudes:

$$X_{jk} \sim \frac{a_{jk}}{p_{jk}} \text{Bern}(p_{jk}), \text{ where } p_{jk} = \max \left\{ p \left(\frac{a_{jk}}{b} \right)^2, \sqrt{p \left(\frac{a_{jk}}{b} \right)^2 \times (8 \log n)^4 / n} \right\}.$$

Achlioptas and McSherry determine that, with probability at least $1 - \exp(-19(\log n)^4)$,

$$\|\mathbf{A} - \mathbf{X}\| < 4b\sqrt{n/p}.$$

Further, the expected number of nonzero entries in $\mathbf{X}$ is less than

$$pmn \times \text{Avg}[(a_{jk}/b)^2] + m(8 \log n)^4. \quad (5)$$

Their choice of p_{jk} , in particular the insertion of the $(8 \log n)^4/n$ factor, is an artifact of their method of proof. Instead, we consider a scheme which compares the magnitudes of a_{jk} and b to determine p_{jk} . Introduce the quantity $R = \max_{a_{jk} \neq 0} b/|a_{jk}|$ to measure the spread of the entries in $\mathbf{A}$, and take

$$X_{jk} \sim \begin{cases} \frac{a_{jk}}{p_{jk}} \text{Bern}(p_{jk}), & \text{where } p_{jk} = \frac{pa_{jk}^2}{pa_{jk}^2 + b^2}, \quad a_{jk} \neq 0 \\ 0, & a_{jk} = 0. \end{cases}$$

With this scheme, $\text{Var}(X_{jk}) = 0$ when $a_{jk} = 0$, otherwise $\text{Var}(X_{jk}) = b^2/p$. Likewise, $\mathbb{E}(X_{jk} - a_{jk})^4 = 0$ if $a_{jk} = 0$, otherwise

$$\mathbb{E}(X_{jk} - a_{jk})^4 \leq \text{Var}(X_{jk}) \|X_{jk} - a_{jk}\|_\infty^2 = \frac{b^2}{p} \max \left\{ |a_{jk}|, |a_{jk}| \left(\frac{pa_{jk}^2 + b^2}{pa_{jk}^2} - 1 \right) \right\}^2 \leq \frac{b^4}{p^2} R^2,$$

so

$$\mathbb{E}\|\mathbf{A} - \mathbf{X}\| \leq C \left(b\sqrt{\frac{n}{p}} + b\sqrt{\frac{m}{p}} + b\sqrt{\frac{R}{p}} \sqrt[4]{mn} \right) \leq C(2 + \sqrt{R})b\sqrt{\frac{n}{p}}.$$

Applying Corollary 4, we find that with probability at least $1 - \exp(-C^2(2 + \sqrt{R})^2 pn/16)$,

$$\|\mathbf{A} - \mathbf{X}\| \leq 2C(2 + \sqrt{R})b\sqrt{\frac{n}{p}}.$$

Thus, Theorem 5 and Achlioptas and McSherry's scheme-specific analysis yield results of the same order in n and p . As before, we see that our bound holds with higher probability and over a larger range of n . Furthermore, since the expected number of nonzero entries in $\mathbf{X}$ satisfies

$$\sum_{jk} p_{jk} = \sum_{jk} \frac{pa_{jk}^2}{pa_{jk}^2 + b^2} \leq pnm \times \text{Avg} \left[\left(\frac{a_{jk}}{b} \right)^2 \right],$$

we have established a smaller limit on the expected number of nonzero entries.

7.1.3. A scheme which simultaneously sparsifies and quantizes. Finally, we use Theorem 6 to estimate the error in using the scheme from [AHK06] which simultaneously quantizes and sparsifies. Fix $\delta > 0$ and consider

$$X_{jk} = \begin{cases} \text{sgn}(a_{jk}) \frac{\delta}{\sqrt{n}} \text{Bern} \left(\frac{|a_{jk}| \sqrt{n}}{\delta} \right), & |a_{jk}| \leq \frac{\delta}{\sqrt{n}} \\ a_{jk}, & \text{otherwise.} \end{cases}$$

Then $\text{Var}(X_{jk}) = 0$ if $|a_{jk}| \geq \delta/\sqrt{n}$, otherwise

$$\text{Var}(X_{jk}) = |a_{jk}|^3 \frac{\sqrt{n}}{\delta} - 2a_{jk}^2 + |a_{jk}| \frac{\delta}{\sqrt{n}} \leq \frac{\delta^2}{n}.$$

Also the fourth moment term is zero for large enough a_{jk} , otherwise

$$\mathbb{E}(X_{jk} - a_{jk})^4 = |a_{jk}|^5 \frac{\sqrt{n}}{\delta} - 4a_{jk}^4 + 6|a_{jk}|^3 \frac{\delta}{\sqrt{n}} - 4a_{jk}^2 \frac{\delta^2}{n} + |a_{jk}| \left(\frac{\delta}{\sqrt{n}} \right)^3 \leq 8 \frac{\delta^4}{n^2}.$$

This gives the estimates

$$\mathbb{E} \|\mathbf{A} - \mathbf{X}\| \leq C \left(\sqrt{n} \frac{\delta}{\sqrt{n}} + \sqrt{m} \frac{\delta}{\sqrt{n}} + 2 \frac{\delta}{\sqrt{n}} \sqrt[4]{mn} \right) \leq 4C\delta$$

and

$$\mathbb{P}(\|\mathbf{A} - \mathbf{X}\| > 4C\delta(\gamma + 1)) \leq e^{-\gamma^2 C^2 n}.$$

Taking $\gamma = 1$, we see that with probability at least $1 - \exp(-C^2 n)$,

$$\|\mathbf{A} - \mathbf{X}\| \leq 8C\delta.$$

Let $S = \sum_{j,k} |A_{jk}|$, then appealing to Lemma 1 in [AHK06], we find that $\mathbf{X}$ has $O\left(\frac{\sqrt{n}S}{\gamma}\right)$ nonzero entries with probability at least $1 - \exp\left(-\Omega\left(\frac{\sqrt{n}S}{\gamma}\right)\right)$.

Arora, Hazan, and Kale establish that this scheme guarantees $\|\mathbf{A} - \mathbf{X}\| \leq O(\delta)$ with probability at least $1 - \exp(-\Omega(n))$, so we see that our general bound recovers a bound of the same order.

In conclusion, we see that the bound on expected spectral error in Theorem 6 in conjunction with the deviation result in Corollary 4 provide guarantees comparable to those derived with scheme-specific analyses. We anticipate that the flexibility demonstrated here will make these useful tools for analyzing and guiding the design of novel sparsification schemes.

REFERENCES

- [AHK06] Sanjeev Arora, Elad Hazan, and Satyen Kale, *A Fast Random Sampling Algorithm for Sparsifying Matrices*, Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques, Springer Berlin, 2006, pp. 272–279.
- [AM01] Dimitris Achlioptas and Frank McSherry, *Fast Computation of Low Rank Matrix Approximations*, Proceedings of the 33rd annual ACM symposium on Theory of Computing, 2001, pp. 611–618.
- [AM07] ———, *Fast Computation of Low Rank Matrix Approximations*, Journal of the ACM **54** (2007), no. 2.
- [AN04] Noga Alon and Assaf Naor, *Approximating the Cut-Norm via Grothendieck’s inequality*, Proceedings of the 36th Annual ACM symposium on Theory of Computing, 2004, pp. 72–80.
- [BLM03] Stéphane Boucheron, Gábor Lugosi, and Pascal Massart, *Concentration inequalities using the entropy method*, The Annals of Probability **31** (2003), 1583–1614.
- [BSS09] Joshua Batson, Daniel Spielman, and Nikhil Srivastava, *Twice-Ramanujan Sparsifiers*, Proceedings of the 41st Annual ACM Symposium on Theory of Computing, 2009, pp. 255–262.
- [DKM06a] Petros Drineas, Ravi Kannan, and Michael W. Mahoney, *Fast Monte Carlo Algorithms for Matrices I: Approximating Matrix Multiplication*, SIAM Journal of Computing **36** (2006), no. 1, 132–157.
- [DKM06b] ———, *Fast Monte Carlo Algorithms for Matrices II: Computing Low-Rank Approximations to a Matrix*, SIAM Journal of Computing **36** (2006), no. 1, 158–183.
- [DKM06c] ———, *Fast Monte Carlo Algorithms for Matrices III: Computing an Efficient Approximate Decomposition of a Matrix*, SIAM Journal of Computing **36** (2006), no. 1, 184–206.
- [FKV98] Alan Frieze, Ravi Kannan, and Santosh Vempala, *Fast Monte-Carlo Algorithms for finding low-rank approximations*, Proceedings of the 39th Annual Symposium on Foundations of Computer Science, 1998, pp. 378–390.
- [FKV04] ———, *Fast monte-carlo algorithms for finding low-rank approximations*, Journal of the ACM **51** (2004), 1025–1041.
- [Kar94a] David R. Karger, *Random sampling in cut, flow, and network design problems*, Proceedings of the 26th Annual ACM Symposium on Theory of Computing, May 1994, pp. 648–657.
- [Kar94b] ———, *Using randomized sparsification to approximate minimum cuts*, Proceedings of the 5th Annual ACM-SIAM Symposium on Discrete Algorithms, January 1994, pp. 424–432.
- [Kar95] ———, *Random Sampling in Graph Optimization Problems*, Ph.D. thesis, Stanford University, 1995.
- [Kar96] ———, *Approximating s - t minimum cuts in $\tilde{O}(n^2)$ time*, Proceedings of the 28th Annual ACM Symposium on Theory of Computing, May 1996, pp. 47–55.
- [Lat05] Rafal Łatała, *Some estimates of norms of random matrices*, Proceedings of the American Mathematical Society **133** (2005), 1273–1282.
- [LT91] Michel Ledoux and Michel Talagrand, *Probability in Banach Spaces*, Springer-Verlag, 1991.
- [Roh00] J. Rohn, *Computing the Norm $\|A\|_{\infty \rightarrow 1}$ is NP-Hard*, Linear and Multilinear Algebra **47** (2000), 195–204.
- [RV07] Mark Rudelson and Roman Vershynin, *Sampling from large matrices: An approach through geometric functional analysis*, Journal of the ACM **54** (2007), no. 4.
- [Seg00] Yoav Seginer, *The Expected Norm of Random Matrices*, Combinatorics, Probability and Computing **9** (2000), no. 2, 149–166.
- [SS08] Daniel A. Spielman and Nikhil Srivastava, *Graph sparsification by effective resistances*, Proceedings of the 40th annual ACM symposium on Theory of computing, 2008, pp. 563–568.
- [Sza76] S. J. Szarek, *On the best constants in the Khintchin-inequality*, Studia Math **58** (1976), 197–208.
- [Tro09] Joel Tropp, *Column subset selection, matrix factorization, and eigenvalue optimization*, Proceedings of the 19th Annual ACM-SIAM Symposium on Discrete Algorithms, 2009, pp. 978–986.
- [vW96] Aad W. van der Vaart and Jon A. Wellner, *Weak Convergence and Empirical Processes: With Applications to Statistics*, Springer, 1996.