

# MONTE CARLO METHODS FOR ESTIMATING THE DIAGONAL OF A REAL SYMMETRIC MATRIX\*

ERIC HALLMAN<sup>†</sup>, ILSE C. F. IPSEN<sup>†</sup>, AND ARVIND K. SAIBABA<sup>†</sup>

**Abstract.** For real symmetric matrices that are accessible only through matrix vector products, we present Monte Carlo estimators for computing the diagonal elements. Our probabilistic bounds for normwise absolute and relative errors apply to Monte Carlo estimators based on random Rademacher, sparse Rademacher, and normalized and unnormalized Gaussian vectors and to vectors with bounded fourth moments. The novel use of matrix concentration inequalities in our proofs represents a systematic model for future analyses. Our bounds mostly do not depend explicitly on the matrix dimension, target different error measures than existing work, and imply that the accuracy of the estimators increases with the diagonal dominance of the matrix. Applications to derivative-based global sensitivity metrics and node centrality measures in network science corroborate this, as do numerical experiments on synthetic test matrices. We recommend against the use in practice of sparse Rademacher vectors, which are the basis for many randomized sketching and sampling algorithms, because they tend to deliver barely a digit of accuracy even under large sampling amounts.

**Key words.** concentration inequalities, Monte Carlo methods, relative error, Rademacher random vectors, Gaussian random vectors

**MSC codes.** 15A15, 65C05, 65F50, 60G50, 68W20

**DOI.** 10.1137/22M1476277

**1. Introduction.** We compute the diagonal elements of symmetric matrices  $\mathbf{A} \in \mathbb{R}^{n \times n}$  with Monte Carlo estimators of the form

$$\hat{\mathbf{A}} = \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top,$$

where  $\mathbf{w}_k$  are independent random vectors. This approach is crucial when the elements of  $\mathbf{A}$  are available only implicitly via matrix vector products.

Estimating the diagonal elements of a matrix is important in many areas of science and engineering: In electronic structure calculations, one computes the diagonal elements of a projector onto the smallest eigenvectors of a Hamiltonian matrix [4]. In statistics, leverage scores for column subset selection can be computed from the diagonals of the projector onto the column space. In Bayesian inverse problems, the diagonal elements of the posterior covariance are computed with matrix-free estimators. Diagonal, or Jacobi, preconditioners can accelerate the convergence of iterative linear solvers [26]. More recently, diagonal estimators have been used to accelerate second-order optimization techniques for machine learning [27]. In network science, subgraph centrality measures rank the importance of the network nodes based on the diagonal of a scaled exponential of the adjacency matrix. In sensitivity analy-

\* Received by the editors February 7, 2022; accepted for publication (in revised form) November 2, 2022; published electronically March 10, 2023.

<https://doi.org/10.1137/22M1476277>

**Funding:** The work of the authors was partially supported by National Science Foundation (NSF) grant DMS-1745654. The work of the second author was also partially supported by NSF grant DMS-1760374 and Department of Energy grant DE-SC0022085. The work of the third author was also partially supported by NSF grant DMS-1845406.

<sup>†</sup>Department of Mathematics, North Carolina State University, Raleigh, NC 27695 USA (eric.r.hallman@gmail.com, ipsen@ncsu.edu, asaibab@ncsu.edu).

sis, Monte Carlo diagonal estimators efficiently compute the derivative-based global sensitivity metrics [7, 13].

Diagonal estimation is related to trace estimation. Once the diagonal elements are known, the trace can be computed from their sum. Therefore, estimators for the diagonal of a matrix can be easily adapted to trace estimators. Monte Carlo methods were first proposed by Hutchinson [11] and subsequently improved and expanded to different distributions [2, 8, 19]. Applications of trace estimators, reviewed in [24], include estimating density of states, log determinants, and Schatten  $p$ -norms.

Diagonal estimation has also been applied to the inverse of a matrix, with methods based on domain decomposition [21] and deterministic probing vectors derived from graph coloring [22].

*Literature review.* To our knowledge, Monte Carlo diagonal estimators were first proposed by Bekas, Kokiopoulou, and Saad [4], who gave a sufficient condition for a Monte Carlo estimator to be unbiased. They also pointed out that large off-diagonal elements can result in large relative errors and developed probing methods to mitigate the effects of the off-diagonal elements. This idea is further explored in [12, 14, 22].

We are aware of a recent paper [3] as the only other work to analyze the sampling amount required to achieve a user-specified relative error with high probability. In contrast to [3], our proofs are the first to exploit matrix concentration inequalities to impose a systematic structure that can serve as a model for future analyses and allow us to analyze the normwise errors in different norms. We analyze more general distributions such as random vectors with bounded fourth moments and sparse Rademacher vectors with a user-specified sparsity parameter, and—in contrast to [3]—focus on unnormalized estimators. Most of our bounds do not show an explicit dependence on the matrix dimension, which is desirable for large-scale problems.

**1.1. Contributions and overview.** After introducing notation, relevant concentration inequalities, and the setup for our analysis (section 2.1), we derive normwise error bounds for Monte Carlo estimators based on Rademacher vectors (section 3), random vectors with bounded fourth moments and Gaussian vectors (section 4), and componentwise bounds for Rademacher and Gaussian vectors (section 5) and apply Monte Carlo estimators to derivative-based global sensitivity metrics (section 6). Numerical experiments (section 7) and implementations (Algorithms 7.1 and 7.2) illustrate the accuracy and cost effectiveness of the Monte Carlo estimators. The novel features of our contributions are the following:

1. Most of our bounds exhibit no explicit dependence on the matrix dimension  $n$ , including the normwise (section 3) and elementwise (section 5) bounds for Rademacher-based estimators. Moreover, our bounds hold for all symmetric matrices, whether positive definite or not.
2. Our normwise bounds suggest that Rademacher-based Monte Carlo estimators are more accurate for matrices that are more strongly diagonally dominant (Theorem 3.2). In particular, the least sampling amount required for the Monte Carlo estimators to achieve a user-specified relative error decreases with increasing diagonal dominance of  $\mathbf{A}$  in the relative sense (Corollaries 3.3 and 4.2).
3. We introduce Monte Carlo estimators based on Rademacher vectors that are parameterized in terms of sparsity levels (Definition 3.4) and show that they lose accuracy with increasing sparsity (Theorem 3.5, Corollary 3.6). Numerical experiments (section 7) confirm that, even for large sampling amounts, the

estimators barely achieve a single digit of accuracy. Therefore, we recommend against their use in practice.

4. Our componentwise bounds suggest that the accuracy for computing a diagonal element  $a_{ii}$  depends only on the diagonal dominance of column/row  $i$  of  $\mathbf{A}$  (Corollaries 5.4 and 5.6).
5. In the context of derivative-based global sensitivity metrics, we design and analyze Monte Carlo estimators based on random vectors from a problem-specific probability distribution (Theorem 6.1, Corollary 6.2). Furthermore, we illustrate the accuracy of Rademacher-based Monte Carlo estimators for node centrality measures in network science (section 7.8).

**2. Background.** After reviewing notation (section 2.1) and relevant concentration inequalities (section 2.2), we present the setup for our analysis (section 2.3).

**2.1. Notation.** The *Schur product* (or *Hadamard*, or *elementwise*, *product*) of  $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{m \times n}$  is denoted by  $\mathbf{C} = \mathbf{A} \circ \mathbf{B} \in \mathbb{R}^{m \times n}$  and has elements

$$c_{ij} = a_{ij}b_{ij} \quad 1 \leq i \leq m, \quad 1 \leq j \leq n.$$

For  $\mathbf{A}, \mathbf{B}, \mathbf{C} \in \mathbb{R}^{m \times n}$ , the Schur product is commutative and distributive,

$$\mathbf{A} \circ \mathbf{B} = \mathbf{B} \circ \mathbf{A}, \quad \mathbf{A} \circ (\mathbf{B} + \mathbf{C}) = \mathbf{A} \circ \mathbf{B} + \mathbf{A} \circ \mathbf{C}.$$

Following MATLAB convention, we define  $\text{diag}(\mathbf{A}) = [a_{11} \ \cdots \ a_{nn}]^\top \in \mathbb{R}^n$  as the column vector of diagonal elements of  $\mathbf{A} \in \mathbb{R}^{n \times n}$ . The operator  $\text{diag}$  is overloaded, and  $\text{diag}(\mathbf{x}) \in \mathbb{R}^{n \times n}$  represents a diagonal matrix whose diagonal elements are the elements of the vector  $\mathbf{x} \in \mathbb{R}^n$ . In particular,

$$(2.1) \quad \mathcal{D}(\mathbf{A}) \equiv \text{diag}(\text{diag}(\mathbf{A})) = \mathbf{I} \circ \mathbf{A} = \begin{bmatrix} a_{11} & & \\ & \ddots & \\ & & a_{nn} \end{bmatrix} \in \mathbb{R}^{n \times n}.$$

In other words,  $\mathbf{I} \circ \mathbf{A}$  zeros out the off-diagonal elements of  $\mathbf{A}$ .

If the first factor in a Schur product is a square matrix  $\mathbf{M} \in \mathbb{R}^{n \times n}$  and the second factor an outer product involving  $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ , then

$$(2.2) \quad \mathbf{M} \circ (\mathbf{x}\mathbf{y}^\top) = \text{diag}(\mathbf{x}) \mathbf{M} \text{diag}(\mathbf{y}).$$

For symmetric matrices  $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$ , the partial order  $\mathbf{A} \preceq \mathbf{B}$  or, equivalently,  $\mathbf{B} \succeq \mathbf{A}$  says that  $\mathbf{B} - \mathbf{A}$  is positive semidefinite. If  $\mathbf{A}$  and  $\mathbf{B}$  are positive semidefinite, then  $\mathbf{A} \preceq \mathbf{B}$  implies that  $\mathbf{A}^{1/2} \preceq \mathbf{B}^{1/2}$ .

The *intrinsic dimension* of a nonzero symmetric positive semidefinite matrix  $\mathbf{A} \in \mathbb{R}^{n \times n}$  is

$$\text{intdim}(\mathbf{A}) \equiv \frac{\text{trace}(\mathbf{A})}{\|\mathbf{A}\|_2} \quad \text{with} \quad 1 \leq \text{intdim}(\mathbf{A}) \leq \text{rank}(\mathbf{A}) \leq n.$$

If, additionally,  $\mathbf{A}$  is a diagonal matrix, then

$$\text{intdim}(\mathbf{A}) = \frac{\sum_{i=1}^n a_{ii}}{\max_{1 \leq i \leq n} a_{ii}}.$$

The columns of  $\mathbf{A} = [\mathbf{a}_1 \ \cdots \ \mathbf{a}_n] \in \mathbb{R}^{m \times n}$  are  $\mathbf{a}_j \in \mathbb{R}^m$ ,  $1 \leq j \leq n$ , and the columns of the identity  $\mathbf{I} = [\mathbf{e}_1 \ \cdots \ \mathbf{e}_n] \in \mathbb{R}^{n \times n}$  are  $\mathbf{e}_j \in \mathbb{R}^n$ . The transpose of  $\mathbf{A}$  is  $\mathbf{A}^\top$ .

We denote by  $\mathbb{P}[\mathcal{E}]$  the probability of an event  $\mathcal{E}$  and by  $\mathbb{E}[Z]$  the expectation of a random variable  $Z$ , which can be a scalar-, vector-, or matrix-valued.

**2.2. Concentration inequalities.** We rely on two scalar and two matrix concentration inequalities.

Markov's inequality [18, section 3.1] bounds the probability that a random variable exceeds a constant.

**THEOREM 2.1** (Markov's inequality). *If  $Z$  is a nonnegative random variable, then for  $t > 0$ ,*

$$\mathbb{P}[Z \geq t] \leq \frac{\mathbb{E}[Z^2]}{t^2}.$$

Hoeffding's inequality for general bounded random variables [25, Theorem 2.2.6] bounds the probability that a sum of scalar random variables exceeds its mean.

**THEOREM 2.2** (Scalar Hoeffding inequality). *Let  $Z_1, \dots, Z_N$  be independent random variables, bounded by  $m_k \leq Z_k \leq M_k$ ,  $1 \leq k \leq N$ , with sum  $Z \equiv \sum_{k=1}^N Z_k$ . Then for  $t > 0$ ,*

$$\mathbb{P}[|Z - \mathbb{E}[Z]| \geq t] \leq 2 \exp \left( \frac{-2t^2}{\sum_{k=1}^N (M_k - m_k)^2} \right).$$

Next are two bounds for sums of independent symmetric matrix-valued random variables. The first is a matrix Bernstein concentration inequality [23, Theorems 7.3.1 and 7.7.1] for sums of independent, symmetric, bounded, zero-mean random matrices.

**THEOREM 2.3** (Matrix Bernstein inequality). *Let  $\mathbf{S}_1, \dots, \mathbf{S}_N \in \mathbb{R}^{n \times n}$  be independent symmetric random matrices with*

$$\mathbb{E}[\mathbf{S}_k] = \mathbf{0}, \quad \|\mathbf{S}_k\|_2 \leq L, \quad 1 \leq k \leq N.$$

*Let  $\mathbf{S} \equiv \sum_{k=1}^N \mathbf{S}_k$  have a matrix-valued variance that is majorized by  $\mathbf{V} \in \mathbb{R}^{n \times n}$ :*

$$\mathbf{V} \succeq \text{Var}(\mathbf{S}) = \mathbb{E}[\mathbf{S}^2] = \sum_{k=1}^N \mathbb{E}[\mathbf{S}_k^2].$$

*Abbreviate  $\nu \equiv \|\mathbf{V}\|_2$  and  $d \equiv \text{intdim}(\mathbf{V})$ . Then for  $t > 0$ ,*

$$(2.3) \quad \mathbb{P}[\|\mathbf{S}\|_2 \geq t] \leq 8d \exp \left( \frac{-t^2}{2(\nu + Lt/3)} \right).$$

*Proof.* In [23, Theorems 7.3.1 and 7.7.1], it is shown that (2.3) holds, provided  $t \geq \sqrt{\nu} + \frac{L}{3}$ . If  $t < \sqrt{\nu} + L/3$ . Then the bound holds vacuously because the right-hand side of (2.3) is less than 1 if and only if

$$t > \frac{L\ell}{3} + \sqrt{(L\ell/3)^2 + 2\ell\nu}, \quad \text{where } \ell \equiv \ln(8d).$$

Since  $\ell > 1$ , this lower bound on  $t$  is strictly greater than  $\sqrt{\nu} + L/3$ .  $\square$

The second matrix concentration inequality [6, Theorem 3.2] bounds the mean of the squared norm of the sum of symmetric random matrices.

**THEOREM 2.4.** *Let  $\mathbf{S}_1, \dots, \mathbf{S}_N \in \mathbb{R}^{n \times n}$  with  $n \geq 3$  be independent centered symmetric random matrices with zero mean. Then*

$$\mathbb{E} \left[ \left\| \sum_{k=1}^N \mathbf{S}_k \right\|_2^2 \right]^{1/2} \leq \sqrt{2e \ln n} \left\| \left( \sum_{k=1}^N \mathbb{E}[\mathbf{S}_k^2] \right)^{1/2} \right\|_2 + 4e \ln n \left( \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{S}_k\|_2^2 \right] \right)^{1/2}.$$

**2.3. Setup for the analysis.** Our Monte Carlo estimators compute the diagonal elements of a symmetric matrix  $\mathbf{A} \in \mathbb{R}^{n \times n}$  by means of matrix vector products with  $\mathbf{A}$ . They sample  $N$  independent random vectors  $\mathbf{w}_k \in \mathbb{R}^n$  and approximate the vector of diagonal elements  $\text{diag}(\mathbf{A}) \in \mathbb{R}^n$  by the mean

$$\text{diag}(\hat{\mathbf{A}}) = \frac{1}{N} \sum_{k=1}^N ((\mathbf{A}\mathbf{w}_k) \circ \mathbf{w}_k) \in \mathbb{R}^n, \quad \text{where} \quad \hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{A}\mathbf{w}_k\mathbf{w}_k^\top \in \mathbb{R}^{n \times n}.$$

To see this, write the  $i$ th diagonal element of the estimator as

$$\begin{aligned} \hat{A}_{ii} &= \frac{1}{N} \sum_{k=1}^N (\mathbf{A}\mathbf{w}_k\mathbf{w}_k^\top)_{ii} = \frac{1}{N} \sum_{k=1}^N (\mathbf{A}\mathbf{w}_k)_i (\mathbf{w}_k^\top)_i \\ &= \frac{1}{N} \sum_{k=1}^N ((\mathbf{A}\mathbf{w}_k) \circ \mathbf{w}_k)_i, \quad 1 \leq i \leq n. \end{aligned}$$

We measure the cost of a diagonal estimator by the number  $N$  of samples and assess the accuracy with normwise and componentwise relative errors. Section 7.1 presents pseudocodes for the estimators and a discussion of the computational cost.

**3. Normwise bounds for Rademacher random vectors.** We present normwise bounds for Monte Carlo estimators based on standard (section 3.1) and on sparse Rademacher vectors (section 3.2).

**3.1. Standard Rademacher vectors.** After defining Rademacher vectors (Definition 3.1) and discussing their properties (Remarks 3.1 and 3.2), we present a normwise absolute error bound (Theorem 3.2) and a bound on the minimal sampling amount to achieve a user-specified relative error (Corollary 3.3).

**DEFINITION 3.1.** A Rademacher random variable takes on the values  $\pm 1$  with equal probability  $1/2$ . A Rademacher vector is a random vector whose elements are independent Rademacher random variables.

Standard Rademacher vectors have the advantage of cheap matrix vector products and the ability to immediately recover diagonal matrices.

*Remark 3.1.* The elements  $w_j$  of a Rademacher vector  $\mathbf{w} \in \mathbb{R}^n$  have the following properties:

1. Zero mean:  $\mathbb{E}[w_j] = 0$ ,  $1 \leq j \leq n$ .
2. Constant square:  $w_j^2 = 1$ ,  $1 \leq j \leq n$ .
3. Independence:  $\mathbb{E}[w_j w_i] = 0$  for  $i \neq j$ .

*Remark 3.2.* Standard Rademacher vectors recover a diagonal matrix with a single sample,  $N = 1$ .

To see this, let  $\mathbf{A} = \mathcal{D}(\mathbf{A}) \in \mathbb{R}^{n \times n}$  be diagonal, and  $\mathbf{w} \in \mathbb{R}^n$  be a Rademacher vector. Remark 3.1 implies that  $\mathbf{A}\mathbf{w}\mathbf{w}^\top \in \mathbb{R}^{n \times n}$  has diagonal elements  $a_{ii}w_i^2 = a_{ii}$ ,  $1 \leq i \leq n$ .

As a consequence, we can focus the analysis of standard Rademacher-based estimators on nondiagonal matrices. The  $p$ -norm bound below is a special case of the one for sparse Rademacher vectors in section 3.2.

**THEOREM 3.2.** Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric and

$$\begin{aligned} K_1 &\equiv \|\mathcal{D}(\mathbf{A}^2) - \mathcal{D}(\mathbf{A})^2\|_p, & K_2 &\equiv \|\mathbf{A} - \mathcal{D}(\mathbf{A})\|_\infty, \\ d &\equiv \text{trace}(\mathcal{D}(\mathbf{A}^2) - \mathcal{D}(\mathbf{A})^2) / K_1. \end{aligned}$$

If  $\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top$  is a Monte Carlo estimator with independent Rademacher vectors  $\mathbf{w}_k \in \mathbb{R}^n$ ,  $1 \leq k \leq N$ , then the probability that the absolute error exceeds  $t > 0$  is at most

$$\mathbb{P} \left[ \|\mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}})\|_p \geq t \right] \leq 8d \exp \left( \frac{-Nt^2}{2(K_1 + tK_2/3)} \right).$$

*Proof.* This is the special case  $s = 1$  of Theorem 3.5.  $\square$

*Interpretation.* The constants  $K_1$ ,  $K_2$ , and  $d$  represent the deviation of  $\mathbf{A}$  from diagonality. More specifically,  $K_1$  and  $K_2$  represent the degree of diagonal dominance of  $\mathbf{A}$  in the absolute sense and  $d$  in the relative sense. Theorem 3.2 implies that the Rademacher estimator has a small absolute error when applied to strongly diagonally dominant matrices. In other words, the normwise absolute error in the Rademacher estimator decreases with increasing diagonal dominance of  $\mathbf{A}$ .

Remark 3.4 shows that  $K_1$  and  $d$  are equal to quantities in the probability (2.3), while only  $K_2$  represents a bound.

We determine the least sampling amount required for a Rademacher-based Monte Carlo estimator to achieve a user-specified normwise relative error  $\epsilon$  with a user-specified success probability  $1 - \delta$ . For convenience, Corollary 3.3 is expressed in terms of a single norm.

**COROLLARY 3.3.** Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric. Let

$$K_1 \equiv \|\mathcal{D}(\mathbf{A}^2) - \mathcal{D}(\mathbf{A})^2\|_\infty,$$

$$\Delta_1 \equiv \frac{K_1}{\|\mathcal{D}(\mathbf{A})\|_\infty^2}, \quad \Delta_2 \equiv \frac{\|\mathbf{A} - \mathcal{D}(\mathbf{A})\|_\infty}{\|\mathcal{D}(\mathbf{A})\|_\infty}, \quad d \equiv \frac{\text{trace}(\mathcal{D}(\mathbf{A}^2) - \mathcal{D}(\mathbf{A})^2)}{K_1},$$

and let  $\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top$  be a Monte Carlo estimator with independent Rademacher vectors  $\mathbf{w}_k \in \mathbb{R}^n$ ,  $1 \leq k \leq N$ . Pick  $\epsilon > 0$  and  $0 < \delta < 1$ . If the sampling amount is at least

$$(3.1) \quad N \geq \frac{2}{3\epsilon^2} (3\Delta_1 + \epsilon\Delta_2) \ln(8d/\delta),$$

then  $\|\mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}})\|_\infty \leq \epsilon \|\mathcal{D}(\mathbf{A})\|_\infty$  holds with probability at least  $1 - \delta$ .

*Proof.* This is the special case  $s = 1$  of Corollary 3.6.  $\square$

*Interpretation.* The constants  $\Delta_1$  and  $\Delta_2$  in Corollary 3.3 represent the respective relative counterparts of  $K_1$  and  $K_2$  in Theorem 3.2: They represent the relative deviation of  $\mathbf{A}$  from diagonality and more specifically the degree of diagonal dominance of  $\mathbf{A}$  in the relative sense. Corollary 3.3 implies that if  $\mathbf{A}$  is strongly diagonally dominant in the relative sense, then a small sampling amount suffices to achieve a normwise relative error with user-specified probability. As with many randomized sampling algorithms, the lower bound for  $N$  is proportional to  $1/\epsilon^2$ .

**3.2. Sparse Rademacher vectors.** For Rademacher vectors that are parameterized in terms of sparsity (Definition 3.4, Remark 3.3), we derive a normwise absolute error bound (Theorem 3.5) and comment on the constants and the sparsity (Remarks 3.4 and 3.5), followed by the minimal sampling amount to achieve a user-specified error relative error (Corollary 3.6).

The random vectors in [1] have elements that assume values from the discrete distribution  $\{-\sqrt{3}, 0, \sqrt{3}\}$  with respective probability  $\{\frac{1}{6}, \frac{2}{3}, \frac{1}{6}\}$ . This concept was extended in [17, equation (2)] to Rademacher vectors that are parameterized in terms of a sparsity parameter  $s$ .

DEFINITION 3.4. A sparse Rademacher random variable with parameter  $s \geq 1$  takes the values  $\{-\sqrt{s}, 0, \sqrt{s}\}$  with probability  $\{\frac{1}{2s}, 1 - \frac{1}{s}, \frac{1}{2s}\}$ , respectively.

A sparse Rademacher vector is a random vector whose elements are independent sparse Rademacher random variables.

The properties of sparse Rademacher vectors are almost the same as those of the standard Rademacher vectors in Remark 3.1.

Remark 3.3. The elements of a sparse Rademacher vector  $\mathbf{w} \in \mathbb{R}^n$  with parameter  $s \geq 1$  have the following properties:

1. Zero mean:  $\mathbb{E}[w_j] = 0$ ,  $1 \leq j \leq n$ .
2. Unit variance:  $\mathbb{E}[w_j^2] = 1$ ,  $1 \leq j \leq n$ .
3. Fourth moment:  $\mathbb{E}[w_j^4] = s$ ,  $1 \leq j \leq n$ .
4. Independence: For  $i \neq j$  and integer  $\ell \geq 1$ .

$$\mathbb{E}[w_i^2 w_j^2] = 1, \quad \mathbb{E}[w_i^\ell w_j] = \mathbb{E}[w_i w_j^\ell] = 0.$$

The case  $s = 1$  corresponds to the standard Rademacher vectors (Definition 3.1), while  $s = 3$  corresponds to the choice in [1].

Below is the extension of the  $p$ -norm bound in Theorem 3.2 to sparse Rademacher vectors with integer parameters  $s$ .

THEOREM 3.5. Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric, and let

$$K_1(s) \equiv \|\mathcal{D}(\mathbf{A}^2) + (s-2)\mathcal{D}(\mathbf{A})^2\|_p, \quad K_2(s) \equiv \|s\mathbf{A} - \mathcal{D}(\mathbf{A})\|_\infty, \\ d(s) \equiv \text{trace}(\mathcal{D}(\mathbf{A}^2) + (s-2)\mathcal{D}(\mathbf{A})^2) / K_1(s).$$

Let  $\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{j=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top$  be a Monte Carlo estimator with independent sparse Rademacher random vectors  $\mathbf{w}_k \in \mathbb{R}^n$  with integer parameter  $s \geq 1$ . Then the probability that the absolute error exceeds  $t > 0$  is at most

$$\mathbb{P} \left[ \|\mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}})\|_p \geq t \right] \leq 8d(s) \exp \left( \frac{-Nt^2}{2(K_1(s) + tK_2(s)/3)} \right).$$

Proof. Define the random diagonal matrices

$$(3.2) \quad \mathbf{S}_k \equiv \frac{1}{N} (\mathbf{I} \circ (\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top) - \mathbf{I} \circ \mathbf{A}), \quad 1 \leq k \leq N,$$

and their sum

$$(3.3) \quad \mathbf{Z} \equiv \sum_{k=1}^N \mathbf{S}_k = \mathbf{I} \circ \hat{\mathbf{A}} - \mathbf{I} \circ \mathbf{A} = \mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}}).$$

Before applying Theorem 2.3, we need to verify the assumptions for the Bernstein inequality.

1. *Expectation.* Remark 3.3 implies that  $\mathbb{E}[\mathbf{w}_k \mathbf{w}_k^\top] = \mathbf{I}$ , and from the linearity of the expectation follows

$$\mathbb{E}[\mathbf{S}_k] = \mathbb{E} \left[ \frac{1}{N} (\mathcal{D}(\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top) - \mathcal{D}(\mathbf{A})) \right] \\ = \frac{1}{N} (\mathcal{D}(\mathbf{A} \mathbb{E}[\mathbf{w}_k \mathbf{w}_k^\top]) - \mathcal{D}(\mathbf{A})) = \mathbf{0}, \quad 1 \leq k \leq N.$$

2. *Boundedness.* From (3.2) follows that the diagonal matrices  $\mathbf{S}_k$  have diagonal elements

$$(3.4) \quad (\mathbf{S}_k)_{ii} = \frac{1}{N} ((\mathbf{A}\mathbf{w}_k)_i (\mathbf{w}_k)_i - a_{ii}) = \frac{1}{N} \left( (\mathbf{w}_k)_i \sum_{j=1}^n a_{ij} (\mathbf{w}_k)_j - a_{ii} \right) \\ = \frac{1}{N} \left( a_{ii} ((\mathbf{w}_k)_i^2 - 1) + (\mathbf{w}_k)_i \sum_{j \neq i} a_{ij} (\mathbf{w}_k)_j \right), \quad 1 \leq i \leq n, \quad 1 \leq k \leq N.$$

Because  $s \geq 1$  is an integer, we have  $|((\mathbf{w}_k)_i^2 - 1)a_{ii}| \leq (s-1)|a_{ii}|$  and

$$|(\mathbf{S}_k)_{ii}| \leq \frac{1}{N} \left( (s-1)|a_{ii}| + s \sum_{j \neq i} |a_{ij}| \right) \\ = \frac{1}{N} (s \|\mathbf{a}_i\|_1 - |a_{ii}|), \quad 1 \leq i \leq n, \quad 1 \leq k \leq N.$$

Since  $\mathbf{S}_k$  is diagonal, its 2-norm is bounded by

$$\|\mathbf{S}_k\|_2 = \max_{1 \leq i \leq n} |(\mathbf{S}_k)_{ii}| \leq \frac{1}{N} \max_{1 \leq i \leq n} (s \|\mathbf{a}_i\|_1 - |a_{ii}|) \\ = \frac{1}{N} \|s\mathbf{A} - \mathcal{D}(\mathbf{A})\|_\infty = \frac{K_2(s)}{N}, \quad 1 \leq k \leq N.$$

Set  $L(s) \equiv K_2(s)/N$ , where  $K_2(s) > 0$  since  $\mathbf{A}$  is not diagonal.

3. *Variance.* Since  $\mathbf{S}_k$  is diagonal, so is  $\mathbf{S}_k^2$ . Unraveling the individual diagonal elements in (3.4) and abbreviating  $\mathbf{u} \equiv \mathbf{w}_k$  gives

$$N^2 (\mathbf{S}_k^2)_{ii} = \left( u_i \sum_{j=1}^n a_{ij} u_j - a_{ii} \right)^2 = u_i^2 \sum_{j=1}^n \sum_{j'=1}^n a_{ij} a_{ij'} u_j u_{j'} - 2a_{ii} u_i \sum_{j=1}^n a_{ij} u_j + a_{ii}^2.$$

Repeated application of the properties in Remark 3.3 leads to

$$\mathbb{E}[N^2 (\mathbf{S}_k^2)_{ii}] = \mathbb{E} \left[ a_{ii}^2 u_i^4 + \sum_{j \neq i} a_{ij}^2 u_i^2 u_j^2 - 2a_{ii}^2 u_i^2 + a_{ii}^2 \right] \\ = s a_{ii}^2 + \sum_{j \neq i} a_{ij}^2 - 2a_{ii}^2 + a_{ii}^2 \\ = \|\mathbf{a}_i\|_2^2 + (s-2)a_{ii}^2.$$

Sum up the individual variances:

$$\mathbf{V}(s) \equiv \text{Var}[\mathbf{Z}] = \sum_{k=1}^N \mathbb{E}[\mathbf{S}_k^2] = \frac{1}{N^2} \sum_{k=1}^N (\mathcal{D}(\mathbf{A}^2) + (s-2)\mathbf{D}^2) \\ = \frac{1}{N} (\mathcal{D}(\mathbf{A}^2) + (s-2)\mathbf{D}^2).$$

Since  $\mathbf{V}(s)$  is diagonal, all  $p$ -norms are the same,

$$\nu(s) \equiv \|\mathbf{V}(s)\|_2 = \frac{1}{N} \|\mathcal{D}(\mathbf{A}^2) + (s-2)\mathbf{D}^2\|_p = \frac{K_1(s)}{N},$$



where  $K_1(s) > 0$  since  $\mathbf{A}$  is not diagonal. The intrinsic dimension of  $\mathbf{V}(s)$  is

$$d(s) \equiv \text{intdim}(\mathbf{V}) = \frac{\text{trace}(\mathcal{D}(\mathbf{A}^2) + (s-2)\mathcal{D}^2)}{N\nu(s)}.$$

4. *Apply Theorem 2.3.* Substituting  $L(s) = K_2(s)/N$  and  $\nu(s) = K_1(s)/N$  into Theorem 2.3, exploiting the fact that the  $p$ -norms of a diagonal matrix are all the same, and remembering that the sum  $\mathbf{Z}$  in (3.3) has zero mean gives

$$\mathbb{P}[\|\mathbf{Z}\|_p \geq t] = \mathbb{P}\left[\|\mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}})\|_p \geq t\right] \leq 8d(s) \exp\left(\frac{-Nt^2}{2(K_1(s) + tK_2(s)/3)}\right). \quad \square$$

*Interpretation.* In the special case  $s = 1$  of standard Rademacher vectors, Theorem 3.5 reduces to Theorem 3.2. However, sparse Rademacher Monte Carlo estimators with  $s > 1$  do, in general, not recover a diagonal matrix with a single sample,  $N = 1$ .

As  $s$  increases, so do the constants  $K_1(s)$  and  $K_2(s)$  and the upper bound as a whole. In other words, the sparser the vectors  $\mathbf{w}_k$ , the less accurate the Monte Carlo estimate  $\mathcal{D}(\hat{\mathbf{A}})$ .

*Remark 3.4* (tightness of the bounds). The constants in Theorem 3.5 arise naturally in the concentration inequality in Theorem 3.3. Specifically,  $K_1(s)$  and  $d(s)$  are equal to quantities in the probability (2.3), and only  $K_2(s)$  represents a bound.

To see this, note that  $K_1(s)$  and  $d(s)$  are part of expressions—not bounds—for the variance norm and intrinsic dimension:

$$\begin{aligned} \mathbf{V} &= \text{Var}(\mathbf{S}) = \frac{1}{N} (\mathcal{D}(\mathbf{A}^2) + (s-2)\mathcal{D}(\mathbf{A})^2), \\ \nu &= \|\mathbf{V}\|_p = \frac{1}{N} K_1(s), \\ d(s) &= \text{intdim}(\mathbf{V}) = \frac{1}{K_1(s)} \text{trace}(\mathcal{D}(\mathbf{A}^2) + (s-2)\mathcal{D}(\mathbf{A})^2). \end{aligned}$$

*Remark 3.5* (noninteger sparsity levels). The restriction to integers  $s$  in Theorem 3.5 is relevant only for  $1 < s < 2$ . More generally, Theorem 3.5 holds for  $s = 1$  and any real number  $s \geq 2$ .

The extension below of Corollary 3.3 presents the minimal sampling amount so that the sparse Rademacher Monte Carlo estimator achieves a user-specified normwise relative error  $\epsilon$  at a user-specified success probability  $1 - \delta$ . For ease of understanding, Corollary 3.3 is expressed in terms of a single norm.

**COROLLARY 3.6.** *Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric, and let*

$$\begin{aligned} K_1(s) &\equiv \|\mathcal{D}(\mathbf{A}^2) + (s-2)\mathcal{D}(\mathbf{A})^2\|_\infty, & \Delta_1(s) &\equiv \frac{K_1(s)}{\|\mathcal{D}(\mathbf{A})\|_\infty^2} \\ \Delta_2(s) &\equiv \frac{\|s\mathbf{A} - \mathcal{D}(\mathbf{A})\|_\infty}{\|\mathcal{D}(\mathbf{A})\|_\infty}, & d(s) &\equiv \frac{\text{trace}(\mathcal{D}(\mathbf{A}^2) + (s-2)\mathcal{D}(\mathbf{A})^2)}{K_1(s)}. \end{aligned}$$

Let  $\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{j=1}^N \mathbf{A} \mathbf{w}_j \mathbf{w}_j^\top$  be a Monte Carlo estimator with independent sparse Rademacher random vectors  $\mathbf{w}_k \in \mathbb{R}^n$  with integer parameter  $s \geq 1$ . Pick  $\epsilon > 0$  and  $0 < \delta < 1$ . If the sampling amount is at least

$$(3.5) \quad N \geq \frac{2}{3\epsilon^2} (3\Delta_1(s) + \epsilon\Delta_2(s)) \ln(8d(s)/\delta),$$

then  $\|\mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}})\|_\infty \leq \epsilon \|\mathcal{D}(\mathbf{A})\|_\infty$  holds with probability at least  $1 - \delta$ .

*Proof.* Set  $p = \infty$  for the norms, denote the bound for the failure probability in Theorem 3.5 by

$$\delta \equiv 8d(s) \exp\left(\frac{-Nt^2}{2(K_1(s) + tK_2(s)/3)}\right),$$

and solve it for  $N$ :

$$N \geq \frac{2}{t^2} (K_1(s) + tK_2(s)/3) \ln(8d(s)/\delta).$$

With probability at least  $1 - \delta$ , the normwise absolute error is bounded above by  $\|\mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}})\|_\infty \leq t$ . Set  $t = \epsilon \|\mathcal{D}(\mathbf{A})\|_\infty$  to convert the absolute error to a relative one.  $\square$

*Interpretation.* As in Corollary 3.3, the constants  $\Delta_1(s)$  and  $\Delta_2(s)$  are respective relative counterparts of  $K_1(s)$  and  $K_2(s)$  in Theorem 3.5, namely, the relative deviation of  $\mathbf{A}$  from diagonality and more specifically the degree of diagonal dominance of  $\mathbf{A}$  in the relative sense. Corollary 3.6 implies that if  $\mathbf{A}$  is strongly diagonally dominant in the relative sense, then a small sampling amount suffices to achieve a user-specified error with a user-specified probability.

Furthermore, Corollary 3.6 suggests that increasing the sparsity parameter  $s$  could on the one hand lower the computational cost per sample but on the other increase the sampling amount for the same accuracy.

**4. Gaussian vectors.** We present normwise bounds for random vectors with bounded fourth moment (section 4.1) and standard Gaussian vectors (section 4.2).

**4.1. Random vectors with bounded fourth moment.** We bound the expectation of the squared absolute error (Theorem 4.1) for Monte Carlo estimators based on random vectors  $\mathbf{w}_k$  with independent entries that have zero mean, variance 1, and bounded fourth moment:

$$(4.1) \quad \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right] < +\infty.$$

These include standard Rademacher (section 3.1) and sparse Rademacher vectors (section 3.2) as well as standard Gaussian vectors (section 4.2).

**THEOREM 4.1.** *Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  with  $n \geq 3$  be symmetric, and let*

$$\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{j=1}^N \mathbf{A} \mathbf{w}_j \mathbf{w}_j^\top$$

*be a Monte Carlo estimator with independent random vectors  $\mathbf{w}_k \in \mathbb{R}^n$ ,  $1 \leq k \leq N$ , that have independent elements with zero mean, unit variance, and bounded fourth moment (4.1). Then*

$$\mathbb{E} \left[ \|\mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A})\|_p^2 \right]^{1/2} \leq \|\mathbf{A}\|_\infty \left( \sqrt{\frac{8e \ln n}{N}} + \frac{8e \ln n}{N} \right) \left( \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right] \right)^{1/2}.$$

*Proof.* We make use of matrix concentration inequalities but follow the spirit of the analysis in [6].

1. *Symmetrization.* Write the normwise error in terms of Schur products with diagonal matrices (2.1),

$$\mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A}) = \mathbf{I} \circ \hat{\mathbf{A}} - \mathbf{I} \circ \mathbf{A} = \frac{1}{N} \sum_{k=1}^N (\mathbf{I} \circ (\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top) - \mathbf{I} \circ \mathbf{A}),$$

and take expectations of the squared norms

$$\mathbb{E} \left[ \left\| \mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A}) \right\|_2^2 \right] = \frac{1}{N^2} \mathbb{E} \left[ \left\| \sum_{k=1}^N (\mathbf{I} \circ (\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top) - \mathbf{I} \circ \mathbf{A}) \right\|_2^2 \right].$$

From  $\mathbf{w}_k$  having independent elements with zero mean and unit variance follows that  $\mathbb{E}[\mathbf{I} \circ (\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top)] = \mathbf{I} \circ \mathbf{A}$ . Hence, the matrix random variables

$$\mathbf{X}_k \equiv \mathbf{I} \circ (\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top) - \mathbf{I} \circ \mathbf{A}, \quad 1 \leq k \leq N,$$

have zero mean. We use symmetrization [25, Lemma 6.4.2] to create *symmetric* random variables  $\varepsilon_k \mathbf{X}_k$ , where  $\varepsilon_k$  are independent symmetric Bernoulli random variables, in other words, Rademacher variables in Definition 3.1. The Rademacher variables  $\varepsilon_k$  are independent of each other and also independent of the random vectors  $\mathbf{w}_k$ . Remark 3.1 implies that  $\mathbb{E}[\varepsilon_k] = 0$ . Hence,

$$(4.2) \quad \mathbb{E} \left[ \sum_{k=1}^N \varepsilon_k \mathbf{I} \circ \mathbf{A} \right] = \mathbf{0}.$$

Then the symmetrization [25, Lemma 6.4.2], followed by (4.2) and (2.2), implies that

$$(4.3) \quad \begin{aligned} \mathbb{E} \left[ \left\| \mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A}) \right\|_2^2 \right] &= \frac{1}{N^2} \mathbb{E} \left[ \left\| \sum_{k=1}^N \mathbf{X}_k \right\|_2^2 \right] \leq 2 \mathbb{E} \left[ \left\| \sum_{k=1}^N \varepsilon_k \mathbf{I} \circ (\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top) \right\|_2^2 \right] \\ &= 2 \mathbb{E} \left[ \left\| \sum_{k=1}^N \mathbf{Y}_k \right\|_2^2 \right], \quad \text{where } \mathbf{Y}_k \equiv \varepsilon_k \text{diag}(\mathbf{A} \mathbf{w}_k) \text{diag}(\mathbf{w}_k) \end{aligned}$$

are symmetric random matrices and the second and third expectations above range over all random vectors  $\mathbf{w}_k$  and all Rademacher variables  $\varepsilon_k$ .

2. *Concentration inequality.* Applying the Cauchy–Schwartz inequality and Theorem 2.4 to the sum  $\mathbf{Z} \equiv \sum_{k=1}^N \mathbf{Y}_k$  gives

$$(4.4) \quad \mathbb{E}[\|\mathbf{Z}\|_2^2]^{1/2} \leq \sqrt{2e \ln n} \left\| \left( \sum_{i=1}^N \mathbb{E}[\mathbf{Y}_k^2] \right)^{1/2} \right\|_2 + 4e \ln n \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{Y}_k\|_2^2 \right]^{1/2}.$$

We bound separately the expectations that represent the matrix variance and the maximal 2-norm.

3. *Variance.* As in item 3 of the proof of Theorem 3.5, abbreviate

$$\mathbf{D}_k \equiv \text{diag}(\mathbf{A} \mathbf{w}_k), \quad \mathbf{W}_k \equiv \text{diag}(\mathbf{w}_k), \quad \mathbf{O} \equiv \text{diag}([\|\mathbf{a}_1\|_1 \quad \cdots \quad \|\mathbf{a}_n\|_1])$$

so that  $\mathbb{E}[\mathbf{Y}_k^2] = \mathbb{E}[\mathbf{D}_k^2 \mathbf{W}_k^2]$ . Apply the Hölder inequality [10, equation (2.2.2)] to the diagonal elements:

$$(\mathbf{D}_k^2 \mathbf{W}_k^2)_{ii} = \left( \sum_{j=1}^n a_{ij}(\mathbf{w}_k)_j \right)^2 (\mathbf{w}_k)_i^2 \leq \|\mathbf{a}_i\|_1^2 \|\mathbf{w}_k\|_\infty^4, \quad 1 \leq i \leq n.$$

For the diagonal matrices as a whole, this implies that  $\mathbf{D}_k^2 \mathbf{W}_k^2 \preceq \|\mathbf{w}_k\|_\infty^4 \mathbf{O}^2$ , and the symmetry of  $\mathbf{A}$  gives  $\mathbf{O} \preceq \|\mathbf{A}\|_\infty \mathbf{I}$ . Combine the two inequalities in the expectation,

$$(4.5) \quad \mathbb{E}[\mathbf{Y}_k^2] = \mathbb{E}[\mathbf{D}_k^2 \mathbf{W}_k^2] \preceq \|\mathbf{A}\|_\infty^2 \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right] \mathbf{I};$$

take the square root of the sum,

$$\left( \sum_{k=1}^N \mathbb{E}[\mathbf{Y}_k^2] \right)^{1/2} \preceq \sqrt{N} \|\mathbf{A}\|_\infty \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right]^{1/2} \mathbf{I};$$

and bound the norm,

$$(4.6) \quad \left\| \left( \sum_{k=1}^N \mathbb{E}[\mathbf{Y}_k^2] \right)^{1/2} \right\|_2 \leq \sqrt{N} \|\mathbf{A}\|_\infty \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right]^{1/2}.$$

4. *Maximal 2-norm.* In analogy to (4.5), we derive

$$\mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{Y}_k\|_2^2 \right] \leq \|\mathbf{A}\|_\infty^2 \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right]$$

and its square root

$$(4.7) \quad \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{Y}_k\|_2^2 \right]^{1/2} \leq \|\mathbf{A}\|_\infty \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right]^{1/2}.$$

5. *Putting everything together.* Substitute the variance bound (4.6) and the norm bound (4.7) into the expectation (4.4) for the sum

$$\mathbb{E}[\|\mathbf{Z}\|_2]^{1/2} \leq \left( \sqrt{2e \ln n} N^{1/2} + 4e \ln n \right) \|\mathbf{A}\|_\infty \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right]^{1/2};$$

substitute this, in turn, into the expectation (4.3) for the absolute error,

$$\begin{aligned} \mathbb{E} \left[ \|\mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A})\|_2^2 \right]^{1/2} &\leq \frac{2}{N} \mathbb{E}[\|\mathbf{Z}\|_2]^{1/2} \\ &\leq \frac{2}{N} \left( \sqrt{2e \ln n} N^{1/2} + 4e \ln n \right) \|\mathbf{A}\|_\infty \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right]^{1/2}; \end{aligned}$$

and simplify. Exploit the fact that all the  $p$ -norms are the same for diagonal matrices.  $\square$

*Interpretation.* Theorem 4.1 bounds the expected absolute error of Monte Carlo estimators based on a large class of random vectors. The bound depends logarithmically on the matrix dimension. The expected error in the Monte Carlo estimator decreases with more sampling, a decrease in matrix norm, or a decrease in the fourth moment of the random vectors.

**4.2. Gaussian vectors.** Gaussian vectors are a special case of random vectors (4.1) whose elements have bounded fourth moment. For Monte Carlo estimators based on Gaussian vectors, we determine the minimal sampling amount to achieve a user-specified error  $\epsilon$  at a user-specified success probability  $1 - \delta$ .

COROLLARY 4.2. Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  with  $n \geq 3$  be symmetric, and let

$$\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{j=1}^N \mathbf{A} \mathbf{w}_j \mathbf{w}_j^\top$$

be a Monte Carlo estimator with independent Gaussian random vectors  $\mathbf{w}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$  in  $\mathbb{R}^n$ ,  $1 \leq k \leq N$ . Pick  $\epsilon > 0$  and  $0 < \delta < 1$ . If the sampling amount  $N$  satisfies  $8e \ln n \leq N \leq n$  and is at least

$$(4.8) \quad N \geq \frac{128(e \ln n)^3}{\epsilon^2 \delta} \left( \frac{\|\mathbf{A}\|_\infty}{\|\mathcal{D}(\mathbf{A})\|_\infty} \right)^2,$$

then  $\|\mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A})\|_\infty \leq \epsilon \|\mathcal{D}(\mathbf{A})\|_\infty$  holds with probability at least  $1 - \delta$ .

*Proof.* For Gaussian random vectors  $\mathbf{w}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ , [6, (3.7)] implies that

$$\left( \mathbb{E} \left[ \max_{1 \leq k \leq N} \|\mathbf{w}_k\|_\infty^4 \right] \right)^{1/2} \leq e \ln(nN) \max_{1 \leq i, j \leq n} |\mathbf{I}_{ij}| = e \ln(nN).$$

Set  $p = \infty$ . Substituting this into Theorem 4.1 gives

$$\left( \mathbb{E} \left[ \|\mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A})\|_\infty^2 \right] \right)^{1/2} \leq \left( \sqrt{\frac{8e \ln n}{N}} + \frac{8e \ln n}{N} \right) e \ln(nN) \|\mathbf{A}\|_\infty.$$

Square both sides, and apply Markov's inequality (Theorem 2.1) to the random variable  $Z \equiv \|\mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A})\|_\infty$  using  $t \equiv \epsilon \|\mathcal{D}(\mathbf{A})\|_\infty$ :

$$(4.9) \quad \mathbb{P} \left[ \|\mathcal{D}(\hat{\mathbf{A}}) - \mathcal{D}(\mathbf{A})\|_\infty \geq \epsilon \|\mathcal{D}(\mathbf{A})\|_\infty \right] \leq \left( \sqrt{\frac{8e \ln n}{N}} + \frac{8e \ln n}{N} \right)^2 (e \ln(nN))^2 \frac{\|\mathbf{A}\|_\infty^2}{\epsilon^2 \|\mathcal{D}(\mathbf{A})\|_\infty^2}.$$

Substituting the assumption  $8e \ln n \leq N \leq n$  into the relevant part of the above bound gives

$$\left( \sqrt{\frac{8e \ln n}{N}} + \frac{8e \ln n}{N} \right)^2 (e \ln(nN))^2 \leq \left( 2\sqrt{\frac{8e \ln n}{N}} \right)^2 (2e \ln n)^2 = \frac{128(e \ln n)^3}{N}.$$

Substitute this, in turn, into (4.9); set the failure probability equal to

$$\delta \equiv \frac{128(e \ln n)^3}{\epsilon^2 N} \left( \frac{\|\mathbf{A}\|_\infty}{\|\mathcal{D}(\mathbf{A})\|_\infty} \right)^2;$$

and solve for the sampling amount  $N$ .  $\square$

**5. Elementwise bounds.** We present elementwise bounds for Monte Carlo estimators based on standard Rademacher vectors (section 5.1) as well as on standard Gaussian (section 5.2) and normalized Gaussian vectors (section 5.3).

**5.1. Standard Rademacher vectors.** We start with an elementwise worst-case absolute error bound (Corollary 5.1), consider products of Rademacher variables (Lemma 5.2), and present a bound on the absolute error of individual diagonal elements (Theorem 5.3), followed by a bound on the minimal sampling amount required for a specific diagonal element (Corollary 5.4).

We express Theorem 3.2 as a worst-case elementwise bound.

COROLLARY 5.1. Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric, and let

$$K_1 \equiv \|\mathcal{D}(\mathbf{A}^2) - \mathcal{D}(\mathbf{A})^2\|_p, \quad K_2 \equiv \|\mathbf{A} - \mathcal{D}(\mathbf{A})\|_\infty, \quad d \equiv \frac{\text{trace}(\mathcal{D}(\mathbf{A}^2) - \mathcal{D}(\mathbf{A})^2)}{K_1}.$$

If  $\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top$  is a Monte Carlo estimator with independent Rademacher vectors  $\mathbf{w}_k \in \mathbb{R}^n$ ,  $1 \leq k \leq N$ , then the probability that the absolute error exceeds  $t > 0$  is at most

$$\mathbb{P} \left[ \max_{1 \leq i \leq n} |a_{ii} - \hat{a}_{ii}| \geq t \right] \leq 8d \exp \left( \frac{-Nt^2}{2(K_1 + tK_2/3)} \right).$$

*Proof.* In Theorem 3.2, the  $p$ -norm of the diagonal matrix  $\mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}})$  is a largest magnitude diagonal element.  $\square$

Given  $n$  independent Rademacher variables, pick a  $(n+1)$ st independent variable, and multiply it with each of the  $n$  variables. We show below that the resulting  $n$  products are again independent.

LEMMA 5.2. If  $Z_1, W_1, W_2, \dots, W_n$  are independent Rademacher variables, then the products  $X_1 \equiv ZW_1, \dots, X_n \equiv ZW_n$  are also independent Rademacher variables.

*Proof.* For any  $x_1, \dots, x_n \in \{-1, +1\}$ , the law of total probability implies that the joint probability mass function satisfies

$$\begin{aligned} \mathbb{P}[\cap_{i=1}^n \{X_i = x_i\}] &= \sum_{z \in \{-1, +1\}} \mathbb{P}[\cap_{i=1}^n \{X_i = x_i\} | Z = z] \mathbb{P}[Z = z] \\ &= \frac{1}{2} \mathbb{P}[\cap_{i=1}^n \{W_i = -x_i\}] + \frac{1}{2} \mathbb{P}[\cap_{i=1}^n \{W_i = x_i\}] \\ &= 2^{-n} = \prod_{i=1}^n \mathbb{P}[X_i = x_i]. \end{aligned}$$

This factorization of the joint probability mass function implies that  $X_1, \dots, X_n$  are independent.  $\square$

In contrast to Corollary 5.1, the following bound for the diagonal element  $a_{ii}$  depends only on row/column  $i$ . In the special case where all off-diagonal elements in row/column  $i$  are zero, the estimator recovers  $a_{ii}$  with a single sample.

THEOREM 5.3. Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric, and let

$$\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top$$

be a Monte Carlo estimator with independent Rademacher vectors  $\mathbf{w}_k \in \mathbb{R}^n$ ,  $1 \leq k \leq N$ . The probability that the absolute error exceeds  $t > 0$  is at most

$$(5.1) \quad \mathbb{P}[|\hat{a}_{ii} - a_{ii}| \geq t] \leq 2 \exp \left( \frac{-Nt^2}{2(\|\mathbf{a}_i\|_2^2 - a_{ii}^2)} \right), \quad 1 \leq i \leq n.$$

In the special case where  $a_{ij} = 0$  for  $j \neq i$ , we have  $\hat{a}_{ii} = a_{ii}$  for all  $N \geq 1$ .

*Proof.* Fix  $i$  for some  $1 \leq i \leq n$ . Remark 3.1 allow us to split off the diagonal element from the estimator as follows:

$$\begin{aligned}\hat{a}_{ii} &= \frac{1}{N} \sum_{k=1}^N (\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top)_{ii} = \frac{1}{N} \sum_{k=1}^N \sum_{j=1}^n a_{ij} (\mathbf{w}_k)_j (\mathbf{w}_k)_i \\ &= a_{ii} + \sum_{k=1}^N \sum_{j \neq i} \underbrace{\frac{a_{ij}}{N} (\mathbf{w}_k)_j (\mathbf{w}_k)_i}_{Z_{kj}^{(i)}}.\end{aligned}$$

Thus, if  $a_{ij} = 0$  for  $i \neq j$ , then  $\hat{a}_{ii} = a_{ii}$  for all  $N \geq 1$ .

Lemma 5.2 implies that for fixed  $i$ , the  $Z_{kj}^{(i)}$  are independent, while Remark 3.1 implies that they have zero mean and are bounded by

$$-\frac{|a_{ij}|}{N} \leq Z_{kj}^{(i)} \leq \frac{|a_{ij}|}{N}, \quad 1 \leq k \leq N, \quad 1 \leq j \leq n, \quad j \neq i.$$

Hence, the absolute error  $\hat{a}_{ii} - a_{ii}$  is a sum of independent bounded zero-mean random variables, and we can apply Hoeffding's inequality in Theorem 2.2:

$$\mathbb{P}[|\hat{a}_{ii} - a_{ii}| \geq t] \leq 2 \exp \left( \frac{-2t^2}{\sum_{k=1}^N \sum_{j \neq i} \left( \frac{2}{N} |a_{ij}| \right)^2} \right) = 2 \exp \left( \frac{-Nt^2}{2 \sum_{j \neq i} a_{ij}^2} \right).$$

Since the matrix elements are real, we can write  $\sum_{j \neq i} a_{ij}^2 = \|\mathbf{a}_i\|_2^2 - a_{ii}^2$ .  $\square$

*Interpretation.* Theorem 5.3 implies that the accuracy of any estimated diagonal element  $\hat{a}_{ii}$  depends only on the magnitude of the off-diagonal elements in row/column  $i$ . The smaller the off-diagonal mass in row/column  $i$ , the smaller the probability that  $\hat{a}_{ii}$  has a large error. In the special case where  $a_{ii}$  is the only nonzero element in row/column  $i$ , the estimator recovers it in a single sample.

Note that the estimator is oblivious about the desired diagonal elements and that the bound in Theorem 5.3 applies to all diagonal elements. However, if one wants to estimate only a single diagonal element  $a_{kk}$ , then it can be extracted exactly with the targeted matrix vector product  $\mathbf{A} \mathbf{e}_k$ .

However, if one wants to estimate several diagonal elements, then the sampling amounts required to achieve a user-specified relative error  $\epsilon$  for a specific element depend on its associated off-diagonal mass. The bound below coincides with [3, (40)], but the proof is different.

**COROLLARY 5.4.** *Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric, and let*

$$\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top$$

*be a Monte Carlo estimator with independent Rademacher vectors  $\mathbf{w}_k \in \mathbb{R}^n$ ,  $1 \leq k \leq N$ . Pick a diagonal element  $a_{ii} \neq 0$ ,  $\epsilon > 0$ , and  $0 < \delta < 1$ . If the sampling amount is at least*

$$N \geq \left( \frac{\|\mathbf{a}_i\|_2^2 - a_{ii}^2}{a_{ii}^2} \right) \frac{2 \ln(2/\delta)}{\epsilon^2},$$

*then  $|a_{ii} - \hat{a}_{ii}| \leq \epsilon |a_{ii}|$  holds with probability at least  $1 - \delta$ .*

*Proof.* Define the 2-norm off-diagonal column sum for diagonal element  $a_{ii}$  by

$$\text{off}_i \equiv (\|\mathbf{a}_i\|_2^2 - a_{ii}^2)^{1/2},$$

denote the bound for the failure probability in Theorem 5.3 by

$$\delta \equiv 2 \exp\left(\frac{-Nt^2}{2\text{off}_i^2}\right),$$

and solve it for  $t$ ,

$$t = \sqrt{\frac{2\text{off}_i^2}{N} \ln(2/\delta)}.$$

Restate Theorem 5.3 in terms of the failure probability: With probability at most  $1 - \delta$ , the absolute error in diagonal element  $a_{ii}$  is bounded by

$$|a_{ii} - \hat{a}_{ii}| \leq t = \sqrt{\frac{2\text{off}_i^2}{N} \ln(2/\delta)}.$$

Converting the absolute error to a relative error requires  $t \leq \epsilon|a_{ii}|$ , which implies that

$$N \geq \left(\frac{\text{off}_i}{a_{ii}}\right)^2 \frac{2 \ln(2/\delta)}{\epsilon^2}. \quad \square$$

*Interpretation.* The minimal sampling amount for the estimated  $\hat{a}_{ii}$  depends on the relative off-diagonal mass of column/row  $i$ . The smaller this off-diagonal mass, the lower the sampling amount to achieve the user-specified relative error  $\epsilon$  with the user-specified success probability  $1 - \delta$ .

**5.2. Gaussian vectors.** We present an elementwise absolute error bound (Theorem 5.5) for Gaussian-based Monte Carlo estimators and a bound on the minimal sampling amount to achieve a user-specified relative error at a user-specified probability (Corollary 5.6). The bounds are derived from and identical to bounds for trace estimators in [8].

**THEOREM 5.5.** *Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric,*

$$L_{1i} \equiv |a_{ii}| + \|\mathbf{a}_i\|_2, \quad L_{i2} \equiv |a_{ii}|^2 + \|\mathbf{a}_i\|_2^2, \quad 1 \leq i \leq n,$$

*and let  $\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top \in \mathbb{R}^{n \times n}$  be a Monte Carlo estimator with independent Gaussian vectors  $\mathbf{w}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$  in  $\mathbb{R}^n$ ,  $1 \leq k \leq N$ . If  $t > 0$ , then*

$$(5.2) \quad \mathbb{P}[|\hat{a}_{ii} - a_{ii}| > t] \leq 2 \exp\left(\frac{-Nt^2}{2(L_{i2} + t L_{i1})}\right), \quad 1 \leq i \leq n.$$

*Proof.* Fix  $i$  for some  $1 \leq i \leq n$ . Write the diagonal element as an inner product

$$(\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top)_{ii} = \sum_{j=1}^n a_{ij}(\mathbf{w}_k)_j (\mathbf{w}_k)_i = \mathbf{w}_k^\top \mathbf{B}_i \mathbf{w}_k, \quad 1 \leq k \leq N,$$



involving the symmetric matrix

$$\mathbf{B}_i \equiv \begin{bmatrix} 0 & \frac{1}{2}a_{1i} & 0 \\ \frac{1}{2}a_{i1} & \cdots & a_{ii} & \cdots & \frac{1}{2}a_{in} \\ 0 & \frac{1}{2}a_{ni} & 0 \end{bmatrix} \in \mathbb{R}^{n \times n} \quad \text{with} \quad \text{trace}(\mathbf{B}_i) = a_{ii}.$$

We can interpret

$$\hat{a}_{ii} = \left( \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top \right)_{ii} = \frac{1}{N} \sum_{k=1}^N (\mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top)_{ii} = \frac{1}{N} \sum_{k=1}^N \mathbf{w}_k^\top \mathbf{B}_i \mathbf{w}_k$$

as a Monte Carlo estimator for  $\text{trace}(\mathbf{B}_i)$  and apply the bound for Gaussian trace estimators [8, Theorem 1]

$$\mathbb{P}[|\hat{a}_{ii} - a_{ii}| \geq t] \leq 2 \exp \left( \frac{-Nt^2}{4\|\mathbf{B}_i\|_F^2 + 4t\|\mathbf{B}_i\|_2} \right),$$

where

$$\|\mathbf{B}_i\|_F^2 = \frac{1}{2} (|a_{ii}|^2 + \|\mathbf{a}_i\|_2^2) = \frac{L_{i2}}{2}, \quad \|\mathbf{B}_i\|_2 = \frac{1}{2} (|a_{ii}| + \|\mathbf{a}_i\|_2) = \frac{L_{i1}}{2}. \quad \square$$

*Interpretation.* Theorem 5.5 implies that the accuracy of any single diagonal element  $\hat{a}_{ii}$  from a Gaussian-based Monte Carlo estimator depends on the norm of the corresponding row/column. Thus, the smaller the norm of row/column  $i$ , the smaller the probability that  $\hat{a}_{ii}$  has a large error.

In contrast, the bounds for Rademacher-based Monte Carlo estimators in Theorem 5.3 depend only on the magnitude of the off-diagonal elements.

**COROLLARY 5.6.** *Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric,*

$$\Delta_{1i} \equiv 1 + \frac{\|\mathbf{a}_i\|_2}{|a_{ii}|}, \quad \Delta_{i2} \equiv 1 + \left( \frac{\|\mathbf{a}_i\|_2}{|a_{ii}|} \right)^2, \quad 1 \leq i \leq n,$$

and let  $\hat{\mathbf{A}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{A} \mathbf{z}_k \mathbf{z}_k^\top \in \mathbb{R}^{n \times n}$  be a Monte Carlo estimator with independent Gaussian vectors  $\mathbf{z}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$  in  $\mathbb{R}^n$ ,  $1 \leq k \leq N$ . Pick  $a_{ii} \neq 0$ ,  $\epsilon > 0$ , and  $0 < \delta < 1$ . If the sampling amount is at least

$$N \geq (\Delta_{2i} + \Delta_{1i}\epsilon) \frac{2 \ln(2/\delta)}{\epsilon^2},$$

then  $|\hat{a}_{ii} - a_{ii}| \leq \epsilon |a_{ii}|$  holds with probability at least  $1 - \delta$ .

*Proof.* This follows immediately from the lower bound for  $N$  in [8, Theorem 1].  $\square$

The required sampling amount for computing  $a_{ii}$  with the Gaussian Monte Carlo estimator depends on  $\|\mathbf{a}_i\|_2/|a_{ii}|$ , which can be interpreted as the 2-norm deviation of column/row  $i$  of  $\mathbf{A}$  from diagonality. The more diagonal row/column  $i$ , the smaller the sampling amount.

**5.3. Normalized Gaussian vectors.** We extend and complete the analysis in [4] for Monte Carlo estimators based on normalized Gaussian vectors,

$$(5.3) \quad \hat{\mathbf{A}} \equiv \left( \sum_{k=1}^N \mathbf{A} \mathbf{w}_k \mathbf{w}_k^\top \right) \oslash \left( \sum_{k=1}^N \mathbf{w}_k \mathbf{w}_k^\top \right),$$

where  $\mathbf{w}_k \in \mathbb{R}^n$  are independent Gaussian random vectors and  $\oslash$  denotes elementwise division. We derive the distribution of the elementwise absolute errors (Lemma 5.7) followed by a bound (Theorem 5.8).

We represent the distribution of the absolute errors in the diagonal elements in terms of a *Student t-distribution* with  $N \geq 1$  degrees of freedom [16, Definition 7.3.3],

$$(5.4) \quad T_N \equiv \frac{Z}{\sqrt{U/N}},$$

where  $Z$  is a Gaussian  $\mathcal{N}(0,1)$  random variable and  $U$  an independent chi-square random variable with  $N$  degrees of freedom.

LEMMA 5.7. *Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be symmetric and (5.3) be a Monte Carlo estimator with independent Gaussian vectors  $\mathbf{w}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$  in  $\mathbb{R}^n$ ,  $1 \leq k \leq N$ . The absolute errors in the diagonal elements are distributed as*

$$\hat{a}_{ii} - a_{ii} \sim \sqrt{\frac{\|\mathbf{a}_i\|_2^2 - a_{ii}^2}{N}} T_N, \quad 1 \leq i \leq n.$$

*Proof.* Due to the normalization in the denominator, we can extract the diagonal elements of  $\mathbf{A}$  from the diagonal elements of the Monte Carlo estimator  $\hat{\mathbf{A}}$ :

$$\hat{a}_{ii} = \frac{\sum_{k=1}^N \sum_{j=1}^n a_{ij} (\mathbf{w}_k)_i (\mathbf{w}_k)_j}{\sum_{k=1}^N (\mathbf{w}_k)_i^2} = a_{ii} + \frac{\sum_{k=1}^N \sum_{j \neq i} a_{ij} (\mathbf{w}_k)_i (\mathbf{w}_k)_j}{\sum_{k=1}^N (\mathbf{w}_k)_i^2}, \quad 1 \leq i \leq n.$$

Normalize across the  $i$ th elements of the Gaussian vectors  $\mathbf{w}_k$  to unit vectors  $\mathbf{u}_i \in \mathbb{R}^N$  with elements

$$(\mathbf{u}_i)_k \equiv \frac{(\mathbf{w}_k)_i}{\sqrt{\sum_{\ell=1}^N (\mathbf{w}_\ell)_i^2}}, \quad 1 \leq k \leq N, \quad 1 \leq i \leq n.$$

Use the denominator to normalize the  $i$ th component in the  $i$ th absolute error:

$$(5.5) \quad \hat{a}_{ii} - a_{ii} = \frac{\sum_{k=1}^N \sum_{j \neq i} a_{ij} (\mathbf{w}_k)_j (\mathbf{u}_i)_k}{\sqrt{\sum_{\ell=1}^N (\mathbf{w}_\ell)_i^2}}, \quad 1 \leq i \leq n.$$

The rotational invariance of the standard Gaussian distribution guarantees the independence of the direction vectors  $\mathbf{u}_i$  and radial components  $\sqrt{\sum_{k=1}^N (\mathbf{w}_k)_i^2}$ ; see [25, Exercise 3.3.6]. Hence, the numerator and denominator in (5.5) are independent. With

$$Z_i \equiv \frac{1}{\sqrt{N}} \sum_{k=1}^N \sum_{j \neq i} a_{ij} (\mathbf{w}_k)_j (\mathbf{u}_i)_k, \quad U_i \equiv \sum_{k=1}^N (\mathbf{w}_k)_i^2, \quad 1 \leq i \leq n,$$

we can rewrite (5.5) as

$$\hat{a}_{ii} - a_{ii} = \frac{Z_i}{\sqrt{U_i/N}}.$$

The random variable  $U_i$  has a chi-square distribution with  $N$  degrees of freedom. The conditional distribution of  $Z_i$  given  $\mathbf{u}_i$  is Gaussian [25, Exercise 3.3.3(a)] with zero mean and variance

$$\frac{1}{N} \sum_{k=1}^N \sum_{j \neq i} a_{ij}^2 (\mathbf{u}_i)_k^2 = \frac{1}{N} \left( \sum_{j \neq i} a_{ij}^2 \right) \left( \sum_{k=1}^N (\mathbf{u}_i)_k^2 \right) = \frac{1}{N} \sum_{j \neq i} a_{ij}^2 = \frac{1}{N} (\|\mathbf{a}_i\|_2^2 - a_{ii}^2).$$

Therefore,  $Z_i | \mathbf{u}_i \sim \mathcal{N}(0, \frac{1}{N} (\|\mathbf{a}_i\|_2^2 - a_{ii}^2))$ . However, the conditional distribution is independent of  $\mathbf{u}_i$ , so this is also the unconditional distribution. The claim then follows from (5.4).  $\square$

*Remark 5.1.* The squared error  $(\hat{a}_{ii} - a_{ii})^2$  has a scaled  $F$ -distribution [3]. Note that the squared Student  $t$ -distribution is specifically a scaled  $F$ -distribution with one degree of freedom in the numerator. Moreover, for a single sample  $N = 1$ , the error has a Cauchy distribution, which has undefined mean and variance.

If  $N$  is large, then the  $t$ -distribution  $T_N$  can be approximated by a standard normal distribution. However,  $T_N$  has wider tails and thus somewhat weaker tail bounds. Existing tail bounds for the Student  $t$ -distribution imply the following concentration inequality.

**THEOREM 5.8.** *Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be symmetric and (5.3) be a Monte Carlo estimator with independent Gaussian vectors, where  $\mathbf{w}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ ,  $1 \leq k \leq N$ . The probability that the absolute error exceeds  $t > 0$  is at most*

$$\mathbb{P}[|\hat{a}_{ii} - a_{ii}| > t] \leq \sqrt{\frac{2(\|\mathbf{a}_i\|_2^2 - a_{ii}^2)}{\pi N}} \frac{1}{t} \left( 1 + \frac{t^2}{\|\mathbf{a}_i\|_2^2 - a_{ii}^2} \right)^{-\frac{N-1}{2}}, \quad 1 \leq i \leq n.$$

*Proof.* The probability density function of  $T_N$  is [20]

$$f_N(x) = c_N \left( 1 + \frac{x^2}{N} \right)^{-\frac{N+1}{2}}, \quad \text{where } c_N \equiv \frac{\Gamma((N+1)/2)}{\Gamma(N/2)\sqrt{N\pi}}$$

and  $1/\pi \leq c_N \leq 1/\sqrt{2\pi}$ . Let  $F_N(x)$  be the cumulative distribution function for  $T_N$ . This implies with [20, Theorem 3.1],

$$\begin{aligned} \mathbb{P}\left[|\hat{a}_{ii} - a_{ii}| > \sqrt{\frac{\|\mathbf{a}_i\|_2^2 - a_{ii}^2}{N}} x\right] &= 1 - F_N(x) < \frac{f_N(x)}{x} \left( 1 + \frac{x^2}{N} \right) \\ &= \frac{c_N}{x} \left( 1 + \frac{x^2}{N} \right)^{-\frac{N-1}{2}}, \quad 1 \leq i \leq n. \end{aligned}$$

Setting  $x = t \sqrt{\frac{N}{\|\mathbf{a}_i\|_2^2 - a_{ii}^2}}$  and bounding the upper tail with  $c_N \leq \sqrt{\frac{1}{2\pi}}$  gives

$$\mathbb{P}[|\hat{a}_{ii} - a_{ii}| > t] \leq \sqrt{\frac{\|\mathbf{a}_i\|_2^2 - a_{ii}^2}{\pi N}} \frac{1}{\sqrt{2}t} \left( 1 + \frac{t^2}{\|\mathbf{a}_i\|_2^2 - a_{ii}^2} \right)^{-\frac{N-1}{2}}, \quad 1 \leq i \leq n.$$

Since  $T_N$  is symmetric about the origin, the lower tail has the same bound. Now take a union bound over the two tails.  $\square$

We determine a sampling amount sufficient to make the normalized Gaussian Monte Carlo estimator a componentwise  $(\epsilon, \delta)$  estimator.

COROLLARY 5.9. Let  $\mathbf{A} \in \mathbb{R}^{n \times n}$  be nondiagonal symmetric; let

$$\Psi_i \equiv \frac{|a_{ii}|}{(\|\mathbf{a}_i\|_2^2 - a_{ii}^2)^{1/2}}, \quad 1 \leq i \leq n;$$

and let  $\hat{\mathbf{A}}$  be defined as in (5.3). Pick  $\epsilon > 0$  and a diagonal element  $a_{ii} \neq 0$  of  $\mathbf{A}$ . For any  $0 < \delta < 1$ , if the sampling amount is positive and at least

$$N \geq 1 + 2 \ln \left( \frac{\sqrt{2/\pi}}{\delta \epsilon \Psi_i} \right) / \ln(1 + \epsilon^2 \Psi_i^2),$$

then  $|\hat{a}_{ii} - a_{ii}| \leq \epsilon |a_{ii}|$  holds with probability at least  $1 - \delta$ .

*Proof.* In Theorem 5.8, set  $t = \epsilon |a_{ii}|$ . If the sampling number satisfies the desired bound, it follows that

$$\begin{aligned} \mathbb{P}[(\hat{a}_{ii} - a_{ii}) > t] &\leq \sqrt{\frac{2(\|\mathbf{a}_i\|_2^2 - a_{ii}^2)}{\pi N}} \frac{1}{t} \left( 1 + \frac{t^2}{\|\mathbf{a}_i\|_2^2 - a_{ii}^2} \right)^{-\frac{N-1}{2}} \\ &= \frac{\sqrt{2/\pi}}{\epsilon \Psi_i \sqrt{N}} (1 + \epsilon^2 \Psi_i^2)^{-\frac{N-1}{2}} \\ &\leq \frac{\sqrt{2/\pi}}{\epsilon \Psi_i} (1 + \epsilon^2 \Psi_i^2)^{-\frac{N-1}{2}}, \end{aligned}$$

where the final inequality holds since  $N \geq 1$  by assumption. Set the failure probability  $\delta$  to the right-hand side as

$$\delta \equiv \frac{\sqrt{2/\pi}}{\epsilon \Psi_i} (1 + \epsilon^2 \Psi_i^2)^{-\frac{N-1}{2}},$$

and solve for  $N$ . □

The larger the value of  $\Psi_i$  (the same measure of diagonal dominance that appears in Corollary 5.4), the smaller the sampling amount for the  $(\epsilon, \delta)$  estimator.

**6. Application: Monte Carlo estimators for a derivative-based global sensitivity metric.** We bound the absolute error (Theorem 6.1) in a Monte Carlo estimator for global sensitivity analysis and more specifically for a derivative-based global sensitivity metric (DGSM) of a function  $f: \mathbb{R}^n \rightarrow \mathbb{R}$ , whose partial derivatives are square integrable with respect to a probability density function  $\rho_{\mathbf{x}}(\mathbf{x})$ . The DGSM is equal to the diagonal  $\mathcal{D}(\mathbf{C})$  of the matrix

$$\mathbf{C} = \int_{\mathcal{X}} \nabla f(\mathbf{x}) [\nabla f(\mathbf{x})]^\top \rho_{\mathbf{x}}(\mathbf{x}) d\mathbf{x}.$$

The matrix  $\mathbf{C}$  is well-defined and symmetric positive semidefinite and can be interpreted as a second moment matrix of the gradient. We compute the DGSM with the Monte Carlo estimator

$$(6.1) \quad \hat{\mathbf{C}} \equiv \frac{1}{N} \sum_{k=1}^N \mathbf{w}_k \mathbf{w}_k^\top, \quad \text{where } \mathbf{w}_k \equiv \nabla f(\mathbf{x}_k), \quad 1 \leq k \leq N,$$

and  $\mathbf{x}_k$  are independent samples from the distribution  $\rho_{\mathbf{x}}(\mathbf{x})$ . Below is a normwise bound for the error in the DGSM computed by the Monte Carlo estimator (6.1). Its derivation is related to the analysis in [15, section 4].

THEOREM 6.1. Let  $f : \mathbb{R}^n \rightarrow \mathbb{R}$  have square integrable partial derivatives with respect to the probability density function  $\rho_{\mathbf{x}}(\mathbf{x})$ ,  $\|\nabla f\|_{\infty} \leq \beta$  almost surely;  $\widehat{\mathbf{C}}$  be the Monte Carlo estimator in (6.1), and

$$\begin{aligned} c_{\max} &\equiv \|\mathcal{D}(\mathbf{C})\|_p, & S_1 &\equiv \|\mathcal{D}(\mathbf{C})(\beta^2 \mathbf{I} - \mathcal{D}(\mathbf{C}))\|_p, \\ S_2 &\equiv c_{\max} + \beta^2, & d &\equiv \frac{\sum_{i=1}^n c_{ii}(\beta^2 - c_{ii})}{S_1}. \end{aligned}$$

If  $c_{\max} > 0$  and  $S_1 > 0$ , then

$$\mathbb{P}\left[\|\mathcal{D}(\mathbf{C}) - \mathcal{D}(\widehat{\mathbf{C}})\|_p \geq t\right] \leq 8d \exp\left(\frac{-t^2/2}{S_1 + S_2 t/3}\right).$$

*Proof.* Before applying the matrix Bernstein inequality in Theorem 2.3, we need to verify the assumptions. The Monte Carlo estimate  $\mathcal{D}(\widehat{\mathbf{C}})$  is an unbiased estimator of the DGSM  $\mathcal{D}(\mathbf{C})$  whose largest diagonal element is

$$c_{\max} = \|\mathcal{D}(\mathbf{C})\|_2 = \max_{1 \leq i \leq n} |c_{ii}| = \max_{1 \leq i \leq n} \mathbb{E}[(\nabla f(\mathbf{x}))_i^2] \leq \beta^2.$$

The absolute error in the DGSM computed by the Monte Carlo estimator (6.1) is

$$\mathbf{Z} = \mathcal{D}(\mathbf{C}) - \mathcal{D}(\widehat{\mathbf{C}}) = \sum_{k=1}^N \mathbf{S}_k, \quad \text{where } \mathbf{S}_k \equiv \frac{1}{N} (\mathcal{D}(\mathbf{C}) - \mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top)).$$

The summands  $\mathbf{S}_k$  have zero mean and are bounded by

$$\begin{aligned} \|\mathbf{S}_k\|_2 &\leq \frac{1}{N} (\|\mathcal{D}(\mathbf{C})\|_2 + \|\mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top)\|_2) \\ &\leq \frac{\|\mathcal{D}(\mathbf{C})\|_2 + \beta^2}{N} = \frac{c_{\max} + \beta^2}{N} = \frac{S_2}{N}, \quad 1 \leq k \leq N. \end{aligned}$$

We let  $L = S_2/N$  so that  $\|\mathbf{S}_k\|_2 \leq L$ . The variance is

$$(6.2) \quad \mathbb{V}\text{ar}[\mathbf{Z}] = \sum_{k=1}^N \mathbb{E}[\mathbf{S}_k^2] = \frac{1}{N^2} \sum_{k=1}^N \mathbb{E}\left[(\mathcal{D}(\mathbf{C}) - \mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top))^2\right].$$

Linearity of the expectation, the majorization  $\mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top) \preceq \beta^2 \mathbf{I}$ , the expectation  $\mathbb{E}[\mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top)] = \mathcal{D}(\mathbf{C})$ , and commutativity of diagonal matrices imply for the summands

$$\begin{aligned} \mathbb{E}\left[(\mathcal{D}(\mathbf{C}) - \mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top))^2\right] &= \mathbb{E}\left[\mathcal{D}(\mathbf{C})^2 + (\mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top))^2 - 2\mathcal{D}(\mathbf{C})\mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top)\right] \\ &= \mathbb{E}\left[(\mathcal{D}(\mathbf{w}_k \mathbf{w}_k^\top))^2\right] - \mathcal{D}(\mathbf{C})^2 \\ &\preceq \beta^2 \mathcal{D}(\mathbf{C}) - \mathcal{D}(\mathbf{C})^2, \quad 1 \leq k \leq N. \end{aligned}$$

Substitute the above into (6.2),

$$\mathbb{V}\text{ar}[\mathbf{Z}] \preceq \mathbf{V} \equiv \frac{1}{N} (\beta^2 \mathcal{D}(\mathbf{C}) - \mathcal{D}(\mathbf{C})^2),$$

and apply Theorem 2.3 with  $\nu = \|\mathbf{V}\|_2 = S_1/N$  and  $d = \text{intdim}(\mathbf{V})$ . Finally, use the fact that all the  $p$ -norms are all the same for diagonal matrices.  $\square$

Below is the minimal sampling amount required for the Monte Carlo estimator (6.1) to achieve a user-specified relative error  $\epsilon$  with a user-specified success probability  $1 - \delta$ .

COROLLARY 6.2. Let  $f: \mathbb{R}^n \rightarrow \mathbb{R}$  have square-integrable partial derivatives with respect to the probability density function  $\rho_{\mathbf{x}}(\mathbf{x})$ ,  $\|\nabla f\|_{\infty} \leq \beta$  almost surely;  $\hat{\mathbf{C}}$  be the Monte Carlo estimator in (6.1), and

$$\begin{aligned} c_{\max} &\equiv \|\mathcal{D}(\mathbf{C})\|_{\infty}, & S_1 &\equiv \|\mathcal{D}(\mathbf{C})(\beta^2 \mathbf{I} - \mathcal{D}(\mathbf{C}))\|_{\infty}, \\ S_2 &\equiv c_{\max} + \beta^2, & d &\equiv \frac{\sum_{i=1}^n c_{ii}}{S_1}. \end{aligned}$$

Pick  $\epsilon > 0$  and  $0 < \delta < 1$ . If  $c_{\max} > 0$  and  $S_1 > 0$  and if the sampling amount is at least

$$N \geq \frac{S_2}{3\epsilon^2} \left( 2\epsilon + \frac{6S_1}{c_{\max} S_2} \right) \ln(8d/\delta),$$

then  $\|\mathcal{D}(\mathbf{C}) - \mathcal{D}(\hat{\mathbf{C}})\|_{\infty} \leq \epsilon \|\mathcal{D}(\mathbf{C})\|_{\infty}$  holds with probability at least  $1 - \delta$ .

*Proof.* The proof is similar to that of Corollary 3.3 but is based instead on Theorem 6.1 with  $p = \infty$ .  $\square$

**6.1. Illustrative examples.** We determine the constants in Corollary 6.2 for two different functions  $f: \mathbb{R}^n \rightarrow \mathbb{R}$  and random variables  $\mathbf{X} \in \mathbb{R}^n$  from a uniform distribution over  $\mathcal{X} = [-1, 1]^n$ .

*Linear Function.* Let  $f(\mathbf{x}) = \mathbf{h}^{\top} \mathbf{x}$  with  $\mathbf{h} \in \mathbb{R}^n$ . Then  $\|\nabla f\|_{\infty} \leq \beta \equiv \|\mathbf{h}\|_{\infty}$  almost surely. The second moment matrix  $\mathbf{C} = \mathbf{h}\mathbf{h}^{\top}$  has a largest diagonal entry  $c_{\max} = \|\mathbf{h}\|_{\infty}^2$ . Let  $\mathbf{v} \in \mathbb{R}^n$  have entries  $v_i = h_i^2(\|\mathbf{h}\|_{\infty}^2 - h_i^2)$  for  $1 \leq i \leq n$ . The constants in Corollary 6.2 are

$$S_1 = \|\mathbf{v}\|_{\infty}, \quad S_2 = 2\|\mathbf{h}\|_{\infty}^2, \quad d = \frac{\sum_{i=1}^n v_i}{S_1}.$$

*Quadratic function.* Let  $f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^{\top} \mathbf{S} \mathbf{x}$ , where  $\mathbf{S} \in \mathbb{R}^{n \times n}$  with  $\mathbf{S} = \mathbf{S}^{\top}$  is a symmetric square root of the positive semidefinite matrix  $\mathbf{M} = \mathbf{S}^2$ . Then  $\|\nabla f\|_{\infty} \leq \beta \equiv \|\mathbf{S}\|_{\infty}$  almost surely. The second moment matrix  $\mathbf{C} = \frac{1}{3} \mathbf{M}$  has a largest diagonal element  $c_{\max} = \frac{1}{3} \|\mathcal{D}(\mathbf{M})\|_{\infty}$ . Let  $\mathbf{v} \in \mathbb{R}^n$  have entries  $v_i = \frac{1}{3} m_{ii}(\|\mathbf{S}\|_{\infty}^2 - \frac{1}{3} m_{ii})$  for  $1 \leq i \leq n$ . The constants in Corollary 6.2 are

$$S_1 = \|\mathbf{v}\|_{\infty}, \quad S_2 = \frac{1}{3} \|\mathcal{D}(\mathbf{M})\|_{\infty} + \|\mathbf{S}\|_{\infty}^2, \quad d = \frac{\sum_{i=1}^n v_i}{S_1}.$$

**7. Numerical Experiments.** After presenting our algorithm (section 7.1) and describing our test matrices (section 7.2), we present four different types of numerical experiments to illustrate the accuracy of the Monte Carlo estimators: Rademacher Monte Carlo estimators applied to the test matrices (section 7.3), accuracy of different Monte Carlo estimators (section 7.4), effect of the sparsity on the accuracy of Rademacher Monte Carlo estimators (section 7.5), and accuracy of the DGSM Monte Carlo estimator (section 7.6). Its applications to circuit models (section 7.7) and an application to node centrality (section 7.8).

**7.1. Algorithm.** We present pseudocodes for Monte Carlo estimators based on general random vectors (Algorithm 7.1) and for estimators based on normalized Gaussian vectors (Algorithm 7.2).

Algorithm 7.1 applies to standard and sparse Rademacher vectors and to vectors with bounded fourth moments, including Gaussian vectors.

Algorithm 7.2 requires a slight modification for normalized Gaussians.

---

**Algorithm 7.1** Monte Carlo diagonal estimation.

---

**Input:** Matrix  $\mathbf{A} \in \mathbb{R}^{n \times n}$ , sampling amount  $N$ , distribution  $\mathcal{S}$  on  $\mathbb{R}^n$

**Output:** Diagonal estimator  $\hat{\mathbf{a}}_{\text{diag}} \in \mathbb{R}^n$

```

1: Initialize  $\hat{\mathbf{a}}_{\text{diag}} = \mathbf{0}$ 
2: for  $k = 1 : N$  do
3:   Sample  $\mathbf{w}_k$  from  $\mathcal{S}$ 
4:    $\hat{\mathbf{a}}_{\text{diag}} = \hat{\mathbf{a}}_{\text{diag}} + (\mathbf{A}\mathbf{w}_k) \circ \mathbf{w}_k$ 
5: end for
6:  $\hat{\mathbf{a}}_{\text{diag}} = \frac{1}{N} \hat{\mathbf{a}}_{\text{diag}}$ 

```

---



---

**Algorithm 7.2** Monte Carlo diagonal estimation: Normalized Gaussian.

---

**Input:** Matrix  $\mathbf{A} \in \mathbb{R}^{n \times n}$ , sampling amount  $N$

**Output:** Diagonal estimator  $\hat{\mathbf{a}}_{\text{diag}} \in \mathbb{R}^n$

```

1: Initialize  $\hat{\mathbf{a}}_{\text{diag}} = \mathbf{0}$  and  $\hat{\mathbf{z}} = \mathbf{0}$ 
2: for  $k = 1 : N$  do
3:   Sample  $\mathbf{w}_k$  from  $\mathcal{N}(\mathbf{0}, \mathbf{I})$ 
4:    $\hat{\mathbf{a}}_{\text{diag}} = \hat{\mathbf{a}}_{\text{diag}} + (\mathbf{A}\mathbf{w}_k) \circ \mathbf{w}_k$ 
5:    $\hat{\mathbf{z}} = \hat{\mathbf{z}} + \mathbf{w}_k \circ \mathbf{w}_k$ 
6: end for
7:  $\hat{\mathbf{a}}_{\text{diag}} = \hat{\mathbf{a}}_{\text{diag}} \oslash \hat{\mathbf{z}}$ 

```

---

The computational cost of Algorithms 7.1 and 7.2 is similar. Denote by  $T_{\mathbf{A}}$  the cost in floating point operations (flops) of a matrix-vector product with  $\mathbf{A}$  and by  $T_{\mathcal{S}}$  the cost of generating a random vector from the distribution  $\mathcal{S}$  on  $\mathbb{R}^n$ . Both algorithms require  $N(T_{\mathbf{A}} + T_{\mathcal{S}} + \mathcal{O}(n))$  flops.

In contrast, the naïve approach, which computes matrix-vector products with the canonical basis vectors  $\mathbf{e}_j$  and extracts diagonal element  $a_{jj}$  from  $\mathbf{A}\mathbf{e}_j$ ,  $1 \leq j \leq n$ , requires  $nT_{\mathbf{A}}$  flops. Algorithms 7.1 and 7.2 are faster if  $N \ll n$  and  $T_{\mathcal{S}} \ll T_{\mathbf{A}}$ .

**7.2. Test Matrices.** We perform numerical experiments on three symmetric test matrices from [19] of dimension  $n = 100$  that depend on a parameter  $\theta$ . We use sampling amounts up to  $N = 10n$  samples merely to illustrate the features of the analysis and the algorithms. We present more realistic test cases in sections 7.7 and 7.8:

1. Identity plus rank-1,

$$\mathbf{A} = \mathbf{I} + \theta \mathbf{e} \mathbf{e}^{\top}, \quad \text{where } .01 \leq \theta \leq 0.1,$$

where  $\mathbf{e} \in \mathbb{R}^n$  is a vector of ones. The constants in Corollary 3.3 are

$$K_1 = (n-1)\theta^2, \quad K_2 = (n-1)\theta, \quad \|\mathcal{D}(\mathbf{A})\|_{\infty} = 1 + \theta$$

so that

$$\Delta_1 = \frac{K_1}{(1+\theta)^2}, \quad \Delta_2 = \frac{(n-1)\theta}{1+\theta}, \quad d = n.$$

2. Rank-1 with decaying elements

$$\mathbf{A} = \frac{\mathbf{x} \mathbf{x}^{\top}}{\|\mathbf{x}\|_2^2}, \quad \text{where } \mathbf{x}_j = e^{-j(1-\theta)}, \quad 1 \leq j \leq n, \quad 0.1 \leq \theta \leq 1.$$

The constants in Corollary 3.3 are

$$K_1 = \left( \frac{x_1}{\|\mathbf{x}\|_2} \right)^2 \left( 1 - \left( \frac{x_1}{\|\mathbf{x}\|_2} \right)^2 \right), \quad K_2 = \frac{x_1}{\|\mathbf{x}\|_2^2} \sum_{j>1} x_j,$$

and  $\|\mathcal{D}(\mathbf{A})\|_\infty = \left( \frac{x_1}{\|\mathbf{x}\|_2} \right)^2$  so that

$$\Delta_1 = \left( \frac{\|\mathbf{x}\|_2}{x_1} \right)^2 - 1, \quad \Delta_2 = \frac{\sum_{j>1} x_j}{x_1}, \quad d = \frac{\sum_{i=1}^n x_i^2 (\|\mathbf{x}\|_2^2 - x_i^2)}{x_1^2 (\|\mathbf{x}\|_2^2 - x_1^2)}.$$

3. Tridiagonal Toeplitz matrix,

$$\mathbf{A} = \begin{bmatrix} 1 & \theta & & \\ \theta & 1 & \ddots & \\ & \ddots & \ddots & \theta \\ & & \theta & 1 \end{bmatrix}, \quad \text{where } 0.1 \leq \theta \leq 1.$$

The constants in Corollary 3.3 are

$$K_1 = 2\theta^2, \quad K_2 = 2\theta, \quad \|\mathcal{D}(\mathbf{A})\|_\infty = 1$$

so that

$$\Delta_1 = 2\theta^2, \quad \Delta_2 = 2\theta, \quad d = \frac{2(n-1)\theta^2}{2\theta^2} = (n-1).$$

For all the test matrices, the constants  $\Delta_1$  and  $\Delta_2$  increase with increasing  $\theta$  as the off-diagonal elements become larger in magnitude relative to the diagonal elements. Therefore, we expect the Rademacher Monte Carlo estimators to lose accuracy with increasing  $\theta$ , as measured by the normwise relative error (NRE) in the computed diagonal  $\mathcal{D}(\hat{\mathbf{A}})$ ,

$$\text{NRE} \equiv \frac{\|\mathcal{D}(\mathbf{A}) - \mathcal{D}(\hat{\mathbf{A}})\|_\infty}{\|\mathcal{D}(\mathbf{A})\|_\infty},$$

in Figures 7.1–7.5.

**7.3. Experiment 1: Accuracy of Rademacher Monte Carlo estimator on test matrices.** Figures 7.1–7.3 show the NRE of the Rademacher Monte Carlo estimator applied to the test matrices in section 7.2 and the bounds from Corollary 3.3.

The big left panel displays the NRE versus the sampling amount  $N$ . This NRE represents the average of the NREs over 10 different independent runs. The small panels on the right show the bound  $\epsilon$  for the normwise estimators from Corollary 3.3 with failure probability  $\delta = 10^{-16}$ .

For Corollary 3.3, we solve for  $\epsilon$  from the simpler bound

$$N \geq \frac{\Delta_2}{3\epsilon^2} (2 + 6\Delta_3) \ln(8d/\delta), \quad \Delta_3 \equiv \frac{\Delta_1}{\Delta_2},$$

to obtain

(7.1)

$$\epsilon = \sqrt{\frac{\Delta_2}{3N} (2 + 6\Delta_3) \ln(8d/\delta)}, \quad \Delta_2 = \frac{\|\mathbf{A} - \mathcal{D}(\mathbf{A})\|_\infty}{\|\mathcal{D}(\mathbf{A})\|_\infty}, \quad \Delta_3 = \frac{K_1}{\|\mathbf{A} - \mathcal{D}(\mathbf{A})\|_\infty}.$$

The big left panels illustrate that, for a fixed sampling amount  $N$ , the NRE for Test Matrices 1 and 3 increases with  $\theta$ . This is because the off-diagonals become more dominant as  $\theta$  becomes larger.



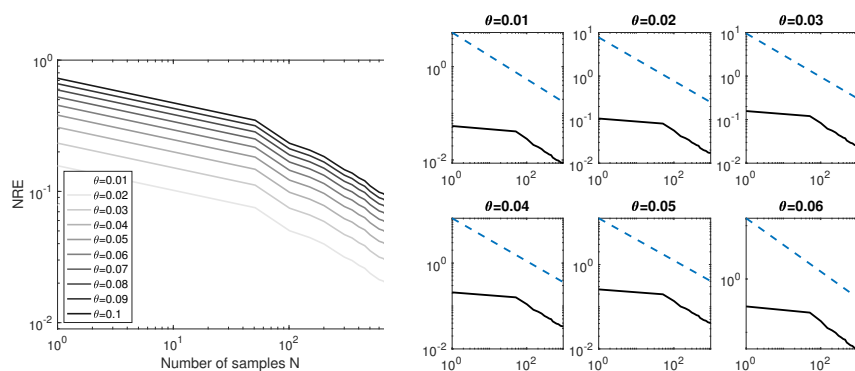


FIG. 7.1. Rademacher Monte Carlo estimator applied to Test Matrix 1. Big left panel: NRE for different values of  $\theta$  versus sampling amount  $N$ . Small panels on the right: NRE (solid black line) and bound (7.1) (blue dotted line) versus sampling amount  $N$  with failure probability  $\delta = 10^{-16}$ .

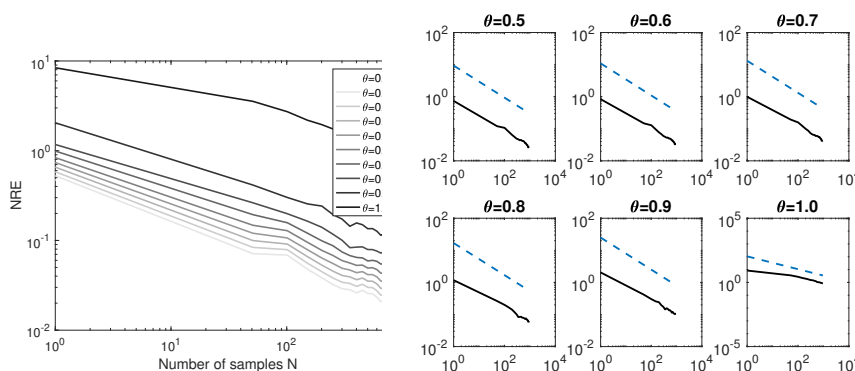


FIG. 7.2. Rademacher Monte Carlo estimator applied to Test Matrix 2. Big left panel: NRE for different values of  $\theta$  versus sampling amount  $N$ . Small panels on the right: NRE (solid black line) and bound (7.1) (blue dotted line) versus sampling amount  $N$  with failure probability  $\delta = 10^{-16}$ .

**7.4. Experiment 2: Different Monte Carlo estimators.** We compare the accuracy of the following Monte Carlo estimators on Test Matrix 1 with  $\theta = 0.01$ : Rademacher, Gaussian, sparse Rademacher with  $s = 3$ , and normalized Gaussian.

For each estimator, Figure 7.4 shows the mean of the NRE and variance over 100 runs with the shaded regions representing the 2.5% and 97.5% quantiles.

The normalized Gaussian estimator is about as accurate as the Rademacher estimator, while the sparse Rademacher with  $s = 3$  is about as accurate as the Gaussian estimator. The Gaussian and sparse Rademacher estimators are less accurate than the Rademacher and normalized Gaussian estimators. The shaded regions illustrate that, as expected, the sample variance of all estimators decreases with increasing sampling amount  $N$ .

**7.5. Experiment 3: Effect of sparsity in Rademacher vectors.** We apply the Rademacher Monte Carlo estimator to Test Matrix 1 with  $\theta = .01$  with four different sparsity levels:  $s = 1$  (standard Rademacher),  $s = 3$  [1],  $s = 10$ , and  $s = 50$ .

For each sampling amount  $N$ , Figure 7.5 shows the mean and the variance of the NRE over 100 runs. It suggests that sparse Rademacher estimators ( $s > 1$ ) may not be able to achieve a single digit of accuracy unless the sampling amount is so large as to exceed the matrix dimension.

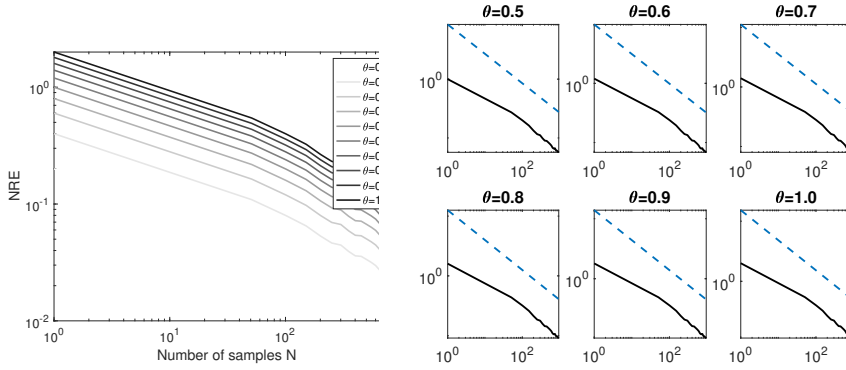


FIG. 7.3. Rademacher Monte Carlo estimator applied to Test Matrix 3. Big left panel: NRE for different values of  $\theta$  versus sampling amount  $N$ . Small panels on the right: NRE (solid black line) and bound (7.1) (blue dotted line) versus sampling amount  $N$  with failure probability  $\delta = 10^{-16}$ .

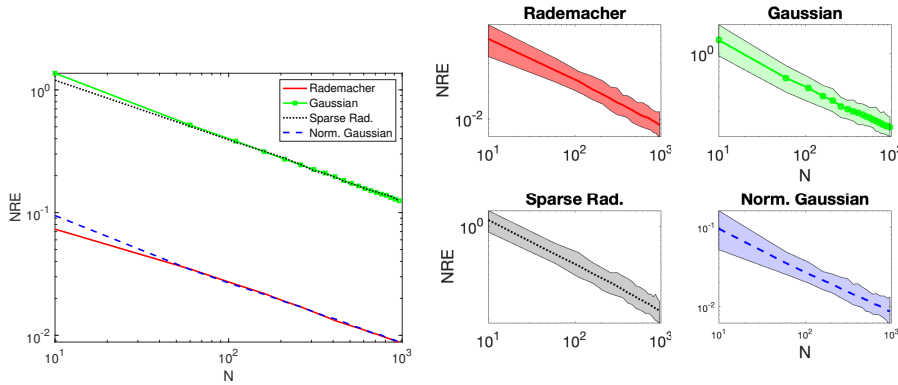


FIG. 7.4. Rademacher, Gaussian, sparse Rademacher with  $s = 3$ , and normalized Gaussian Monte Carlo estimators applied to Test Matrix 1 with  $\theta = 0.01$ . Big left panel: NRE mean versus sampling amount  $N$  for different estimators. Small right panels: NRE mean (styled lines) and 2.5% and 97.5% quantiles (shaded regions) versus sampling amount  $N$ .

**7.6. Example 4: Bounds for DGSM Monte Carlo estimator.** We apply the DGSM Monte Carlo estimator (6.1) to the diagonal matrix

$$(7.2) \quad \mathbf{S} \equiv \text{diag}(\mathbf{s}) \in \mathbb{R}^{n \times n}, \quad s_j \equiv \exp(-10j/n), \quad 1 \leq j \leq n,$$

from the quadratic function in section 6.1 for  $n = 100$  and illustrate the accuracy of Corollary 6.2.

The left panel of Figure 7.6 shows the normwise relative error

$$\text{NRE} \equiv \frac{\|\mathcal{D}(\mathbf{C}) - \mathcal{D}(\hat{\mathbf{C}})\|_2}{\|\mathcal{D}(\mathbf{C})\|_2},$$

which represents the the average of the NREs over 100 independent runs.

For Corollary 6.2, we fix the sample size  $N$  and solve for  $\epsilon$  from the simpler bound

$$N \geq \frac{S_2}{3\epsilon^2} (2 + 6S_3) \ln(8d/\delta), \quad S_3 \equiv \frac{S_1}{c_{\max} S_2},$$

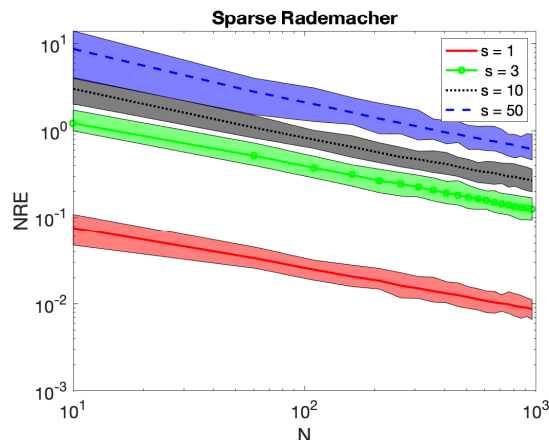


FIG. 7.5. Sparse Rademacher Monte Carlo estimators with sparsity levels  $s = 1, 3, 10, 50$  applied to Test Matrix 1 with  $\theta = 0.01$ . NRE (dotted lines) and 2.5% and 97.5% quantiles (shaded regions).

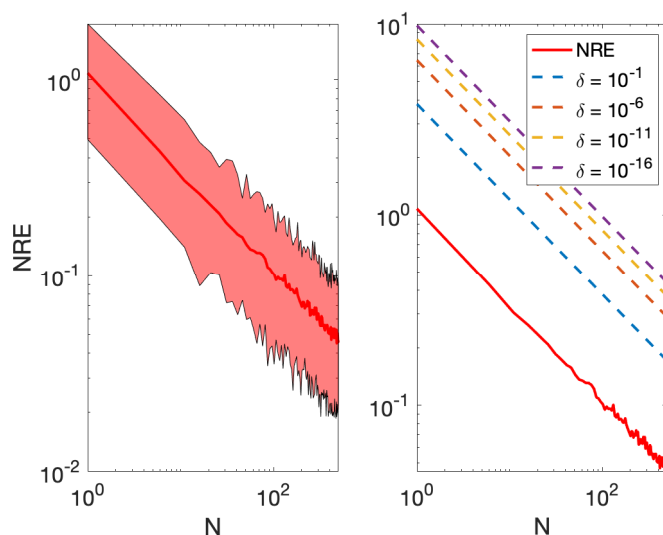


FIG. 7.6. DGSM Monte Carlo estimator (6.1) applied to  $\mathbf{S} \in \mathbb{R}^{100 \times 100}$  in (7.2). Left panel: NRE mean (solid line) and 2.5% and 97.5% quantiles (shaded regions) versus sampling amount  $N$ . Right panel: NRE and bounds (7.3) for different failure probabilities  $\delta$  versus sampling amount.

to obtain

$$(7.3) \quad \epsilon = \sqrt{\frac{S_2}{3N} (2 + 6S_3) \ln(8d/\delta)}.$$

The expressions for  $S_1, S_2, c_{\max}$ , and  $d$  for this example have been derived in section 6.1.

The right panel of Figure 7.6 illustrates that with less stringent failure probabilities  $\delta$ , the relative bounds (7.3) move closer to the NRE.

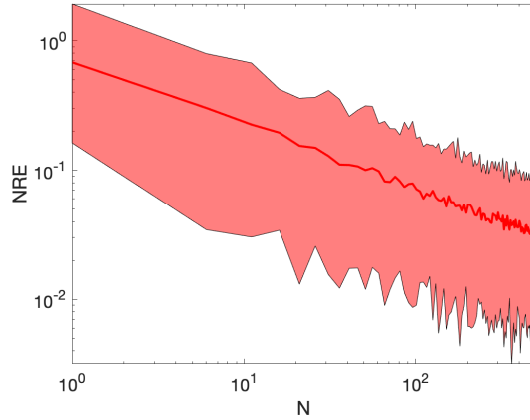


FIG. 7.7. DGSM Monte Carlo estimator (6.1) applied to the circuit model. NRE mean (solid line) and 2.5% and 97.5% quantiles (shaded regions) versus sampling amount  $N$ .

**7.7. Example 5: DGSM on the circuit model.** We apply the Monte Carlo DGSM estimator (6.1) to the *circuit model* from [7]. The quantity of interest is the midpoint voltage of a transformerless push-pull circuit, which depends on six parameters through a nonlinear closed-form algebraic expression.

As in [7], we normalize the parameter space to  $\mathcal{X} = [-1, 1]^6$  and scale the partial derivatives appropriately.<sup>1</sup>

Figure 7.7 shows the mean of the NRE and variances over 100 independent runs per sampling amount  $N$  with the shaded regions representing 2.5% and 97.5% quantiles. Since the exact expressions for the DGSMs are unavailable, we use as the exact value a tensor product Gauss–Legendre quadrature-based approximation with 15 points per dimension (i.e.,  $15^n$  total points).

**7.8. Example 6: Node centrality.** In network science, the centrality of a node quantifies its relative importance and can serve as a criterion for ranking the nodes. Many notions of centrality can be computed from the diagonal elements of a function of a matrix [5]. Here we consider the “resolvent subgraph centrality” measure for a graph represented by a symmetric adjacency matrix  $\mathbf{A}$  and compute the diagonal elements of

$$\mathbf{K}(\alpha) \equiv (\mathbf{I} - \alpha \mathbf{A})^{-1}, \quad \text{where} \quad 0 < \alpha < 1/\|\mathbf{A}\|_2$$

is a resolvent parameter. A Neumann series expansion of  $\mathbf{K}$  shows that  $\mathbf{K} \rightarrow \mathbf{I}$  as  $\alpha \rightarrow 0$ .

*Experimental setup.* We test the accuracy of the Monte Carlo estimators on the matrices in Table 7.1. They correspond to all the symmetric networks in the SNAP (Stanford Network Analysis Platform) collection of the SuiteSparse matrix collection [9]. The resolvent parameters are

$$(7.4) \quad \alpha_1 = \frac{0.9}{\|\mathbf{A}\|_2 + 1} \quad \text{and} \quad \alpha_2 = \frac{0.5}{\|\mathbf{A}\|_2 + 1}$$

with  $\|\mathbf{A}\|_2$  estimated via the MATLAB command `normest`.

<sup>1</sup>MATLAB codes are available at <https://bitbucket.org/paulcon/global-sensitivity-metrics-from-active-subspaces/src/master/>.

TABLE 7.1

Relative errors for resolvent-based centrality measures on 100 randomly selected nodes. The columns represent: adjacency matrix of the graph, its node and edge counts, average and maximal node degree, and relative errors (7.5) for the two different resolvent parameters in (7.4).

| Name        | Nodes   | Edges             | $d_{\text{avg}}$ | $d_{\text{max}}$ | RelErr ( $\alpha_1$ ) | RelErr ( $\alpha_2$ ) |
|-------------|---------|-------------------|------------------|------------------|-----------------------|-----------------------|
| as-735      | 7716    | $1.3 \times 10^4$ | 3.4              | 1459             | $2.5 \times 10^{-2}$  | $5.1 \times 10^{-3}$  |
| as-Skitter  | 1696415 | $1.1 \times 10^7$ | 13.1             | 35455            | $1.2 \times 10^{-2}$  | $1.5 \times 10^{-3}$  |
| ca-AstroPh  | 18772   | $1.9 \times 10^5$ | 21.1             | 504              | $2.7 \times 10^{-2}$  | $1.1 \times 10^{-3}$  |
| ca-CondMat  | 23133   | $9.3 \times 10^4$ | 8.1              | 280              | $5.3 \times 10^{-2}$  | $2.5 \times 10^{-2}$  |
| ca-GrQc     | 5242    | $1.5 \times 10^4$ | 5.5              | 81               | $4.2 \times 10^{-2}$  | $8.0 \times 10^{-3}$  |
| ca-HepPh    | 12008   | $1.2 \times 10^5$ | 19.7             | 491              | $5.5 \times 10^{-2}$  | $5.0 \times 10^{-3}$  |
| ca-HepTh    | 9877    | $2.6 \times 10^4$ | 5.3              | 65               | $1.9 \times 10^{-2}$  | $8.0 \times 10^{-3}$  |
| email-Enron | 36692   | $1.8 \times 10^5$ | 10.0             | 1383             | $1.1 \times 10^{-2}$  | $1.2 \times 10^{-2}$  |
| roadNet-CA  | 1971281 | $2.7 \times 10^6$ | 2.8              | 12               | $7.0 \times 10^{-2}$  | $4.6 \times 10^{-2}$  |
| roadNet-PA  | 1090920 | $1.5 \times 10^6$ | 2.83             | 9                | $1.1 \times 10^{-1}$  | $4.8 \times 10^{-2}$  |
| roadNet-TX  | 1393383 | $1.9 \times 10^6$ | 2.8              | 12               | $1.2 \times 10^{-1}$  | $3.9 \times 10^{-2}$  |

The Rademacher-based Monte Carlo estimators use  $N = 100$  samples. Matrix vectors products with  $\mathbf{K}$  are computed with the conjugate gradient algorithm with tolerance  $10^{-6}$  and maximal iteration count 128. Since estimating all diagonal elements is time consuming, we determine the relative error

$$(7.5) \quad \text{RelErr} \equiv \frac{\max_{i \in \mathcal{I}} |\mathbf{K}_{ii} - \widehat{\mathbf{K}}_{ii}|}{\max_{i \in \mathcal{I}} |\mathbf{K}_{ii}|}$$

only for 100 diagonal elements with randomly selected indices in the set  $\mathcal{I}$ .

Table 7.1. For the parameter  $\alpha_1$ , all errors are on the order of  $10^{-1}$  or  $10^{-2}$ , meaning one or two accurate digits regardless of the connectivity or size of the networks.

For the parameter  $\alpha_2$ , all errors are on the order of  $10^{-2}$  or  $10^{-3}$ , meaning two or three accurate digits. The consistently smaller errors compared to those for the larger parameter  $\alpha_1$  confirm that accuracy increases with diagonal dominance of  $\mathbf{K}$  and closeness to  $\mathbf{I}$ .

**8. Conclusion and future work.** Our normwise and elementwise probabilistic bounds for Rademacher- and Gaussian-based Monte Carlo diagonal estimators suggest that the accuracy increases with the diagonal dominance of the matrix and that sparse random vectors deliver less accuracy.

Avenues for future work include (i) diagonal estimators for matrix functions that are approximated by polynomials or rational functions and (ii) extension to Monte Carlo estimators for selected matrix elements, including off-diagonal elements.

**9. Acknowledgements.** The authors thank Alen Alexanderian for helpful discussions and the Associate Editor and two reviewers for their constructive recommendations.

## REFERENCES

- [1] D. ACHLIOPTAS, *Database-friendly random projections: Johnson-Lindenstrauss with binary coins*, J. Comput. System Sci., 66 (2003), pp. 671–687.
- [2] H. AVRON AND S. TOLEDO, *Randomized algorithms for estimating the trace of an implicit symmetric positive semi-definite matrix*, J. ACM, 58 (2011), pp. 1–34.

- [3] R. A. BASTON AND Y. NAKATSUKASA, *Stochastic Diagonal Estimation: Probabilistic Bounds and an Improved Algorithm*, preprint, arXiv:2201.10684, 2022.
- [4] C. BEKAS, E. KOKIOPOULOU, AND Y. SAAD, *An estimator for the diagonal of a matrix*, Appl. Numer. Math., 57 (2007), pp. 1214–1229.
- [5] M. BENZI AND P. BOITO, *Matrix functions in network analysis*, GAMM-Mitt., 43 (2020), e202000012, p. 36.
- [6] R. Y. CHEN, A. GITTENS, AND J. A. TROPP, *The masked sample covariance estimator: An analysis using matrix concentration inequalities*, Inf. Inference, 1 (2012), pp. 2–20.
- [7] P. G. CONSTANTINE AND P. DIAZ, *Global sensitivity metrics from active subspaces*, Reliab. Eng. Syst. Safe, 162 (2017), pp. 1–13.
- [8] A. CORTINOVIS AND D. KRESSNER, *On randomized trace estimates for indefinite matrices with an application to determinants*, Found. Comput. Math., (2021), pp. 1–29.
- [9] T. A. DAVIS AND Y. HU, *The University of Florida Sparse Matrix Collection*, ACM Trans. Math. Softw., 38 (2011), pp. 1–25.
- [10] G. H. GOLUB AND C. F. VAN LOAN, *Matrix Computations*, 4th ed., Johns Hopkins Studies in the Mathematical Sciences, Johns Hopkins University Press, Baltimore, 2013.
- [11] M. F. HUTCHINSON, *A stochastic estimator of the trace of the influence matrix for Laplacian smoothing splines*, Comm. Statist. Simulation Comput., 18 (1989), pp. 1059–1076.
- [12] B. J. KAPERICK, *Diagonal Estimation with Probing Methods*, Master's thesis, Virginia Polytechnic Institute and State University, 2019.
- [13] S. KUCHERENKO AND B. IOOSS, *Derivative-Based Global Sensitivity Measures*, Springer International Publishing, Cham, Switzerland, 2017, 1241–1263.
- [14] J. H. LAEUCHLI, *Methods for Estimating the Diagonal of Matrix Functions*, Ph.D. thesis, College of William and Mary, 2016.
- [15] R. R. LAM, O. ZAHM, Y. M. MARZOUK, AND K. E. WILLCOX, *Multifidelity dimension reduction via active subspaces*, SIAM J. Sci. Comput., 42 (2020), pp. A929–A956.
- [16] J. L. LARSEN AND L. M. MORRIS, *An Introduction to Mathematical Statistics and Its Applications*, 4th ed., Pearson Prentice Hall, Upper Saddle River, NJ, 2006.
- [17] P. LI, T. J. HASTIE, AND K. W. CHURCH, *Very sparse random projections*, in Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD '06, Association for Computing Machinery, New York, 2006, pp. 287–296.
- [18] M. MITZENMACHER AND E. UPFAL, *Probability and Computing*, Cambridge University Press, Cambridge, 2005.
- [19] F. ROOSTA-KHORASANI AND U. ASCHER, *Improved bounds on sample size for implicit matrix trace estimators*, Found. Comput. Math., 15 (2015), pp. 1187–1212.
- [20] A. P. SOMS, *An asymptotic expansion for the tail area of the t-distribution*, J. Amer. Statist. Assoc., 71 (1976), pp. 728–730.
- [21] J. M. TANG AND Y. SAAD, *Domain-decomposition-type methods for computing the diagonal of a matrix inverse*, SIAM J. Sci. Comput., 33 (2011), pp. 2823–2847.
- [22] J. M. TANG AND Y. SAAD, *A probing method for computing the diagonal of a matrix inverse*, Numer. Linear Algebra Appl., 19 (2012), pp. 485–501.
- [23] J. A. TROPP, *An introduction to matrix concentration inequalities*, Found. Trends Mach. Learn., 8 (2015), pp. 1–230.
- [24] S. UBARU AND Y. SAAD, *Applications of trace estimation techniques*, in High Performance Computing in Science and Engineering, T. Kozubek, M. Čermák, P. Tichý, R. Blaheta, J. Šístek, D. Lukáš, and J. Jaroš eds., Springer International Publishing, Cham, Switzerland, 2018, pp. 19–33.
- [25] R. VERSHYNIN, *High-Dimensional Probability: An Introduction with Applications in Data Science*, Vol. 47, Cambridge University Press, Cambridge, 2018.
- [26] A. J. WATHEN, *Preconditioning*, Acta Numer., 24 (2015), pp. 329–376.
- [27] Z. YAO, A. GHOLAMI, S. SHEN, M. MUSTAFA, K. KEUTZER, AND M. W. MAHONEY, *ADAHESIAN: An Adaptive Second Order Optimizer for Machine Learning*, preprint, arXiv:2006.00719, 2020.