

Applied Random Matrix Theory

Joel A. Tropp[†]

Abstract. Random matrices now play a role in many parts of computational mathematics. To advance these applications, it is desirable to have tools that are flexible, easy to use, and powerful. Over the last 25 years, researchers have developed a remarkable family of results, called matrix concentration inequalities, that meet the criteria. This paper offers an invitation to the field of matrix concentration and its multifarious applications.

Key words. graph algorithms, high-dimensional statistics, numerical analysis, numerical linear algebra, quantum information, random matrix theory

MSC codes. 15B52, 60B20, 65-02

DOI. 10.1137/25M1806922

1 Motivation. Random matrix theory emerged from applications in *statistics* (sample covariance matrices [58]) and in *nuclear physics* (Hamiltonians of heavy nuclei [57]). The subject established itself within the mathematical firmament through deep connections to *number theory* (zeros of the Riemann zeta function [35]), *operator algebras* (free product factors [55]), *combinatorics* (longest increasing subsequence [3]), and beyond. Over the last generation, random matrices have attracted new attention in computational mathematics, starting with work in *algorithms* [33] and in *quantum information* [1], and expanding in waves that have touched the farthest shores.

The classic literature emphasizes random matrices that enjoy a harmonious mathematical structure, such as the independent and identically distributed (i.i.d.) entries of a Wigner matrix [57]. By now, we understand these models comprehensively [2, 38]. In contrast, contemporary applications often lead to random matrices of more baroque construction, such as a random set of columns drawn from a fixed data matrix [48, sect. 5.2]. The standard methods for studying random matrices have relatively little to say about these strange examples.

To advance into this frontier, researchers have developed new tools for applied random matrix theory, collectively called *matrix concentration inequalities* [48]. These results offer several key benefits:

1. *Flexibility.* They apply to a large family of random matrices.
2. *Ease of use.* They reduce the analysis to a calculation of simple summary statistics.
3. *Power.* They accurately describe features of the random matrix model.

Matrix concentration has roots in Banach space geometry [45, 34, 40]; it also owes a heavy debt to research on quantum information theory [1]. The core results crystallized in 2010 through the efforts of Oliveira [37] and the author [46, 47]. My monograph [48] collected many applications and popularized the theory. Matrix concentration soon became textbook material [54, 56], and it has now informed thousands of research papers.

The purpose of this memoir is to introduce the fundamental matrix concentration inequalities for independent sums and martingales. I will illustrate these tools through eight contemporary applications in uncertainty quantification, numerical linear algebra, spectral graph theory, high-dimensional statistics, and quantum information science. Parts of this work are adapted from my monograph [48] and lecture notes [51].

1.1 Concentration of random matrices. A *random matrix* is a matrix whose entries are random variables, not necessarily independent from each other. The distinctive concern of random matrix theory is the action of the random matrix as a random linear map between linear spaces. In that vein, we can investigate its geometric properties (reflected in operator norms or its action on sets) and its spectral features (singular values and singular vectors, eigenvalues and eigenvectors in the square case).

[†]Department of Computing and Mathematical Sciences, Caltech, Pasadena, CA 91125 USA (jtropp@caltech.edu, <https://tropp.caltech.edu/>).

Matrix concentration studies the deviation of a random matrix $\mathbf{S} \in \mathbb{C}^{d_1 \times d_2}$ from its expectation $\mathbb{E}[\mathbf{S}]$, as measured in the spectral norm $\|\cdot\|$. For levels $t \geq 0$, we would like to control the tail probability

$$(1.1) \quad \mathbb{P}\{\|\mathbf{S} - \mathbb{E}[\mathbf{S}]\| \geq t\} \leq \text{???}.$$

The expectation of the random matrix is computed componentwise; we always assume that this expectation is defined and finite. The spectral norm is also called the ℓ_2 operator norm; it coincides with the largest singular value. If we control the deviation $\|\mathbf{S} - \mathbb{E}[\mathbf{S}]\|$, we also control other characteristics of the random matrix:

1. *Singular values.* The singular values of $\mathbf{S}$ and $\mathbb{E}[\mathbf{S}]$ are close.
2. *Singular vectors.* The singular vectors of $\mathbf{S}$ and $\mathbb{E}[\mathbf{S}]$ are close for isolated singular values.
3. *Linear functionals.* All linear functionals of $\mathbf{S}$ and $\mathbb{E}[\mathbf{S}]$ are comparable.

Many practical questions reduce to one of these considerations.

1.2 Random matrix models. We cannot hope to make progress on the question (1.1) without imposing some restrictions on the random matrix. The challenge is to identify models that capture a wide range of examples, yet offer enough scaffolding to support strong theoretical guarantees. This paper focuses on two models inspired by the most basic scalar stochastic processes: independent sums and martingales. Our attention to these templates will be rewarded by a host of applications. Section 6 outlines recent research and alternative models.

1.2.1 Independent sums. A random matrix $\mathbf{S} \in \mathbb{C}^{d_1 \times d_2}$ follows the *independent sum model* when it admits the decomposition

$$\mathbf{S} = \sum_{k=1}^n \mathbf{X}_k \quad \text{where the family } (\mathbf{X}_1, \dots, \mathbf{X}_n) \text{ is statistically independent.}$$

We use the term “statistical independence” to mark a contrast against linear independence. Let us emphasize that the entries within any particular matrix $\mathbf{X}_k$ may exhibit dependencies, but $\mathbf{X}_k$ cannot provide any information about events involving $(\mathbf{X}_j : j \neq k)$. Moreover, there may be many distinct decompositions of a random matrix as an independent sum. Inter alia, the independent sum model describes a random Monte Carlo approximation of a fixed matrix as an average of simple unbiased estimates [48, sect. 6.2].

We would like to capture information about the concentration of the independent sum $\mathbf{S}$ through summary statistics. In the scalar setting, we can achieve this goal with the Bernstein inequality [10, Thm. 2.10], which is arguably the most useful probability inequality for an independent sum of scalars. The matrix Bernstein inequality [47, Thm. 6.2] offers a perfect analogue in the matrix setting, where it may be the single most productive tool for studying random matrices. We present this result as Theorem 2.1. Section 3 showcases applications to active subspace methods, stochastic rounding, graph sparsification, and quantum state tomography.

1.2.2 Martingales. A *matrix martingale* is a sequence $(\mathbf{S}_0, \mathbf{S}_1, \mathbf{S}_2, \dots)$ of random matrices with common dimension $d_1 \times d_2$ that satisfies the “status quo” property:

$$\mathbb{E}[\mathbf{S}_{k+1} | \mathbf{S}_0, \dots, \mathbf{S}_k] = \mathbf{S}_k \quad \text{for each } k = 0, 1, 2, \dots$$

Given what we know so far, our best estimate for the next matrix $\mathbf{S}_{k+1}$ is the current matrix $\mathbf{S}_k$. As a simple example, the partial sums of an independent family of centered random matrices compose a matrix martingale. More generally, matrix martingales can model complicated sequences of random matrices that are revealed one step at a time, such as the iterates of a randomized algorithm that performs a linear algebra computation.

One of the most powerful tools for studying scalar martingale sequences is the Freedman inequality [18], which is the martingale analogue of the Bernstein inequality. The matrix Freedman inequality [37, 46] generalizes the scalar result to matrix martingales. It allows us to treat sequences of random matrices that evolve adaptively, and it provides uniform control on the whole trajectory. We present this result as Theorem 4.4. Section 5 highlights applications to online covariance estimation, Cholesky decomposition with stochastic rounding, fast graph Laplacian solvers, and randomized approximation of quantum Hamiltonians.

1.3 Notation. The Pascal notation $:=$ or $=:$ generates a definition. The symbols $\vee$ and $\wedge$ denote the infix maximum and minimum. We work over a scalar field $\mathbb{F} \in \{\mathbb{R}, \mathbb{C}\}$, which is real or complex. The star $*$ delivers the (conjugate) transpose of a vector or matrix. The linear space $\mathbb{F}^d$ is equipped with the standard inner product $\langle \mathbf{b}, \mathbf{a} \rangle := \mathbf{b}^* \mathbf{a}$ and the associated ℓ_2 -norm $\|\cdot\|$.

In the space $\mathbb{F}^d$, the standard basis elements are $\mathbf{e}_1, \dots, \mathbf{e}_d$, and $\mathbf{1}$ is the vector of ones. The standard basis elements in a matrix space, such as $\mathbb{F}^{d_1 \times d_2}$, are $\mathbf{E}_{ij}$. The identity matrix is $\mathbf{I}$. Dimensions are determined by context. We write a_{ij} for the entries of a matrix $\mathbf{A}$, and we use the colon operator $\mathbf{A}(i, :)$ or $\mathbf{A}(:, j)$ to extract the i th row or j th column. For matrices, $\|\cdot\|$ is the spectral norm (also known as the Schatten ∞ -norm), $\|\cdot\|_F$ denotes the Frobenius norm (or Schatten 2-norm), and $\|\cdot\|_1$ refers to the trace norm (or Schatten 1-norm).

The linear space $\mathbb{M}_d(\mathbb{F})$ consists of $d \times d$ matrices with entries in $\mathbb{F}$, on which the operator Tr computes the trace. The real-linear subspace $\mathbb{H}_d(\mathbb{F})$ consists of the $d \times d$ Hermitian (i.e., conjugate-symmetric) matrices. The cone $\mathbb{H}_d^+(\mathbb{F})$ contains the $d \times d$ positive-semidefinite (psd) matrices, and $\mathbb{H}_d^{++}(\mathbb{F})$ is the cone of positive-definite matrices. For Hermitian matrices, the semidefinite partial order $\mathbf{A} \preceq \mathbf{H}$ signifies that $\mathbf{H} - \mathbf{A}$ is psd, while the maps $\lambda_{\max}$ and $\lambda_{\min}$ return the (algebraic) maximum and minimum eigenvalues.

The operator $\mathbb{E}$ computes the expectation of a random variable, and $\mathbb{P}$ reports the probability of an event. The symbol $\sim$ means “has the distribution.” For a real random variable X , the function $\sup X$ specifies its least upper bound on the probability space. Within a master probability space, the function $\sigma(\cdot)$ returns the sigma-algebra generated by its arguments; the empty argument returns the trivial sigma-algebra.

The big- $\mathcal{O}$ notation is interpreted per the conventions of computer science.

2 The matrix Bernstein inequality. Among matrix concentration inequalities, the single most important result is the matrix extension of the scalar Bernstein inequality [10, Thm. 2.10]. The matrix Bernstein inequality was established independently by Oliveira [37] in late 2009 and the author [47] in early 2010. We state the version of this result from [48, Thm. 6.1.1].

THEOREM 2.1 (Matrix Bernstein). *Consider a statistically independent family $(\mathbf{X}_1, \dots, \mathbf{X}_n)$ of $d_1 \times d_2$ random matrices, real or complex, that are centered and uniformly bounded: $\mathbb{E}[\mathbf{X}_k] = \mathbf{0}$ and $\|\mathbf{X}_k\| \leq B$ for each index k . Form the sum, and calculate its matrix variance statistic:*

$$(2.1) \quad \mathbf{S} := \sum_{k=1}^n \mathbf{X}_k \quad \text{and} \quad v(\mathbf{S}) := \|\mathbb{E}[\mathbf{S}\mathbf{S}^*]\| \vee \|\mathbb{E}[\mathbf{S}^*\mathbf{S}]\|.$$

Then the spectral norm of the sum satisfies the probability inequalities

$$(2.2) \quad \mathbb{E} \|\mathbf{S}\| \leq \sqrt{2v(\mathbf{S}) \log(d_1 + d_2)} + \frac{1}{3} B \log(d_1 + d_2),$$

$$(2.3) \quad \mathbb{P} \{\|\mathbf{S}\| \geq t\} \leq (d_1 + d_2) \exp \left(\frac{-t^2/2}{v(\mathbf{S}) + Bt/3} \right) \quad \text{for } t \geq 0.$$

The proof of Theorem 2.1 parallels the proof of the scalar Bernstein inequality, but it requires sophisticated tools from matrix analysis. The argument will occupy the rest of this section. First, we elaborate on the meaning of the result.

The significant outcome of Theorem 2.1 is the expectation bound (2.2). In contrast, the tail bound (2.3) is usually somewhat loose; it can be improved by combining the expectation bound (2.2) with scalar concentration inequalities for the norm of an independent sum, such as [10, Prob. 13.33].

We remark that Theorem 2.1 reproduces the scalar Bernstein inequality when it is applied to random scalars (i.e., 1×1 random matrices), so the numerical constants are sharp. The dimensional factor $(d_1 + d_2)$ is a new feature in the matrix setting, and it is a necessary component of the bound; see subsection 2.6.

Like the scalar variance, the matrix variance $v(\mathbf{S})$ reflects the magnitude of the expected “square” of the centered random matrix, where we keep in mind that the non-Hermitian matrix $\mathbf{S}$ has two distinct squares $\mathbf{S}\mathbf{S}^*$ and $\mathbf{S}^*\mathbf{S}$. Both terms in the matrix variance are required, although they coincide when the sum is Hermitian. It is often convenient to express the matrix variance directly in terms of the summands:

$$(2.4) \quad v(\mathbf{S}) = \left\| \sum_{k=1}^n \mathbb{E}[\mathbf{X}_k \mathbf{X}_k^*] \right\| \vee \left\| \sum_{k=1}^n \mathbb{E}[\mathbf{X}_k^* \mathbf{X}_k] \right\|.$$

The last identity exploits independence and centering. For simplicity, Theorem 2.1 assumes that the matrix $\mathbf{S}$ is centered; otherwise, note that

$$\|\mathbf{S}\| \leq \|\mathbb{E}[\mathbf{S}]\| + \|\mathbf{S} - \mathbb{E}[\mathbf{S}]\|.$$

The theorem now applies to the second term.

2.1 Reduction to the Hermitian case. Let us commence with the proof of Theorem 2.1. It suffices to consider the complex case. Introduce the *Hermitian dilation*:

$$(2.5) \quad \mathcal{H}: \mathbb{C}^{d_1 \times d_2} \rightarrow \mathbb{H}_{d_1+d_2}(\mathbb{C}) \quad \text{where} \quad \mathcal{H}(\mathbf{A}) := \begin{bmatrix} \mathbf{0} & \mathbf{A} \\ \mathbf{A}^* & \mathbf{0} \end{bmatrix}.$$

The map $\mathcal{H}$ is real-linear, and it preserves spectral data in the sense that $\lambda_{\max}(\mathcal{H}(\mathbf{A})) = -\lambda_{\min}(\mathcal{H}(\mathbf{A})) = \|\mathbf{A}\|$. As a consequence, to prove Theorem 2.1, we can pass to the random Hermitian matrix

$$\mathcal{H}(\mathbf{S}) = \sum_{k=1}^n \mathcal{H}(\mathbf{X}_k) \quad \text{where} \quad \mathbb{E}[\mathcal{H}(\mathbf{X}_k)] = \mathbf{0} \quad \text{and} \quad \|\mathcal{H}(\mathbf{X}_k)\| \leq B.$$

Through this device, we can simply instantiate the assumption that the summands are Hermitian matrices.

2.2 The matrix Laplace transform method. To extract information about the distribution of a bounded random *Hermitian* matrix $\mathbf{S} \in \mathbb{H}_d(\mathbb{C})$, define the logarithm of the *matrix moment generating function* (briefly, the *matrix log-mgf*):

$$\Xi_{\mathbf{S}}(\theta) := \log \mathbb{E} e^{\theta \mathbf{S}} \in \mathbb{H}_d(\mathbb{C}) \quad \text{for a parameter } \theta \in \mathbb{R}.$$

We apply a scalar function to an Hermitian matrix by applying the function to each eigenvalue without modifying the associated eigenspace. Ahlswede and Winter [1] formulated the ideas in this section, while Oliveira [37] and the author [47, 48] crystallized the arguments.

Even in the matrix setting, we can pursue the same strategy that leads to exponential tail bounds in the scalar setting [10, Chap. 2]. When the parameter $\theta > 0$,

$$(2.6) \quad \begin{aligned} \mathbb{P}\{\lambda_{\max}(\mathbf{S}) \geq t\} &= \mathbb{P}\left\{e^{\theta \lambda_{\max}(\mathbf{S})} \geq e^{\theta t}\right\} \leq e^{-\theta t} \cdot \mathbb{E} e^{\theta \lambda_{\max}(\mathbf{S})} \\ &= e^{-\theta t} \cdot \mathbb{E} \lambda_{\max}(e^{\theta \mathbf{S}}) \leq e^{-\theta t} \cdot \mathbb{E} \operatorname{Tr} e^{\theta \mathbf{S}} = e^{-\theta t} \cdot \operatorname{Tr} \exp(\Xi_{\mathbf{S}}(\theta)). \end{aligned}$$

The first inequality is Markov's. Afterward, invoke the spectral mapping theorem to draw the eigenvalue map through the exponential, and note that the trace of a psd matrix dominates its maximum eigenvalue. The emergence of the trace exponential of the matrix log-mgf unlocks powerful tools for trace functions.

A similar calculation furnishes a bound for the expectation of the maximum eigenvalue. For each $\theta > 0$,

$$(2.7) \quad \mathbb{E} \lambda_{\max}(\mathbf{S}) = \theta^{-1} \log \exp(\mathbb{E} \lambda_{\max}(\theta \mathbf{S})) \leq \theta^{-1} \log \mathbb{E} e^{\lambda_{\max}(\theta \mathbf{S})} \leq \theta^{-1} \log \mathbb{E} \operatorname{Tr} e^{\theta \mathbf{S}} = \theta^{-1} \log \operatorname{Tr} \exp(\Xi_{\mathbf{S}}(\theta)).$$

The first inequality is Jensen's. Once again, the matrix log-mgf controls the maximum eigenvalue.

2.3 Subadditivity of the matrix log-mgf. In the scalar setting, the log-mgf of a sum of independent real random variables agrees with the sum of the log-mgfs:

$$\log \mathbb{E} e^{\theta \sum_{k=1}^n X_k} = \sum_{k=1}^n \log \mathbb{E} e^{\theta X_k} \quad \text{where } (X_1, \dots, X_n) \in \mathbb{R}^n \text{ is an independent family.}$$

The latter formula depends on the fact $e^{a+b} = e^a e^b$ for $a, b \in \mathbb{R}$, a property that catastrophically fails for matrices.

Remarkably, there is a substitute. I established a *subadditivity rule* for the matrix log-mgf [47, Lem. 3.4]:

$$(2.8) \quad \operatorname{Tr} \exp(\Xi_{\mathbf{S}}(\theta)) \leq \operatorname{Tr} \exp\left(\sum_{k=1}^n \Xi_{\mathbf{X}_k}(\theta)\right) \quad \text{where } \mathbf{S} := \sum_{k=1}^n \mathbf{X}_k.$$

In this expression, $(\mathbf{X}_1, \dots, \mathbf{X}_n)$ is a statistically independent family of bounded, Hermitian random matrices of the same dimension. In the proof of (2.8), the key ingredient is a deep concavity theorem of Lieb [32, Thm. 6], which is one of the crown jewels of matrix analysis. See [48, Chap. 8] for a detailed proof of Lieb's result.

2.4 The matrix Bernstein log-mgf bound. The subadditivity rule (2.8) reduces the bound on the matrix log-mgf of a sum to individual bounds on the matrix log-mgfs of the summands. Classic scalar concentration inequalities [10, Chap. 2] provide inspiration about fruitful strategies for controlling the matrix log-mgfs [47, 48].

The matrix Bernstein inequality depends on the same type of log-mgf bound as the scalar Bernstein inequality [10, Thm. 2.10]. Indeed, the result [48, Lem. 6.1] states that the matrix log-mgf of a bounded, centered, random Hermitian matrix $\mathbf{X} \in \mathbb{H}_d(\mathbb{C})$ satisfies the relation

$$(2.9) \quad \Xi_{\mathbf{X}}(\theta) \preceq \frac{\theta^2/2}{1 - B|\theta|/3} \cdot \mathbb{E}[\mathbf{X}^2] \quad \text{where } \mathbb{E}[\mathbf{X}] = \mathbf{0} \text{ and } \|\mathbf{X}\| \leq B \text{ and } |\theta| < 3/B.$$

The semidefinite partial order $\mathbf{A} \preceq \mathbf{H}$ on Hermitian matrices means that $\mathbf{H} - \mathbf{A}$ is psd. The proof of (2.9) tracks the scalar argument. Expand the exponential function in the matrix log-mgf as a Taylor series, use the centering to remove the term with degree one, and apply the norm bound to control the terms with degree two and higher. We also employ the fact that the matrix logarithm respects the semidefinite partial order [48, Prop. 8.4.4].

2.5 Assembly line. We are prepared to establish Theorem 2.1. Consider the sum $\mathbf{S} := \sum_{k=1}^n \mathbf{X}_k$ of independent, random Hermitian matrices with dimension d that satisfy $\mathbb{E}[\mathbf{X}_k] = \mathbf{0}$ and $\|\mathbf{X}_k\| \leq B$. Sequence the displays (2.6), (2.8), and (2.9) to arrive at the tail inequality

$$(2.10) \quad \begin{aligned} \mathbb{P}\{\lambda_{\max}(\mathbf{S}) \geq t\} &\leq e^{-\theta t} \cdot \text{Tr} \exp \left(\frac{\theta^2/2}{1 - B|\theta|/3} \sum_{k=1}^n \mathbb{E}[\mathbf{X}_k^2] \right) \\ &= e^{-\theta t} \cdot \text{Tr} \exp \left(\frac{\theta^2/2}{1 - B|\theta|/3} \mathbb{E}[\mathbf{S}^2] \right) \leq d \cdot e^{-\theta t} \cdot \exp \left(\frac{v(\mathbf{S})\theta^2/2}{1 - B|\theta|/3} \right). \end{aligned}$$

The parameter is restricted to the interval $|\theta| < B/3$. The first inequality depends on the fact that the trace exponential respects the semidefinite partial order [48, Ex. 8.1]. The identity exploits the centering and independence of the summands. Afterward, bound the trace exponential by the dimension d times the maximum eigenvalue, and use spectral mapping to recognize the matrix variance $v(\mathbf{S}) = \lambda_{\max}(\mathbb{E}[\mathbf{S}^2])$.

Finally, in (2.10), set the parameter $\theta = t/(v(\mathbf{S}) + Bt/3)$ to reach the finished tail inequality

$$(2.11) \quad \mathbb{P}\{\lambda_{\max}(\mathbf{S}) \geq t\} \leq d \cdot \exp \left(\frac{-t^2/2}{v(\mathbf{S}) + Bt/3} \right).$$

To establish the analogous tail bound (2.3) for a general non-Hermitian sum, apply the Hermitian dilation (2.5) and invoke (2.11). The bound (2.2) on the *expected* norm of the sum follows a similar argument starting from (2.7).

2.6 Optimality. As noted, the expectation bound (2.2) is the most significant outcome from Theorem 2.1. We can strengthen this bound, but not by very much. Introduce the *tail content* statistic:

$$B_2 := (\mathbb{E} \max_k \|\mathbf{X}_k\|^2)^{1/2}.$$

The tail content satisfies $B_2 \leq B$; the slackness of this inequality depends on the particular choice of summands. Using another style of argument based on symmetrization [49], we can obtain a two-sided bound of the form

$$(2.12) \quad \sqrt{v(\mathbf{S})} + B_2 \lesssim \mathbb{E} \|\mathbf{S}\| \lesssim \sqrt{v(\mathbf{S}) \log(d_1 + d_2)} + B_2 \log(d_1 + d_2).$$

The paper [49] provides examples to show that each term in (2.12) is necessary, so we cannot improve the bounds without a substantial amount of extra information (about the covariance structure of the random matrix). Section 6 mentions some situations where improvements are possible.

3 Independent sum model: Applications. In this section, we present some contemporary applications of the matrix Bernstein inequality in several areas of computational mathematics. As a caveat, these stylized applications may not reflect the intricacies of each problem. It is sometimes possible to find a simpler argument by invoking a more suitable matrix concentration inequality, and we can occasionally obtain sharper analyses using more complicated tools; see section 6.

3.1 Active subspace methods. In engineering design, researchers build computer models of complex physical systems that are governed by several input parameters. Key tasks include optimization of parameters, analysis of sensitivity to changes in parameters, and quantification of uncertainty about the system output. For example, an aeronautics engineer designs an airfoil by modifying its geometry to reach a target lift and drag coefficient. It is challenging to explore the parameter space of a complicated, nonlinear model, so it is common to seek reduced models. One basic methodology is to restrict our attention to the most salient directions in the input space, which we can identify through a random sampling procedure. The analysis of this approach depends on matrix concentration inequalities. This vignette is adapted from Constantine's book [16, Chap. 3]. The mathematics are similar with the classic problem of covariance estimation; see subsection 5.1.

Let $f: \mathbb{R}^d \rightarrow \mathbb{R}$ be an L -Lipschitz, differentiable function that models a complicated system, and assume that we can evaluate its gradient ∇f at each point in the parameter space—but the computational cost is significant. As part of the model, introduce a random vector $\mathbf{z} \in \mathbb{R}^d$ that describes a distribution over the parameter space; this distribution is often interpreted as a Bayesian prior. To capture the variability of the function, introduce the psd *sensitivity matrix*

$$(3.1) \quad \Sigma := \mathbb{E}[(\nabla f(\mathbf{z}))(\nabla f(\mathbf{z}))^*] \in \mathbb{H}_d^+(\mathbb{R}).$$

The quadratic form in Σ reflects the sensitivity of the function to directional perturbations in the input space:

$$\mathbf{a}^* \Sigma \mathbf{a} = \mathbb{E}[|\langle \nabla f(\mathbf{z}), \mathbf{a} \rangle|^2] \quad \text{for each } \mathbf{a} \in \mathbb{R}^d.$$

The subspace spanned by the leading eigenvectors of Σ is called an *active subspace* of the model, because it captures the most salient directions in the input space.

We cannot evaluate the sensitivity matrix Σ explicitly, but we can construct a Monte Carlo approximation:

$$(3.2) \quad \hat{\Sigma}_n := \frac{1}{n} \sum_{k=1}^n (\nabla f(\mathbf{z}_k))(\nabla f(\mathbf{z}_k))^* \quad \text{where } \mathbf{z}_k \sim \mathbf{z} \text{ i.i.d.}$$

How many gradient samples suffice to approximate the sensitivity matrix Σ ? Can we determine the directions in which the model varies substantially?

THEOREM 3.1 (Sampling active subspaces [16, Thm. 3.7]). *Let $f: \mathbb{R}^d \rightarrow \mathbb{R}$ be L -Lipschitz and differentiable. As in (3.1) and (3.2), consider the sensitivity matrix $\Sigma \in \mathbb{H}_d(\mathbb{R})$ and the empirical approximation $\hat{\Sigma}_n \in \mathbb{H}_d(\mathbb{R})$ determined by n gradient samples. The expectation of the relative approximation error satisfies*

$$\mathbb{E} \left[\frac{\|\hat{\Sigma}_n - \Sigma\|}{\|\Sigma\|} \right] \leq \sqrt{2\beta_n} + \frac{1}{3}\beta_n \quad \text{where} \quad \beta_n := \frac{L^2 \log(2d)}{\|\Sigma\|} \cdot \frac{1}{n}.$$

For $\varepsilon \in (0, 1)$, the sample complexity $n \geq 4\varepsilon^{-2} L^2 \log(2d) / \|\Sigma\|$ results in expected relative error at most ε .

The ratio $r := L^2 / \|\Sigma\|$ reflects the number of energetic dimensions in the model. Given $n = \mathcal{O}(r \log d)$ gradient samples, the estimator $\hat{\Sigma}_n$ for the sensitivity matrix Σ allows us to find this active subspace. To reduce the relative error to a level ε , the estimator requires an additional factor of ε^{-2} samples, so the empirical approximation does not attain high accuracy. This is a familiar weakness of Monte Carlo methods.

Proof sketch. The error in the Monte Carlo approximation is an instance of the independent sum model. Abbreviate $\mathbf{y}_k := \nabla f(\mathbf{z}_k)$ for each index k . Since the function f is L -Lipschitz, the random vectors satisfy the uniform bound $\|\mathbf{y}_k\| \leq L$. Introduce the matrix approximation error

$$\mathbf{S} := \hat{\Sigma}_n - \Sigma = \sum_{k=1}^n \frac{1}{n} (\mathbf{y}_k \mathbf{y}_k^* - \mathbb{E}[\mathbf{y}_k \mathbf{y}_k^*]) =: \sum_{k=1}^n \mathbf{X}_k.$$

The summands $(\mathbf{X}_k)$ compose an independent family of bounded, centered, random Hermitian matrices with common dimension d . The statistics appearing in the matrix Bernstein inequality (Theorem 2.1) satisfy

$$B = \sup \|\mathbf{X}_k\| \leq \frac{L^2}{n} \quad \text{and} \quad v(\mathbf{S}) = \|\mathbb{E}[\mathbf{S}^2]\| = n \cdot \|\mathbb{E}[\mathbf{X}_1^2]\| \leq \frac{1}{n} \|\mathbb{E}[(\mathbf{y}_1 \mathbf{y}_1^*)^2]\| \leq \frac{L^2}{n} \cdot \|\Sigma\|.$$

The stated result follows from (2.2). □

3.2 Stochastic rounding. Conventional computer architectures offer floating-point numbers with 16 decimal digits of precision or more. Driven by contemporary applications, computer engineers have started to develop new architectures with far lower precision—sometimes as few as 4 or 8 bits (i.e., 1–2 digits). At this extreme, large numerical errors can accumulate as we round successive calculations to the nearest machine-representable number. One way to mitigate these errors is to design computer systems that perform *stochastic rounding*. To analyze linear algebra computations that employ stochastic rounding, we can exploit matrix concentration tools. This section offers an idealized treatment of the simplest problem; see [17] for more texture.

A *floating-point number system* is a finite set of machine-representable numbers $F \subset \mathbb{R}$. Ignoring underflow and overflow, a (deterministic) rounding rule approximates each real number $a \in \mathbb{R}$ by a floating-point number $\text{float}(a) \in F$ that admits the relative-error bound

$$(3.3) \quad |a - \text{float}(a)| \leq u \cdot |a| \quad \text{where } u \in \mathbb{R} \text{ is the unit-roundoff parameter.}$$

In IEEE double-precision arithmetic, the unit-roundoff parameter $u = 2^{-53} \approx 10^{-16}$, and the rounding rule must satisfy additional requirements to meet the standard [17, Tables 1 and 2]. As an alternative, a *stochastic rounding rule* maps each real number $a \in \mathbb{R}$ to a *random* floating-point number $\text{sr}(a) \in F$ that satisfies

$$(3.4) \quad \mathbb{E}[\text{sr}(a)] = a \quad \text{and} \quad |a - \text{sr}(a)| \leq u \cdot |a|.$$

In other words, stochastic rounding is unbiased, and it commits a relative error on the order of the unit roundoff.

Suppose that we round each entry of a real-valued matrix stochastically to the nearest floating-point number. How much damage will we do? How does this bound compare with deterministic rounding?

THEOREM 3.2 (Stochastic rounding). *Consider a matrix $\mathbf{A} \in \mathbb{R}^{d_1 \times d_2}$. For $p \in \{1, 2\}$, define the norms*

$$\|\mathbf{A}\|_{\max} := \max_{ij} |a_{ij}| \quad \text{and} \quad \|\mathbf{A}\|_{\text{rc}, p} := (\max_i \|\mathbf{A}(i, :)\|_p) \vee (\max_j \|\mathbf{A}(:, j)\|_p).$$

Form a random matrix $\text{sr}(\mathbf{A}) \in F^{d_1 \times d_2}$ by independently rounding each entry of $\mathbf{A}$, honoring the rule (3.4). Then

$$\mathbb{E} \|\mathbf{A} - \text{sr}(\mathbf{A})\| \leq u \cdot \left[\sqrt{2\|\mathbf{A}\|_{\text{rc}, 2}^2 \log(d_1 + d_2)} + \frac{1}{3} \|\mathbf{A}\|_{\max} \log(d_1 + d_2) \right].$$

In contrast, the deterministic rule (3.3) only guarantees that $\|\mathbf{A} - \text{float}(\mathbf{A})\| \leq u \cdot \|\mathbf{A}\|_{\text{rc}, 1}$.

To interpret the result, it is helpful to consider a square matrix with dimension d whose entries all have similar magnitudes. In this case, the stochastic rounding bound is governed by the norm $\|\mathbf{A}\|_{\text{rc}, 2}$, which is proportional to $\sqrt{d}$. Meanwhile, the deterministic rounding bound depends on $\|\mathbf{A}\|_{\text{rc}, 1}$, which is proportional to the matrix dimension d . This contrast indicates that stochastic rounding confers a significant benefit.

Proof sketch. The matrix of rounding errors is an instance of the independent sum model:

$$\mathbf{S} := \mathbf{A} - \text{sr}(\mathbf{A}) = \sum_{ij} \varepsilon_{ij} a_{ij} \mathbf{E}_{ij} \quad \text{where } \mathbb{E}[\varepsilon_{ij}] = 0 \text{ and } |\varepsilon_{ij}| \leq u.$$

The summands $\mathbf{X}_{ij} := \varepsilon_{ij} a_{ij} \mathbf{E}_{ij}$ are bounded, centered independent random matrices with dimension $d_1 \times d_2$. A short calculation confirms that the parameters in the matrix Bernstein inequality (Theorem 2.1) satisfy

$$B = \sup \|\mathbf{X}_{ij}\| \leq u \cdot \|\mathbf{A}\|_{\max} \quad \text{and} \quad v(\mathbf{S}) = \|\mathbb{E}[\mathbf{S}\mathbf{S}^*]\| \vee \|\mathbb{E}[\mathbf{S}^*\mathbf{S}]\| \leq u^2 \cdot \|\mathbf{A}\|_{\text{rc}, 2}^2.$$

The bound on the expectation of the error follows directly from (2.2). The error bound for deterministic rounding holds because $\|\mathbf{M}\| \leq \|\mathbf{M}\|_{\text{rc}, 1}$ for any matrix $\mathbf{M}$; see [22, Prob. 5.6P21]. $\square$

In fact, the error estimate for stochastic rounding in Theorem 3.2 admits a modest refinement:

$$\mathbb{E} \|\mathbf{A} - \text{sr}(\mathbf{A})\| \leq 4u \cdot \left[\|\mathbf{A}\|_{\text{rc}, 2} + \|\mathbf{A}\|_{\max} \sqrt{1 + \log(d_1 \wedge d_2)} \right].$$

This bound follows from a specialized result [12, Thm. 1.4] for random matrices with independent entries, plus a symmetrization argument [54, Lem. 6.4.2]. For most practical use cases, the improvement is insubstantial.

3.3 Graph sparsification. A combinatorial graph encodes the pairwise relationships among a family of objects. We may ask whether it is possible to find a simpler graph (with fewer edges) that preserves structural properties of the original graph, including the weights of vertex cuts and the mixing time of a random walk. Spielman and Srivastava [41] showed how to achieve this goal by randomly sampling edges from the graph. We can analyze the procedure with matrix concentration.

Consider a connected, weighted, undirected graph (V, w) . The graph comprises a set $V := \{1, \dots, n\}$ of vertices and a symmetric function $w : V \times V \rightarrow \mathbb{R}_+$ that assigns a nonnegative weight to each pair of vertices. We require that $w_{ii} = 0$ and $w_{ij} = w_{ji}$ for all $i, j \in V$. The number of edges $m := \#\{i < j : w_{ij} > 0\}$.

We can also represent the graph by means of the *graph Laplacian* $L \in \mathbb{H}_n^+(\mathbb{R})$, which is the psd matrix

$$(3.5) \quad L := \sum_{i < j} w_{ij} \cdot (\mathbf{e}_i - \mathbf{e}_j)(\mathbf{e}_i - \mathbf{e}_j)^* =: \sum_{i < j} w_{ij} \cdot \Delta_{ij}.$$

The sum involves m nonzero terms. Since the graph is connected, $\text{null}(L) = \text{span}\{\mathbf{1}\}$. Without further notice, we restrict L to the $(n - 1)$ -dimensional subspace where it is nonsingular. Define the *effective resistance* of each vertex pair in the graph:

$$\varrho_{ij} := \text{Tr}[L^{-1/2} \Delta_{ij} L^{-1/2}] \quad \text{for each vertex pair } (i, j).$$

If the graph models an electrical network whose wires have conductances w_{ij} , then the effective resistance ϱ_{ij} is the voltage required to push one unit of current from vertex i to vertex j . Observe that $\sum_{i < j} w_{ij} \varrho_{ij} = n - 1$. We can approximate all m of the nonzero effective resistances in $\mathcal{O}(m \log n)$ arithmetic operations [41, Thm. 2].

Our goal is to produce a *sparse* graph whose Laplacian matrix $\hat{L} \in \mathbb{H}_n^+(\mathbb{R})$ is comparable with the original Laplacian L in the semidefinite partial order:

$$(3.6) \quad (1 - \varepsilon)L \preceq \hat{L} \preceq (1 + \varepsilon)L \quad \text{for a parameter } \varepsilon \in (0, 1).$$

The spectral equivalence (3.6) ensures that the sparse graph is structurally similar to the original graph. We plan to construct a spectral sparsifier $\hat{L}$ by randomly sampling edges according to a carefully chosen probability distribution. Introduce the random matrix $Y \in \mathbb{H}_d^+(\mathbb{R})$ with distribution

$$\mathbb{P}\left\{Y = \frac{n-1}{\varrho_{ij}} \Delta_{ij}\right\} = \frac{w_{ij} \varrho_{ij}}{n-1} \quad \text{for each } i < j \text{ where } \varrho_{ij} > 0.$$

In other words, Y is the Laplacian of a weighted edge between a pair (i, j) of vertices, chosen with probability proportional to the weight w_{ij} and the effective resistance ϱ_{ij} . It is easy to confirm that $\mathbb{E}[Y] = L$, so the random matrix Y is an unbiased estimator for the Laplacian L of the original graph. By averaging q independent copies of Y , we obtain the Laplacian of a random sparse graph with at most q edges:

$$(3.7) \quad \hat{L} := \frac{1}{q} \sum_{k=1}^q Y_k \quad \text{where } Y_k \sim Y \text{ i.i.d.}$$

How many edges q suffice to achieve the approximation guarantee (3.6)?

THEOREM 3.3 (Graph sparsification [41, Thm. 1]). *Let L be the Laplacian (3.5) of a weighted graph on n vertices, and let $\hat{L}$ be the Laplacian (3.7) of a random sparse graph with at most q edges. With probability at least $1 - \delta$, the random matrix $\hat{L}$ is an ε -approximation of L in the sense of (3.6) provided that $q \geq 3\varepsilon^{-2}n \log(2n/\delta)$.*

The result ensures that we can approximate a graph on n vertices by a sparse graph with $q = \mathcal{O}(n \log n)$ edges, regardless of the number m of edges in the original graph. We pay an extra factor ε^{-2} in the number of edges q to achieve accuracy ε . Although the randomized method is simple and efficient, there is a more expensive deterministic algorithm that finds an ε -spectral sparsifier with just $q = \mathcal{O}(n/\varepsilon^2)$ edges [8].

Proof sketch. Define the linear map $\Phi(A) := L^{-1/2} A L^{-1/2}$ on symmetric matrices that act on $\text{range}(L)$. The spectral equivalence (3.6) is the same as the requirement that $\|\Phi(\hat{L} - L)\| \leq \varepsilon$. The matrix inside the norm can be written as an independent sum of centered random matrices acting on an $(n - 1)$ -dimensional space:

$$S := \Phi(\hat{L} - L) = \sum_{k=1}^q \frac{1}{q} \Phi(Y_k - \mathbb{E}[Y_k]) =: \sum_{k=1}^q X_k.$$

To activate the matrix Bernstein inequality (Theorem 2.1), check that

$$B = \sup \|\mathbf{X}_k\| = \frac{n-1}{q} \quad \text{and} \quad v(\mathbf{S}) \leq \frac{n-1}{q}.$$

The result now follows from the tail inequality (2.3). For details, see [51, sect. 5.2]. $\square$

3.4 Quantum tomography. The state of a finite-dimensional quantum system, such as a register in a quantum computer, is characterized by a finite-dimensional psd matrix. We can probe the system by taking measurements, but the laws of quantum mechanics imply that each measurement returns a random number. As a consequence, methods for reconstructing the state of a quantum system lead to problems involving random matrices. Matrix concentration tools offer a quick way to analyze quantum state estimators. This section summarizes a lecture by Richard Kueng (see [51, Lecture 3]), which distills ideas from a paper of Guță et al. [19]. Related ideas animate the theory of classical shadows, a major recent advance in quantum information science [27].

We model a quantum system by means of its *density matrix*, which is a complex psd matrix with trace one. The convex, compact set $\mathcal{D}_d := \{\boldsymbol{\rho} \in \mathbb{H}_d^+(\mathbb{C}) : \text{Tr}[\boldsymbol{\rho}] = 1\}$ collects the density matrices of dimension d . Next, a *quantum measurement* is a family $\{\mathbf{H}_1, \dots, \mathbf{H}_m\} \subset \mathbb{H}_d^+(\mathbb{C})$ of psd matrices that partitions the identity: $\sum_{j=1}^m \mathbf{H}_j = \mathbf{I}$. When we apply the measurement to a quantum system with density matrix $\boldsymbol{\rho} \in \mathcal{D}_d$, two things happen. First, we sample a discrete random variable J with distribution

$$(3.8) \quad \mathbb{P}\{J = j \mid \boldsymbol{\rho}\} = \text{Tr}[\mathbf{H}_j \boldsymbol{\rho}] \quad \text{for } j = 1, \dots, m.$$

Second, the quantum system “collapses.” Therefore, to characterize a quantum system, we must prepare several copies in the same state and measure them repeatedly to obtain statistical evidence about the state.

For simplicity, we consider a special type of quantum measurement, called a *complex projective 2-design*. Here, each of the measurement matrices has rank one: $\mathbf{H}_j = (d/m)\mathbf{u}_j\mathbf{u}_j^*$ where $\mathbf{u}_j \in \mathbb{C}^d$ is a unit-norm vector. Moreover, these measurements support the reconstruction property

$$(3.9) \quad \frac{1}{m} \sum_{j=1}^m \text{Tr}[\mathbf{u}_j\mathbf{u}_j^* \mathbf{A}] \cdot \mathbf{u}_j\mathbf{u}_j^* = \frac{1}{(d+1)d} (\mathbf{A} + \text{Tr}[\mathbf{A}] \cdot \mathbf{I}) \quad \text{for each matrix } \mathbf{A} \in \mathbb{H}_d(\mathbb{C}).$$

The vectors $(\mathbf{u}_1, \dots, \mathbf{u}_m)$ composing a complex projective 2-design are arranged very regularly over the complex sphere. Examples include systems of “mutually unbiased” orthonormal bases and systems of equiangular lines [51, Ex. 3.3]. This type of measurement system can stably distinguish any pair of density matrices.

To reconstruct a quantum system with density matrix $\boldsymbol{\rho} \in \mathcal{D}_d$, we perform a measurement to sample the random variable J . Then we build a random Hermitian matrix $\mathbf{Y} \in \mathbb{H}_d(\mathbb{C})$ according to the rule

$$\mathbf{Y} = (d+1)\mathbf{u}_J\mathbf{u}_J^* - \mathbf{I}.$$

The random matrix $\mathbf{Y}$ serves as an unbiased estimator for the density matrix:

$$\mathbb{E}[\mathbf{Y}] = \sum_{j=1}^m ((d+1)\mathbf{u}_j\mathbf{u}_j^* - \mathbf{I}) \cdot \mathbb{P}\{J = j \mid \boldsymbol{\rho}\} = \boldsymbol{\rho}.$$

This calculation follows from (3.8) and (3.9). By repeating the measurement on n independent (i.e., unentangled) copies of the quantum system, each with density matrix $\boldsymbol{\rho}$, we can obtain an estimator

$$(3.10) \quad \mathbf{S}_n := \frac{1}{n} \sum_{k=1}^n \mathbf{Y}_k \quad \text{where } \mathbf{Y}_k \sim \mathbf{Y} \text{ i.i.d.}$$

How many independent measurements suffice to approximate the underlying density matrix?

THEOREM 3.4 (Quantum tomography). *Let $\boldsymbol{\rho} \in \mathcal{D}_d$ be a density matrix with dimension d . Using a complex projective 2-design (3.9), perform a quantum measurement on each of n independent quantum systems with common state $\boldsymbol{\rho}$. Form the sample average estimator $\mathbf{S}_n$, as in (3.10). For $\varepsilon \in (0, 1)$,*

$$\mathbb{P}\{\|\mathbf{S}_n - \boldsymbol{\rho}\| \geq \varepsilon\} \leq 2d \cdot \exp\left(\frac{-3n\varepsilon^2}{14d}\right).$$

In particular, the sample complexity $n \geq 5\varepsilon^{-2}d \log(2d/\delta)$ yields a failure probability no greater than $\delta \in (0, 1)$.

Proof sketch. To apply the matrix Bernstein inequality (Theorem 2.1) to the estimator (3.10), simply compute

$$B = \sup \|n^{-1}\mathbf{Y}\| = \frac{d}{n} \quad \text{and} \quad v(\mathbf{S}_n) = \frac{1}{n} \|\mathbb{E}[\mathbf{Y}^2]\| = \frac{1}{n} \|(d-1)\boldsymbol{\rho} + d\mathbf{I}\| \leq \frac{2d-1}{n}.$$

The statement follows quickly from (2.3). $\square$

While the sample average estimator (3.10) is unbiased, it generally does not produce a density matrix. Moreover, the trace norm yields a more interpretable distance between density matrices than the spectral norm. To address these concerns, define the *projected state estimator* using the Frobenius norm:

$$\hat{\boldsymbol{\rho}}_n \in \arg \min_{\boldsymbol{\sigma} \in \mathcal{D}_d} \|\boldsymbol{\sigma} - \mathbf{S}_n\|_F.$$

Lemma 3.9 of [51] ensures that the projected state estimator satisfies the trace-norm bound

$$\|\hat{\boldsymbol{\rho}}_n - \boldsymbol{\rho}\|_1 \leq 4r \|\mathbf{S}_n - \boldsymbol{\rho}\| \quad \text{where } r := \text{rank}(\boldsymbol{\rho}).$$

As a consequence, we arrive at an upper bound for the sample complexity of the projected state estimator:

$$\mathbb{P}\{\|\hat{\boldsymbol{\rho}}_n - \boldsymbol{\rho}\|_1 \geq \varepsilon\} \leq \delta \quad \text{when } n \geq 80\varepsilon^{-2}r^2d \log(2d/\delta).$$

This bound implies that the projected state estimator $\hat{\boldsymbol{\rho}}_n$ nearly achieves the *optimal* sample complexity that is possible in this setting [20].

4 Concentration for matrix martingales. The full majesty of the framework for exponential matrix concentration comes into view when we extend it to dynamically evolving sequences of matrices.

4.1 Matrix martingales. Fix a probability space. A *filtration* is an increasing sequence of sigma-algebras $\mathcal{F}_0 \subseteq \mathcal{F}_1 \subseteq \mathcal{F}_2 \subseteq \dots$ inside the master sigma-algebra. A *matrix martingale* is a sequence $(\mathbf{S}_0, \mathbf{S}_1, \mathbf{S}_2, \dots)$ of complex random matrices with common dimension $d_1 \times d_2$ that satisfies three properties. For each index $k = 0, 1, 2, \dots$, the sequence is

1. *Adapted.* Each random matrix $\mathbf{S}_k$ is measurable with respect to $\mathcal{F}_k$.
2. *Integrable.* The expected norm is finite: $\mathbb{E}\|\mathbf{S}_k\| < +\infty$.
3. *Status quo.* The conditional expectation $\mathbb{E}[\mathbf{S}_{k+1} | \mathcal{F}_k] = \mathbf{S}_k$.

The *difference sequence* consists of the matrices $\mathbf{X}_k := \mathbf{S}_k - \mathbf{S}_{k-1}$ with $k \geq 1$. Each member of the difference sequence is conditionally centered: $\mathbb{E}[\mathbf{X}_k | \mathcal{F}_{k-1}] = \mathbf{0}$. Matrix martingales model a wide range of examples.

Example 4.1 (Adapted sums). Let $\mathbf{S}_0 := \mathbf{A}$ be a fixed matrix. For each index $k \geq 1$, suppose that $\mathbf{S}_k := \mathbf{S}_{k-1} + \mathbf{X}_k$ where the increments are conditionally centered: $\mathbb{E}[\mathbf{X}_k | \mathbf{X}_0, \dots, \mathbf{X}_{k-1}] = \mathbf{0}$. Then $(\mathbf{S}_0, \mathbf{S}_1, \mathbf{S}_2, \dots)$ is a matrix martingale with respect to the filtration $\mathcal{F}_k := \sigma(\mathbf{S}_0, \mathbf{X}_1, \dots, \mathbf{X}_k)$ defined for $k \geq 0$.

Example 4.2 (Adapted products). Let $\mathbf{S}_0 := \mathbf{A}$ be a fixed matrix. Suppose that $\mathbf{S}_k := \mathbf{Y}_k \mathbf{S}_{k-1}$ where the factors satisfy $\mathbb{E}[\mathbf{Y}_k | \mathbf{S}_0, \dots, \mathbf{S}_{k-1}] = \mathbf{I}$. Then $(\mathbf{S}_0, \mathbf{S}_1, \mathbf{S}_2, \dots)$ composes a matrix martingale with respect to the filtration $\mathcal{F}_k := \sigma(\mathbf{S}_0, \mathbf{Y}_1, \dots, \mathbf{Y}_k)$ defined for $k \geq 0$.

Example 4.3 (Lévy–Doob martingale). Fix a filtration $\mathcal{F}_0 \subseteq \mathcal{F}_1 \subseteq \mathcal{F}_2 \subseteq \dots$. Let $\mathbf{S}$ be a random $d_1 \times d_2$ matrix with $\mathbb{E}\|\mathbf{S}\| < +\infty$. Construct the sequence of conditional expectations:

$$\mathbf{S}_k := \mathbb{E}[\mathbf{S} | \mathcal{F}_k] \quad \text{for each } k = 0, 1, 2, \dots$$

Then $(\mathbf{S}_0, \mathbf{S}_1, \mathbf{S}_2, \dots)$ composes a martingale with respect to the filtration.

To avoid technicalities, we only consider finite martingale sequences $(\mathbf{S}_0, \dots, \mathbf{S}_n)$. Furthermore, we require the elements of the martingale to be uniformly bounded in the sense that $\sup \|\mathbf{S}_k\| \leq B_\infty < +\infty$ for each index $k = 0, \dots, n$. The filtration may be suppressed if it is determined by context.

4.2 Matrix Freedman. To analyze a matrix martingale, the most valuable theorem is the matrix Freedman inequality, originally due to Oliveira [37]. The author exploited the subadditivity (2.8) of the matrix log-mgf to refine Oliveira’s result and to establish some other matrix martingale inequalities [46]. The version stated here follows from [46, Cor. 1.3, Thm. 3.1].

THEOREM 4.4 (Matrix Freedman). *Consider a finite matrix martingale sequence $(\mathbf{S}_0, \dots, \mathbf{S}_n)$ consisting of $d_1 \times d_2$ matrices, real or complex. Suppose that the elements of the difference sequence $(\mathbf{X}_1, \dots, \mathbf{X}_n)$ are uniformly bounded: $\|\mathbf{X}_k\| \leq B$. Introduce the conditional quadratic variation sequence:*

$$V_k := \left\| \sum_{j=1}^k \mathbb{E}[\mathbf{X}_j \mathbf{X}_j^* | \mathcal{F}_{j-1}] \right\| \vee \left\| \sum_{j=1}^k \mathbb{E}[\mathbf{X}_j^* \mathbf{X}_j | \mathcal{F}_{j-1}] \right\| \quad \text{for } k = 1, \dots, n.$$

For parameters $v, t \geq 0$, we have the probability inequalities

$$(4.1) \quad \mathbb{P} \left\{ \exists k : \|\mathbf{S}_k - \mathbf{S}_0\| \geq \sqrt{2vt} + Bt/3 \text{ and } V_k \leq v \right\} \leq (d_1 + d_2) e^{-t},$$

$$(4.2) \quad \mathbb{P} \{ \exists k : \|\mathbf{S}_k - \mathbf{S}_0\| \geq t \text{ and } V_k \leq v \} \leq (d_1 + d_2) \exp \left(\frac{-t^2/2}{v + Bt/3} \right).$$

The rest of this section presents a proof of Theorem 4.4 that combines ideas from my paper [46] and from the survey of Howard et al. [23]. First, we dilate on the meaning of the result.

The conditional quadratic variation is the martingale analogue of the matrix variance; cf. (2.4). At each step j of the process, we compute the expected “squares” of the element $\mathbf{X}_j$ in the difference sequence, given data $\mathcal{F}_{j-1}$ about the previous position of the martingale, and we combine the results. Thus, V_k is a random variable that reflects the total observed volatility of the matrix martingale up to time k . When the difference sequence $(\mathbf{X}_1, \dots, \mathbf{X}_n)$ is statistically independent, each V_k reduces to the matrix variance $v(\mathbf{S}_k)$.

The probability inequalities (4.1) and (4.2) state that the norm of the martingale element $\|\mathbf{S}_k\|$ is unlikely to be large when the conditional quadratic variation V_k is bounded above. The first inequality (4.1) formulates a clean bound on the tail probability, while the second inequality provides a direct bound on the level of $\|\mathbf{S}_k\|$. Let us emphasize that both estimates provide uniform control on the *entire trajectory* of the matrix martingale.

4.3 The log-mgf supermartingale. As in the proof of Theorem 2.1, we can and will employ the Hermitian dilation (2.5) to restrict our attention to Hermitian random matrices.

Consider a bounded Hermitian matrix martingale $(\mathbf{S}_0, \dots, \mathbf{S}_n)$ taking values in $\mathbb{H}_d(\mathbb{C})$, with initial condition $\mathbf{S}_0 := \mathbf{0}$ and with difference sequence $(\mathbf{X}_1, \dots, \mathbf{X}_n)$. For a parameter $\theta \in \mathbb{R}$, define the *conditional* log-mgfs and their partial sums:

$$(4.3) \quad \Xi_{\mathbf{X}_k | \mathcal{F}_{k-1}}(\theta) := \log \mathbb{E} [e^{\theta \mathbf{X}_k} | \mathcal{F}_{k-1}] \quad \text{and} \quad \mathbf{W}_k(\theta) := \sum_{j=1}^k \Xi_{\mathbf{X}_j | \mathcal{F}_{j-1}}(\theta).$$

The element $\mathbf{W}_k(\theta) \in \mathbb{H}_d(\mathbb{C})$ of the log-mgf process reflects the total volatility of the martingale up to time k .

To track the evolution of the matrix martingale, let us pass to a real-valued random sequence:

$$(4.4) \quad M_0(\theta) := d \quad \text{and} \quad M_k(\theta) := \text{Tr} e^{\theta \mathbf{S}_k - \mathbf{W}_k(\theta)} \quad \text{where } k = 1, \dots, n.$$

Applying the subadditivity rule (2.8) for the matrix log-mgf conditionally, we quickly verify that $\mathbb{E}[M_k(\theta) | \mathcal{F}_{k-1}] \leq M_{k-1}(\theta)$ for each $k = 1, \dots, n$. In other words, $(M_k(\theta) : k = 0, \dots, n)$ composes a positive supermartingale. We can study the trajectory of the matrix martingale by bounding the supermartingale [37, 46]. The approach here is adapted from Howard et al. [23].

PROPOSITION 4.5 (Matrix martingale: Eigenvalue bounds). *Consider a finite matrix martingale $(\mathbf{S}_0, \dots, \mathbf{S}_n)$ taking values in $\mathbb{H}_d(\mathbb{C})$ with log-mgf process $(\mathbf{W}_1(\theta), \dots, \mathbf{W}_n(\theta))$, as in (4.3). For parameters $\theta, \alpha > 0$,*

$$\mathbb{P} \{ \exists k : \lambda_{\max}(\mathbf{S}_k) \geq \theta^{-1}(\alpha + \lambda_{\max}(\mathbf{W}_k(\theta))) \} \leq d \cdot e^{-\alpha}.$$

Proof. The log-mgf supermartingale (4.4) admits a lower bound involving the eigenvalues of the sequences:

$$\log M_k(\theta) = \log \text{Tr} e^{\theta \mathbf{S}_k - \mathbf{W}_k(\theta)} \geq \theta \lambda_{\max}(\mathbf{S}_k) - \lambda_{\max}(\mathbf{W}_k(\theta)) \quad \text{for each } \theta > 0.$$

For the inequality, note that the trace dominates the maximum eigenvalue; then apply spectral mapping and Weyl’s perturbation theorem [9, Cor. III.2.6]. Ville’s inequality for positive supermartingales [23, Lem. 1] implies

$$\mathbb{P} \{ \exists k : \theta \lambda_{\max}(\mathbf{S}_k) - \lambda_{\max}(\mathbf{W}_k(\theta)) \geq \alpha \} \leq \mathbb{P} \{ \exists k : \log M_k(\theta) \geq \alpha \} \leq e^{-\alpha} \cdot \mathbb{E}[M_0(\theta)] = d \cdot e^{-\alpha}.$$

This is the advertised result. $\square$

4.4 Proof of matrix Freedman. To invoke Proposition 4.5, we seek a bound for the maximum eigenvalue of the log-mgf process $(\mathbf{W}_k(\theta) : k = 0, \dots, n)$. In the setting of Theorem 4.4, each element of the difference sequence admits a bound on the conditional log-mgf of the Bernstein type (2.9):

$$\Xi_{\mathbf{X}_j | \mathcal{F}_{j-1}}(\theta) \preceq g(\theta) \cdot \mathbb{E}[\mathbf{X}_j^2 | \mathcal{F}_{j-1}] \quad \text{where} \quad g(\theta) := \frac{\theta^2/2}{1 - B|\theta|/3} \quad \text{and} \quad |\theta| < 3/B.$$

By Weyl's monotonicity theorem [9, Cor. III.2.3],

$$\lambda_{\max}(\mathbf{W}_k(\theta)) = \lambda_{\max}\left(\sum_{j=1}^k \Xi_{\mathbf{X}_j | \mathcal{F}_{j-1}}(\theta)\right) \leq g(\theta) \cdot \lambda_{\max}\left(\sum_{j=1}^k \mathbb{E}[\mathbf{X}_j^2 | \mathcal{F}_{j-1}]\right) = g(\theta) \cdot V_k.$$

Proposition 4.5 yields

$$\begin{aligned} d \cdot e^{-\alpha} &\geq \mathbb{P}\{\exists k : \lambda_{\max}(\mathbf{S}_k) \geq \theta^{-1}(\alpha + \lambda_{\max}(\mathbf{W}_k(\theta)))\} \geq \mathbb{P}\{\exists k : \lambda_{\max}(\mathbf{S}_k) \geq \theta^{-1}(\alpha + g(\theta)V_k)\} \\ &\geq \mathbb{P}\{\exists k : \lambda_{\max}(\mathbf{S}_k) \geq \theta^{-1}(\alpha + g(\theta)v) \text{ and } V_k \leq v\}. \end{aligned}$$

In the Hermitian setting, we obtain an analogue of the first matrix Freedman tail bound (4.1) by selecting $\alpha = t$ and $\theta = \sqrt{t}/(\sqrt{v/2} + B\sqrt{t}/3)$. The second tail bound (4.2) follows from the first tail bound by inverting the function of t and making a simple bound; see [10, sects. 2.4 and 2.8]. Finally, we extend to a general (non-Hermitian) matrix martingale by applying these two bounds to the Hermitian dilation (2.5).

5 Matrix martingales: Applications. In this section, we outline some applications of the matrix Freedman inequality in computational mathematics. This technique is particularly valuable for handling sequences of matrices that evolve adaptively, and it provides uniform control over the entire trajectory of the process.

5.1 Uniform covariance estimation. Consider a centered random vector $\mathbf{y} \in \mathbb{R}^d$ that is uniformly bounded: $\mathbb{E}[\mathbf{y}] = \mathbf{0}$ and $\|\mathbf{y}\| \leq L$. Its covariance matrix takes the form $\Sigma := \mathbb{E}[\mathbf{y}\mathbf{y}^*]$. Given i.i.d. samples $\mathbf{y}_1, \mathbf{y}_2, \mathbf{y}_3, \dots$ from the distribution, we can construct a sequence of empirical covariance estimates:

$$(5.1) \quad \hat{\Sigma}_n := \frac{1}{n} \sum_{k=1}^n \mathbf{y}_k \mathbf{y}_k^* \quad \text{for each } n = 1, 2, 3, \dots$$

We would like to obtain a confidence region for the entire sequence of covariance estimates, ensuring simultaneous coverage at all times. The following statement is adapted from Howard et al. [24, sect. 4.3].

THEOREM 5.1 (Uniform covariance estimation [24, sect. 4.3]). *Let $\mathbf{y} \in \mathbb{R}^d$ be a bounded, centered random vector with $\|\mathbf{y}\| \leq L$ and covariance matrix Σ . As in (5.1), construct the sequence $(\hat{\Sigma}_n : n \in \mathbb{N})$ of empirical covariance estimates. The following uniform confidence region attains confidence level $1 - \delta \in (0, 1)$:*

$$\mathbb{P}\left\{\forall n \in \mathbb{N} : \frac{\|\hat{\Sigma}_n - \Sigma\|}{\|\Sigma\|} \leq \sqrt{8\beta_n} + \frac{2}{3}\beta_n\right\} \geq 1 - \delta \quad \text{where} \quad \beta_n := \frac{L^2}{\|\Sigma\|} \cdot \frac{\log(2d/\delta) + \log \log(en)}{n}.$$

The ratio $r := L^2/\|\Sigma\|$ reflects the number of dimensions where the random vector has substantial fluctuation. We need $n = \mathcal{O}(r \log d)$ samples to reliably estimate the covariance in these directions. As the number n of samples continues to increase, the size of the confidence region decreases at a rate $\sqrt{n^{-1} \log \log n}$. The log-log factor is required to correct for occasional extreme fluctuations. More precise asymptotic results follow from the law of the iterated logarithm in a Banach space [31, Thm. 8.2].

Proof sketch. For a parameter $N \in \mathbb{N}$, introduce the finite martingale sequence

$$\mathbf{S}_0 := \mathbf{0} \quad \text{and} \quad \mathbf{S}_n := \sum_{k=1}^n (\mathbf{y}_k \mathbf{y}_k^* - \Sigma) =: \sum_{k=1}^n \mathbf{X}_k \quad \text{for } n = 1, \dots, N.$$

The difference sequence $(\mathbf{X}_k)$ consists of i.i.d. centered random matrices with dimension d . A short calculation gives

$$B = \sup \|\mathbf{X}_k\| \leq L^2 \quad \text{and} \quad V_n = \left\| \sum_{k=1}^n \mathbb{E}[\mathbf{X}_k^2] \right\| \leq nL^2 \cdot \|\Sigma\|.$$

The matrix Freedman inequality (4.1) ensures that

$$\mathbb{P}\left\{\exists n \leq N : \|\mathbf{S}_n\| \geq \sqrt{2nL^2\|\Sigma\|t} + L^2t/3\right\} \leq 2d \cdot e^{-t}.$$

Dividing each element $\mathbf{S}_n$ by the factor $n \cdot \|\Sigma\|$ and limiting the index n to the range $[N/2, N]$, we find

$$\mathbb{P}\left\{\exists n \in [N/2, N] : \frac{\|\hat{\Sigma}_n - \Sigma\|}{\|\Sigma\|} \geq \sqrt{\frac{4L^2t}{n\|\Sigma\|}} + \frac{L^2t}{3n\|\Sigma\|}\right\} \leq 2d \cdot e^{-t}.$$

The key insight is to apply the last inequality on dyadic intervals: $N = 2^j$ for $j = 1, 2, 3, \dots$ at the level $t_j = \log(4j^2d/\delta) \leq 2\log(2d\log(en)/\delta)$. Combine the results with a union bound, then take the complement. $\square$

5.2 Cholesky decomposition with stochastic rounding. Most numerical algorithms proceed in stages, and the numerical errors compound at each step of the process. Traditional analyses model these errors deterministically, and error bounds increase linearly with the number of steps in the algorithm. More recently, researchers have started to make a close accounting of error propagation under probabilistic models, where the errors accumulate more slowly because of cancellations. Martingale methods offer an elegant technique for tracking random errors that depend on the past history of the algorithm. Connolly and Higham [15] have employed this approach to give a probabilistic analysis of the Householder QR algorithm. In the same spirit, this section offers a stylized analysis of the Cholesky factorization algorithm under a stochastic rounding model (cf. subsection 3.2).

Let $\mathbf{A} \in \mathbb{H}_d^{++}(\mathbb{R})$ be a positive-definite matrix. For clarity of interpretation, we assume that $\text{diag}(\mathbf{A}) = \mathbf{I}$. Cholesky's method iteratively produces a factorization $\mathbf{A} = \mathbf{C}\mathbf{C}^*$ where $\mathbf{C} \in \mathbb{M}_d(\mathbb{R})$ is lower triangular. Define the initial residual $\mathbf{R}_0 := \mathbf{A}$ and the initial triangular factor $\mathbf{C}_0 := \mathbf{0}$. At the k th step, reduce the residual $\mathbf{R}_{k-1}$ by performing a Schur complement with respect to its k th column:

$$\mathbf{c}_k := \frac{\mathbf{R}_{k-1}(:, k)}{\sqrt{\mathbf{R}_{k-1}(k, k)}} \quad \text{and} \quad \mathbf{C}_k := \mathbf{C}_{k-1} + \mathbf{c}_k \mathbf{e}_k^* \quad \text{and} \quad \mathbf{R}_k := \mathbf{R}_{k-1} - \mathbf{c}_k \mathbf{c}_k^*.$$

This procedure zeroes out the k th row and column from $\mathbf{R}_{k-1}$, while placing $\mathbf{c}_k$ in the k th column of $\mathbf{C}_k$. The matrix $\mathbf{C}_k$ remains lower triangular. At each step, the residual and the partial triangular factorization compose the original matrix: $\mathbf{A} = \mathbf{R}_k + \mathbf{C}_k \mathbf{C}_k^*$. After d steps, $\mathbf{R}_d = \mathbf{0}$ and $\mathbf{C}_d \mathbf{C}_d^* = \mathbf{A}$. Each step requires $\mathcal{O}(d^2)$ arithmetic operations, for a total cost of $\mathcal{O}(d^3)$.

In actual practice, we commit floating-point errors when we reduce the residual—but we can extract a column from the residual without further error. If we employ stochastic rounding, then the residual update becomes

$$\mathbf{R}_k := \text{sr}(\mathbf{R}_{k-1} - \mathbf{c}_k \mathbf{c}_k^*) = \mathbf{R}_{k-1} - \mathbf{c}_k \mathbf{c}_k^* + \mathbf{Y}_k.$$

The stochastic error $\mathbf{Y}_k$ is a random symmetric matrix that depends on the previous residual. The sigma-algebra $\mathcal{F}_{k-1} := \sigma(\mathbf{R}_0, \dots, \mathbf{R}_{k-1})$ captures the history of the algorithm through step $k-1$. As in subsection 3.2, each entry of $\mathbf{Y}_k$ satisfies the rounding properties in (3.4), conditional on $\mathcal{F}_{k-1}$, and each entry is rounded independently of the others (modulo symmetry). We frame the assumptions that

$$(5.2) \quad \mathbb{E}[\mathbf{Y}_k | \mathcal{F}_{k-1}] = \mathbf{0} \quad \text{and} \quad \|\mathbb{E}[\mathbf{Y}_k^2 | \mathcal{F}_{k-1}]\| \leq 2u^2 \cdot \|\mathbf{A}\| \quad \text{and} \quad \|\mathbf{Y}_k\|_{\max} \leq u.$$

The third property follows because $\text{diag}(\mathbf{A}) = \mathbf{I}$ and $\text{diag}(\mathbf{R}_{k-1})$ only decreases at each step. The second property holds if the computed residual $\mathbf{R}_{k-1}$ remains psd and satisfies $\|\mathbf{R}_{k-1}\| \leq 2\|\mathbf{A}\|$. Under mild assumptions, these conditions can be justified with some additional effort. We also demand that no underflow or overflow occurs.

The computed residuals $\mathbf{R}_k$ and the partial triangular factorizations $\mathbf{C}_k \mathbf{C}_k^*$ induce a sequence of approximations of the original matrix:

$$\mathbf{S}_k := \mathbf{R}_k + \mathbf{C}_k \mathbf{C}_k^* = \mathbf{R}_k + \sum_{j=1}^k \mathbf{c}_j \mathbf{c}_j^* \quad \text{for } k = 0, 1, 2, \dots, d.$$

Using these approximations, we can calculate the error in the completed triangular factorization $\mathbf{C} := \mathbf{C}_d$.

$$\mathbf{C}\mathbf{C}^* - \mathbf{A} = \mathbf{S}_d - \mathbf{S}_0 = \sum_{k=1}^d (\mathbf{S}_k - \mathbf{S}_{k-1}) = \sum_{k=1}^d (\mathbf{R}_k - \mathbf{R}_{k-1} + \mathbf{c}_k \mathbf{c}_k^*) = \sum_{k=1}^d \mathbf{Y}_k.$$

In other words, the sequence of approximations composes a matrix martingale with difference sequence $(\mathbf{Y}_k : k = 1, \dots, d)$. We can use this martingale to assess the accumulated rounding error in the Cholesky decomposition.

THEOREM 5.2 (Cholesky decomposition with stochastic rounding). *Let $\mathbf{A} \in \mathbb{H}_d^{++}(\mathbb{R})$ be a $d \times d$ positive-definite matrix with $\text{diag}(\mathbf{A}) = \mathbf{I}$. As described above, compute the Cholesky decomposition $\mathbf{C}\mathbf{C}^*$ with a stochastic rounding procedure that satisfies the assumption (5.2). The resulting factorization admits the error bound*

$$\mathbb{P}\left\{\|\mathbf{C}\mathbf{C}^* - \mathbf{A}\| \geq \mathbf{u} \cdot \left(2\sqrt{d\|\mathbf{A}\|}t + \frac{1}{3}t\right)\right\} \leq 2d \cdot e^{-t}.$$

For the Cholesky decomposition, the classic deterministic rounding error analysis [21, Chap. 10] gives a bound for the entrywise error $\|\mathbf{C}\mathbf{C}^* - \mathbf{A}\|_{\max}$ that is proportional to $\mathbf{u} \cdot d$. For stochastic rounding, the stylized analysis here indicates that the entrywise error is typically on the order of $\mathbf{u} \cdot \sqrt{d\|\mathbf{A}\|}$. The factor $\sqrt{d}$ improvement can be significant. A full comparison with deterministic rounding falls outside our ambit.

Proof sketch. We need to consider a more fine-grained martingale to exploit the fact that matrix entries are rounded independently. Decompose each error matrix as $\mathbf{Y}_k = \sum_{i \leq j} \mathbf{X}_{ijk}$, where $\mathbf{X}_{ijk}$ describes the rounding error in the (i, j) and (j, i) entries of $\mathbf{Y}_k$. For each k , the family $(\mathbf{X}_{ijk} : i \leq j)$ is conditionally independent, given $\mathcal{F}_{k-1}$. According to (5.2), the parameters in the matrix Freedman inequality (Theorem 4.4) satisfy

$$B = \sup \|\mathbf{X}_{ijk}\| = \sup \|\mathbf{Y}_k\|_{\max} \leq \mathbf{u},$$

$$V_d = \left\| \sum_{k=1}^d \sum_{i \leq j} \mathbb{E}[\mathbf{X}_{ijk}^2 | \mathcal{F}_{k-1}] \right\| = \left\| \sum_{k=1}^d \mathbb{E}[\mathbf{Y}_k^2 | \mathcal{F}_{k-1}] \right\| \leq 2\mathbf{u}^2 d \cdot \|\mathbf{A}\|.$$

The result follows from the matrix Freedman inequality (4.1). $\square$

5.3 Fast Laplacian solvers. As in subsection 3.3, consider a weighted, connected, undirected graph on n vertices with Laplacian matrix $\mathbf{L} \in \mathbb{H}_n^+(\mathbb{R})$. A linear system in the Laplacian matrix takes the form

$$(5.3) \quad \mathbf{L}\mathbf{u} = \mathbf{f} \quad \text{where} \quad \mathbf{1}^* \mathbf{u} = \mathbf{1}^* \mathbf{f} = 0.$$

This type of linear system arises in a dizzying range of applications, including numerical discretizations of elliptic PDEs, clustering and partitioning of data, and network flow problems [44]. The basic algorithm for (5.3) starts by computing a Cholesky decomposition $\mathbf{L} = \mathbf{C}\mathbf{C}^*$, and it solves the linear system by two steps of triangular substitution. In general, this approach requires $\mathcal{O}(n^3)$ operations, which is often prohibitive. Can we do better?

Spielman and Teng [42] gave an affirmative answer by devising a theoretical Laplacian solver that runs in time linear in the number m of edges in the graph and polylogarithmic in the number n of vertices. After a train of refinements, Kyng and Sachdeva [29] designed a *practical* Laplacian solver that runs in time $\mathcal{O}(m \log^3 n)$. Their algorithm, **SparseCholesky**, performs an approximate Cholesky decomposition by randomly sparsifying the Schur complement at each step. The analysis relies on matrix martingales. See [51] for another account.

Here is the intuition. When we apply the Cholesky method to a graph Laplacian, each Schur complement step eliminates a vertex from the graph by adding a clique to the remaining vertices of the graph. The method is expensive because the clique involves up to $n - k$ vertices and up to $(n - k)^2$ edges at step k . Yet, as we saw in subsection 3.3, we can dramatically sparsify a graph by randomly sampling edges in proportion to their effective resistances. In this spirit, **SparseCholesky** starts with residual $\mathbf{R}_0 := \mathbf{L}$ and triangular factor $\mathbf{C}_0 := \mathbf{0}$. It performs the steps

$$\mathbf{c}_k := \frac{\mathbf{R}_{k-1}(:, k)}{\sqrt{\mathbf{R}_{k-1}(k, k)}} \quad \text{and} \quad \mathbf{C}_k := \mathbf{C}_{k-1} + \mathbf{c}_k \mathbf{e}_k^* \quad \text{and} \quad \mathbf{R}_k := \mathbf{R}_{k-1} - \text{Sparsify}(\mathbf{c}_k \mathbf{c}_k^*).$$

The **Sparsify** method produces a random, unbiased approximation to the matrix $\mathbf{c}_k \mathbf{c}_k^*$, which contains the Laplacian of the clique. After sparsification, the remaining number of edges equals the number of neighbors of vertex k in the residual graph $\mathbf{R}_{k-1}$. The full algorithm requires several other innovations: it splits edges into smaller pieces to reduce their weights; it maintains estimates for the effective resistances to implement the sparsification step; and it eliminates vertices in a random order to limit the average number of neighbors. As in subsection 5.2, this procedure results in a matrix martingale, namely $\mathbf{S}_k := \mathbf{R}_k + \mathbf{C}_k \mathbf{C}_k^*$ for $k = 0, 1, 2, \dots, d$. We can analyze this martingale using the matrix Freedman inequality (Theorem 4.4). Here is the outcome.

THEOREM 5.3 (Sparse Cholesky approximation [29]). *Let $\mathbf{L} \in \mathbb{H}_d^+(\mathbb{R})$ be a weighted graph Laplacian, as in (3.5). With probability at least $1 - d^{-1}$, the **SparseCholesky** algorithm produces an approximate Cholesky factorization $\mathbf{C}\mathbf{C}^*$ with the spectral guarantee*

$$0.5\mathbf{L} \preceq \mathbf{C}\mathbf{C}^* \preceq 1.5\mathbf{L}.$$

The lower-triangular matrix $\mathbf{C}$ has $\mathcal{O}(m \log^2 n)$ nonzero entries. The expected runtime is $\mathcal{O}(m \log^3 n)$ operations.

Once we have computed the approximate Cholesky decomposition $\mathbf{L} \approx \mathbf{C}\mathbf{C}^*$, we can solve the linear system (5.3) using preconditioned conjugate gradient (PCG) with the triangular preconditioner $\mathbf{C}^{-1}$.

PROPOSITION 5.4 (Fast Laplacian solvers [29]). *Given the triangular matrix $\mathbf{C}$ computed by the **SparseCholesky** algorithm, the PCG algorithm solves the consistent Laplacian system (5.3) to relative error ε in the Dirichlet energy norm after $\mathcal{O}(m \log^2(n) \log(1/\varepsilon))$ arithmetic operations.*

For a sparse graph with $m = \mathcal{O}(n)$ edges, the total cost of solving the Laplacian system (5.3) to fixed accuracy is just $\mathcal{O}(n \log^3 n)$ operations. For a dense graph with $m = \mathcal{O}(n^2)$ edges, the total cost is $\mathcal{O}(n^2 \log^3 n)$ operations. Both results compare favorably with the $\mathcal{O}(n^3)$ cost of the classic method based on a full Cholesky factorization.

5.4 Randomized Trotter formulas. The evolution of a quantum mechanical system is governed by a Hamiltonian matrix $\mathbf{H} \in \mathbb{H}_d(\mathbb{C})$. The state of the system at time $t \geq 0$ takes the form

$$\Phi_t(\rho) := e^{-it\mathbf{H}} \cdot \rho \cdot e^{+it\mathbf{H}} \quad \text{where the initial state } \rho \in \mathcal{D}_n.$$

A basic goal of quantum information science is to simulate the action Φ_t induced by a complicated Hamiltonian through the application of simpler Hamiltonians. Randomized methods can provide effective algorithms for this task. This section summarizes a result of Chen et al. [14] that builds on work of Campbell [13].

Suppose that the Hamiltonian can be expressed as a sum of many “simple” Hermitian terms:

$$(5.4) \quad \mathbf{H} := \sum_{m=1}^M \mathbf{H}_m \quad \text{with} \quad L := \sum_{m=1}^M \|\mathbf{H}_m\| \quad \text{and each } \mathbf{H}_m \in \mathbb{H}_d(\mathbb{C}).$$

The statistic L measures the interaction strength within the Hamiltonian. This model encompasses applications in quantum chemistry and quantum physics. For clarity, fix the time horizon for the simulation at $t = 1$. Given a parameter $n \in \mathbb{N}$, construct a random unitary matrix by drawing a random term from the Hamiltonian according to an importance sampling distribution:

$$(5.5) \quad \mathbf{Y} := e^{-i\mathbf{Z}/n} \quad \text{where} \quad \mathbb{P}\left\{\mathbf{Z} = \frac{L}{\|\mathbf{H}_j\|} \mathbf{H}_j\right\} = \frac{\|\mathbf{H}_j\|}{L} \quad \text{for } j = 1, \dots, M.$$

To approximate the unitary matrix $\mathbf{U} := e^{-i\mathbf{H}}$, form a product of n independent copies of $\mathbf{Y}$:

$$(5.6) \quad \mathbf{Q} := \mathbf{Y}_n \cdots \mathbf{Y}_3 \mathbf{Y}_2 \mathbf{Y}_1 \quad \text{where } \mathbf{Y}_k \sim \mathbf{Y} \text{ i.i.d. for } k = 1, \dots, n.$$

In other words, we simulate the action of the full Hamiltonian $\mathbf{H}$ by performing a series of short simulations with the simpler Hamiltonians $\mathbf{H}_m$. This is a randomized variant of the Lie–Trotter–Suzuki product formulas that are widely used to approximate the matrix exponential of a sum. We can analyze the quality of the simulation using matrix martingale methods.

THEOREM 5.5 (Random product formulas [14, Thm. 1]). *Consider a Hamiltonian of the form (5.4) with interaction strength L , and form the unitary matrix $\mathbf{U} := e^{-i\mathbf{H}}$. Fix parameters $\varepsilon, \delta \in (0, 1)$ where $\varepsilon \leq L$. For $n \geq 40\varepsilon^{-2}L^2 \log(2d/\delta)$, construct a random unitary matrix $\mathbf{Q}$ with n factors via the random product formula (5.6). Then*

$$\mathbb{P}\left\{\sup_{\rho \in \mathcal{D}_n} \|\mathbf{Q} \cdot \rho \cdot \mathbf{Q}^* - \mathbf{U} \cdot \rho \cdot \mathbf{U}^*\|_1 \geq 2\varepsilon\right\} \leq \delta.$$

This result states that, with high probability, the random product formula (5.6) approximates the evolution of the quantum Hamiltonian for *all* initial states. The number n of factors in the product does not depend on the number M of summands in the Hamiltonian (5.4), but it grows quadratically with the interaction strength L . For comparison, the number of factors in a deterministic product formula depends linearly on the number M of summands in the Hamiltonian, but more weakly on the interaction strength L . As a consequence, random product formulas are most effective for short-time simulations of Hamiltonians with many summands.

Proof sketch. The quantity of interest depends on the spectral-norm distance between the two unitaries:

$$E := \sup_{\rho \in \mathcal{D}_n} \|\mathbf{Q} \cdot \rho \cdot \mathbf{Q}^* - \mathbf{U} \cdot \rho \cdot \mathbf{U}^*\|_1 \leq 2 \|\mathbf{Q} - \mathbf{U}\|.$$

This point follows from the triangle inequality, unitary invariance, and the operator ideal property of the trace norm. The right-hand side consists of a random fluctuation term and a deterministic bias:

$$\|\mathbf{Q} - \mathbf{U}\| \leq \|\mathbf{Q} - \mathbb{E}[\mathbf{Q}]\| + \|\mathbb{E}[\mathbf{Q}] - \mathbf{U}\| =: \text{flux} + \text{bias}.$$

By statistical independence of the factors in (5.6), the expectation $\mathbb{E}[\mathbf{Q}] = (\mathbb{E}[\mathbf{Y}])^n$.

The bias term admits a bound that depends on the interaction strength L of the Hamiltonian:

$$\text{bias} = \|(\mathbb{E}[\mathbf{Y}])^n - (\mathrm{e}^{-\mathrm{i}\mathbf{H}/n})^n\| \leq n \cdot \|\mathbb{E}[\mathbf{Y}] - \mathrm{e}^{-\mathrm{i}\mathbf{H}/n}\| \leq n \cdot \left(\frac{L}{n}\right)^2 = \frac{L^2}{n}.$$

The first inequality follows from repeated application of the triangle inequality and unitary invariance, while the second inequality requires a second-order approximation of the matrix exponential. See [14, Lems. 3.5 and 3.6] for the details. We insist that $L^2/n \leq \varepsilon/2$ to control the bias term.

For the fluctuation term, we introduce a Lévy–Doob martingale. Let $\mathcal{F}_k := \sigma(\mathbf{Y}_1, \dots, \mathbf{Y}_k)$, and define

$$\mathbf{S}_k := \mathbb{E}[\mathbf{Q} | \mathcal{F}_k] = (\mathbb{E} \mathbf{Y})^{n-k} \mathbf{Y}_k \dots \mathbf{Y}_1 \quad \text{for } k = 0, \dots, n.$$

Observe that $\mathbf{S}_n = \mathbf{Q}$ and $\mathbf{S}_0 = \mathbb{E}[\mathbf{Q}]$. Construct the difference sequence $\mathbf{X}_k := \mathbf{S}_k - \mathbf{S}_{k-1}$ for $k = 1, \dots, n$. By first-order approximation of the matrix exponential, the difference sequence satisfies uniform bounds

$$\|\mathbf{X}_k\| \leq \|\mathbf{Y}_k - \mathbb{E}[\mathbf{Y}_k]\| \leq \frac{2L}{n} \quad \text{and} \quad \|\mathbb{E}[\mathbf{X}_k \mathbf{X}_k^* | \mathcal{F}_{k-1}]\| \vee \|\mathbb{E}[\mathbf{X}_k^* \mathbf{X}_k | \mathcal{F}_{k-1}]\| \leq \|\mathbf{Y}_k - \mathbb{E}[\mathbf{Y}_k]\|^2 \leq \frac{4L^2}{n^2}.$$

See [14, Prop. 3.3] for the details. The matrix Freedman inequality (4.2) now implies a probability inequality for $\text{flux} = \|\mathbf{S}_n - \mathbf{S}_0\|$ at level $t = \varepsilon/2$. Combine the bias and fluctuation to get a tail bound for the quantity E . $\square$

6 Matrix concentration: Extensions. Over the last few years, researchers have made significant advances in understanding the behavior of the independent sum model. These results require a detour into the study of Gaussian random matrices, which serve as ideal models for independent sums. First, we outline some improvements of matrix concentration for Gaussian random matrices, and then we explain how Gaussian models can capture the spectral features of an independent sum of random matrices. Last, we mention several other random matrix models that remain under active study.

6.1 Variance statistics. We can describe the behavior of random matrix models more fully by employing a broader family of variance statistics. Introduce the real inner product

$$\langle \mathbf{B}, \mathbf{A} \rangle_{\text{Re}} := \text{Re Tr}[\mathbf{B}^* \mathbf{A}] \quad \text{for } \mathbf{A}, \mathbf{B} \in \mathbb{F}^{d_1 \times d_2}.$$

The variance function of a random matrix $\mathbf{S} \in \mathbb{F}^{d_1 \times d_2}$ is defined as

$$(6.1) \quad \text{Var}[\mathbf{S}] : \mathbb{F}^{d_1 \times d_2} \rightarrow \mathbb{R}_+ \quad \text{where} \quad \text{Var}[\mathbf{S}](\mathbf{A}) := \text{Var}[\langle \mathbf{S}, \mathbf{A} \rangle_{\text{Re}}] \quad \text{for } \mathbf{A} \in \mathbb{F}^{d_1 \times d_2}.$$

Variance functions are in correspondence with psd quadratic forms on $\mathbb{F}^{d_1 \times d_2}$, treated as a real linear space.

The variance function $\text{Var}[\mathbf{S}]$ packs up all the second-order statistics of the random matrix. In particular, it completely determines the matrix variance $v(\mathbf{S})$, defined in (2.1). We can also introduce the *weak variance*, which is related to the tail behavior of the sum: $v_*(\mathbf{S}) := \sup_{\|\mathbf{A}\|_1 \leq 1} \text{Var}[\mathbf{S}](\mathbf{A})$. The weak variance $v_*(\mathbf{S})$ is always smaller than the matrix variance $v(\mathbf{S})$, and often much smaller.

6.2 Gaussian matrix models. A random matrix is *Gaussian* when each of its linear marginals is Gaussian. More precisely, $\mathbf{Z} \in \mathbb{F}^{d_1 \times d_2}$ is Gaussian if and only if the random variable $\langle \mathbf{Z}, \mathbf{A} \rangle_{\text{Re}}$ follows a real Gaussian distribution for each matrix $\mathbf{A} \in \mathbb{F}^{d_1 \times d_2}$. A Gaussian random matrix is fully characterized by its expectation $\mathbb{E}[\mathbf{Z}]$ and by its variance function $\text{Var}[\mathbf{Z}]$. For a specified expectation matrix $\mathbf{M} \in \mathbb{F}^{d_1 \times d_2}$ and a variance function $\mathbf{V} : \mathbb{F}^{d_1 \times d_2} \rightarrow \mathbb{R}_+$, we denote the (unique) Gaussian distribution with these statistics by $\text{NORMAL}(\mathbf{M}, \mathbf{V})$.

6.3 Second-order matrix Khinchin inequalities. The norm of a Gaussian matrix $\mathbf{Z}$ satisfies an elegant matrix concentration bound, called the *matrix Khinchin inequality*, originally due to Lust-Piquard [34]:

$$(6.2) \quad \sqrt{(2/\pi)v(\mathbf{Z})} \leq \mathbb{E} \|\mathbf{Z} - \mathbb{E}[\mathbf{Z}]\| \leq \sqrt{2v(\mathbf{Z}) \log(d_1 + d_2)}.$$

The lower bound involves Gaussian isoperimetry [30, Cor. 3]; the upper bound relies on exponential matrix concentration arguments [48, Thm. 4.1.1]. Both inequalities are saturated.

To improve on the matrix Khinchin inequality (6.2), we need extra information about the variance function of the Gaussian matrix. Define the *interaction energy* to be the maximum eigenvalue of the quadratic form:

$$w(\mathbf{Z}) := 2 \sup_{\|\mathbf{A}\|_F \leq 1} \text{Var}[\mathbf{Z}](\mathbf{A}).$$

The interaction energy $w(\mathbf{Z})$ can be much smaller than the matrix variance $v(\mathbf{Z})$, but they are incomparable. Heuristically, the interaction energy is small when $\mathbf{Z}$ is “highly noncommutative.” This statistic supports a second-order refinement of the matrix Khinchin inequality [5, Cor. 2.2 and Lem. 2.5]:

$$(6.3) \quad \mathbb{E} \|\mathbf{Z} - \mathbb{E}[\mathbf{Z}]\| \leq 2\sqrt{v(\mathbf{Z})} + \text{Const} \cdot [v(\mathbf{Z})w(\mathbf{Z}) \log^3(d_1 + d_2)]^{1/4}.$$

When $w(\mathbf{Z}) \ll v(\mathbf{Z})$, then the result (6.3) improves over (6.2). As a simple example, consider a real Ginibre matrix $\mathbf{G} \in \mathbb{M}_d(\mathbb{R})$, which is a $d \times d$ matrix with i.i.d. standardized Gaussian entries. The matrix variance $v(\mathbf{G}) = d$, while the interaction energy $w(\mathbf{G}) = 2$. Thus,

$$(6.2) \quad \text{implies} \quad \mathbb{E} \|\mathbf{G}\| \leq \sqrt{2d \log(2d)};$$

$$(6.3) \quad \text{implies} \quad \mathbb{E} \|\mathbf{G}\| \leq 2\sqrt{d} + \text{Const} \cdot [2d \log^3(2d)]^{1/4}.$$

The latter inequality is the stronger one, and its first term is numerically sharp.

Roughly speaking, the logarithmic factor in (6.2) arises when the Hermitian dilations of two independent copies $\mathcal{H}(\mathbf{Z}), \mathcal{H}(\mathbf{Z}')$ of the Gaussian matrix “almost commute” with each other, in an appropriate sense. To quantify this insight, inspired by free probability, I introduced a subtle statistic called the *matrix alignment* [50, Def. 3.1]. For a special class of Gaussian matrices, I also established a preliminary version of the second-order Khinchin inequality [50, Thm. 3.1] that reveals how the matrix alignment arises.

The finished result (6.3) was obtained through additional insights of Bandeira and colleagues (see [4, 5]). The latter papers use matrix alignment statistics to compare the spectral norm of a Gaussian matrix with the spectral norm of a free probability model (namely, an operator semicircle) that admits explicit formulas. The role of the interaction energy $w(\mathbf{Z})$ is to provide a computable bound for the matrix alignment [4, Rem. 3.3]. Similar results hold for other spectral properties, such as the support of the spectrum and the mean singular value distribution.

6.4 Independent sums: Gaussian comparison. Let us return to a more general setting. Consider an independent sum of bounded, centered random matrices with dimension $d_1 \times d_2$:

$$\mathbf{S} := \sum_{k=1}^n \mathbf{X}_k \quad \text{where } \mathbb{E}[\mathbf{X}_k] = \mathbf{0} \text{ and } \|\mathbf{X}_k\| \leq B.$$

Let $\mathbf{V} := \text{Var}[\mathbf{S}]$ be the variance function of the sum. The multivariate central limit theorem suggests that we should compare $\mathbf{S}$ with the centered Gaussian random matrix $\mathbf{Z} \sim \text{NORMAL}(\mathbf{0}, \mathbf{V})$, but it is not clear how to quantify the difference between their distributions.

Under weak assumptions, we can obtain detailed nonasymptotic comparisons for *spectral properties* of the independent sum $\mathbf{S}$ and the Gaussian model $\mathbf{Z}$. Brailovskaya and van Handel [11, Cor. 2.7] established that

$$(6.4) \quad |\mathbb{E} \|\mathbf{S}\| - \mathbb{E} \|\mathbf{Z}\|| \lesssim [v(\mathbf{S})B \log^2(d)]^{1/3} + B \log(d) + \sqrt{v_*(\mathbf{S}) \log(d)} \quad \text{where } d := d_1 + d_2.$$

In view of the matrix Khinchin inequality (6.2), the Gaussian matrix satisfies $\mathbb{E} \|\mathbf{Z}\| \approx \sqrt{v(\mathbf{Z})} = \sqrt{v(\mathbf{S})}$. The weak variance term $\sqrt{v_*(\mathbf{S})}$ is often negligible. Therefore, in case $B^2 \ll v(\mathbf{S})$, the discrepancy between the expected norms is much smaller than the expected norm of the Gaussian matrix. In other words, when the bound B on the summands is sufficiently small, we can accurately approximate $\mathbb{E} \|\mathbf{S}\|$ by way of estimates for $\mathbb{E} \|\mathbf{Z}\|$. The second-order matrix Khinchin inequality (6.3) provides one such estimate.

Brailovskaya and van Handel [11] developed Gaussian comparisons, similar to (6.4), for other spectral statistics of the independent sum, including the support of the spectrum and the mean distribution of the singular values. Their arguments rely on Stein's method, cumulant expansions, Möbius inversion, and a selection of matrix inequalities. My paper [53] offers an alternative proof of their results, based on the method of exchangeable pairs. I recently obtained another type of Gaussian comparison [52] for the *minimum* eigenvalue of an independent sum of *psd* matrices; the proof exploits a deep result from matrix analysis, called Stahl's theorem [43].

We conclude with a critical observation. For independent sums, the latest matrix concentration results, such as (6.3) and (6.4), offer new insights. Nevertheless, it may be difficult to deploy these inequalities in applications, such as the ones in section 3, because we often lack fine-grained information about the variance function of the random matrix model. As a consequence, classic matrix concentration tools (e.g., Theorem 2.1) and the more recent refinements play complementary roles.

6.5 Other models. While the matrix concentration theory for the independent sum model is rather complete, other types of random matrix models remain mysterious and have continued to attract attention. Citations are limited to a few typical papers, as there is not enough space to do justice to this rich literature.

The nonlinear random matrix model concerns a matrix-valued function of underlying random variables that are usually (but not always) independent. Results for this model include matrix Efron–Stein inequalities [39] and (local) matrix Poincaré inequalities [25, 26]. There are also specialized results for matrix-valued polynomial chaos [7]. A related challenge arises from random matrix models with dependencies, such as the sum of fixed matrices modulated by negatively correlated scalar random variables; some recent progress on these questions appears in [28]. In addition to matrix martingales, researchers have also treated other types of sequential processes, such as matrix-valued Markov chains; for example, see [36]. Another line of work [6] pursues extensions of matrix concentration to random tensors.

Many of the models in the last paragraph emerged from problems in combinatorics and algorithms, and the theoretical results have led to satisfying progress. Thus, research on matrix concentration tools remains active, and it continues to make a profound impact on applications.

Acknowledgments. JAT gratefully acknowledges support under ONR Award N00014-24-1-2223 and the Caltech Carver Mead New Adventures Fund. Many thanks are due to Ethan Epperly for his valuable insights on a preliminary draft of this article.

References

- [1] R. AHLWEDE AND A. WINTER, *Strong converse for identification via quantum channels*, IEEE Trans. Inform. Theory, 48 (2002), pp. 569–579, <https://doi.org/10.1109/18.985947>.
- [2] Z. BAI AND J. W. SILVERSTEIN, *Spectral Analysis of Large Dimensional Random Matrices*, 2nd ed., Springer, New York, 2010, <https://doi.org/10.1007/978-1-4419-0661-8>.
- [3] J. BAIK, P. DEIFT, AND K. JOHANSSON, *On the distribution of the length of the longest increasing subsequence of random permutations*, J. Amer. Math. Soc., 12 (1999), pp. 1119–1178, <https://doi.org/10.1090/S0894-0347-99-00307-0>.
- [4] A. S. BANDEIRA AND M. T. BOEDIHARDJO, *The Spectral Norm of Gaussian Matrices with Correlated Entries*, preprint, <https://arxiv.org/abs/2104.02662>, 2021.
- [5] A. S. BANDEIRA, M. T. BOEDIHARDJO, AND R. VAN HANDEL, *Matrix concentration inequalities and free probability*, Invent. Math., 234 (2023), pp. 419–487, <https://doi.org/10.1007/s00222-023-01204-6>.
- [6] A. S. BANDEIRA, S. GOPI, H. JIANG, K. LUCCA, AND T. ROTHVOSS, *Tensor concentration inequalities: A geometric approach*, in Proceedings of the 57th Annual ACM Symposium on Theory of Computing, STOC '25, Association for Computing Machinery, New York, 2025, pp. 822–832, <https://doi.org/10.1145/3717823.3718188>.

- [7] A. S. BANDEIRA, K. LUCCA, P. NIZIC-NIKOLAC, AND R. VAN HANDEL, *Matrix chaos inequalities and chaos of combinatorial type*, in Proceedings of the 57th Annual ACM Symposium on Theory of Computing, STOC '25, Association for Computing Machinery, New York, 2025, pp. 795–805, <https://doi.org/10.1145/3717823.3718238>.
- [8] J. BATSON, D. A. SPIELMAN, AND N. SRIVASTAVA, *Twice-Ramanujan sparsifiers*, SIAM Rev., 56 (2014), pp. 315–334, <https://doi.org/10.1137/130949117>.
- [9] R. BHATIA, *Matrix Analysis*, Springer-Verlag, New York, 1997, <https://doi.org/10.1007/978-1-4612-0653-8>.
- [10] S. BOUCHERON, G. LUGOSI, AND P. MASSART, *Concentration Inequalities: A Nonasymptotic Theory of Independence*, Oxford University Press, Oxford, 2013, <https://doi.org/10.1093/acprof:oso/9780199535255.001.0001>.
- [11] T. BRAILOVSKAYA AND R. VAN HANDEL, *Universality and sharp matrix concentration inequalities*, Geom. Funct. Anal., 34 (2024), pp. 1734–1838, <https://doi.org/10.1007/s00039-024-00692-9>.
- [12] T. BRAILOVSKAYA AND R. VAN HANDEL, *Extremal random matrices with independent entries and matrix superconcentration inequalities*, Ann. Probab., to appear.
- [13] E. CAMPBELL, *Random compiler for fast Hamiltonian simulation*, Phys. Rev. Lett., 123 (2019), 070503, <https://doi.org/10.1103/PhysRevLett.123.070503>.
- [14] C.-F. CHEN, H.-Y. HUANG, R. KUENG, AND J. A. TROPP, *Concentration for random product formulas*, Phys. Rev. X Quantum, 2 (2021), 040305, <https://doi.org/10.1103/PRXQuantum.2.040305>.
- [15] M. P. CONNOLLY AND N. J. HIGHAM, *Probabilistic rounding error analysis of Householder QR factorization*, SIAM J. Matrix Anal. Appl., 44 (2023), pp. 1146–1163, <https://doi.org/10.1137/22M1514817>.
- [16] P. G. CONSTANTINE, *Active Subspaces: Emerging Ideas for Dimension Reduction in Parameter Studies*, SIAM, Philadelphia, 2015, <https://doi.org/10.1137/1.9781611973860>.
- [17] M. CROCI, M. FASI, N. J. HIGHAM, T. MARY, AND M. MIKAITIS, *Stochastic rounding: Implementation, error analysis and applications*, Roy. Soc. Open Sci., 9 (2022), 211631, <https://doi.org/10.1098/rsos.211631>.
- [18] D. A. FREEDMAN, *On tail probabilities for martingales*, Ann. Probab., 3 (1975), pp. 100–118, <https://doi.org/10.1214/aop/1176996452>.
- [19] M. GUȚĂ, J. KAHN, R. KUENG, AND J. A. TROPP, *Fast state tomography with optimal error bounds*, J. Phys. A, 53 (2020), 204001, <https://doi.org/10.1088/1751-8121/ab8111>.
- [20] J. HAAH, A. W. HARROW, Z. JI, X. WU, AND N. YU, *Sample-optimal tomography of quantum states*, IEEE Trans. Inform. Theory, 63 (2017), pp. 5628–5641, <https://doi.org/10.1109/TIT.2017.2719044>.
- [21] N. J. HIGHAM, *Accuracy and Stability of Numerical Algorithms*, 2nd ed., SIAM, Philadelphia, 2002, <https://doi.org/10.1137/1.9780898718027>.
- [22] R. A. HORN AND C. R. JOHNSON, *Matrix Analysis*, 2nd ed., Cambridge University Press, Cambridge, 2013.
- [23] S. R. HOWARD, A. RAMDAS, J. MCAULIFFE, AND J. SEKHON, *Time-uniform Chernoff bounds via nonnegative supermartingales*, Probab. Surv., 17 (2020), pp. 257–317, <https://doi.org/10.1214/18-PS321>.
- [24] S. R. HOWARD, A. RAMDAS, J. MCAULIFFE, AND J. SEKHON, *Time-uniform, nonparametric, nonasymptotic confidence sequences*, Ann. Statist., 49 (2021), pp. 1055–1080, <https://doi.org/10.1214/20-AOS1991>.
- [25] D. HUANG AND J. A. TROPP, *From Poincaré inequalities to nonlinear matrix concentration*, Bernoulli, 27 (2021), pp. 1724–1744, <https://doi.org/10.3150/20-bej1289>.

- [26] D. HUANG AND J. A. TROPP, *Nonlinear matrix concentration via semigroup methods*, Electron. J. Probab., 26 (2021), 8, <https://doi.org/10.1214/20-EJP578>.
- [27] H.-Y. HUANG, R. KUENG, AND J. PRESKILL, *Predicting many properties of a quantum system from very few measurements*, Nature Phys., 16 (2020), pp. 1050–1057, <https://doi.org/10.1038/s41567-020-0932-7>.
- [28] T. KAUFMAN, R. KYNG, AND F. SOLDÁ, *Scalar and matrix Chernoff bounds from ℓ_∞ -independence*, in Proceedings of the 2022 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA), SIAM, Philadelphia; ACM, New York, 2022, pp. 3732–3753, <https://doi.org/10.1137/1.9781611977073.147>.
- [29] R. KYNG AND S. SACHDEVA, *Approximate Gaussian elimination for Laplacians—fast, sparse, and simple*, in 57th Annual IEEE Symposium on Foundations of Computer Science—FOCS 2016, IEEE Computer Society, Los Alamitos, CA, 2016, pp. 573–582, <https://doi.org/10.1109/FOCS.2016.68>.
- [30] R. LATAŁA AND K. OLESZKIEWICZ, *Gaussian measures of dilations of convex symmetric sets*, Ann. Probab., 27 (1999), pp. 1922–1938, <https://doi.org/10.1214/aop/1022677554>.
- [31] M. LEDOUX AND M. TALAGRAND, *Probability in Banach Spaces: Isoperimetry and Processes*, Springer-Verlag, Berlin, 1991, <https://doi.org/10.1007/978-3-642-20212-4>.
- [32] E. H. LIEB, *Convex trace functions and the Wigner–Yanase–Dyson conjecture*, Advances in Math., 11 (1973), pp. 267–288, [https://doi.org/10.1016/0001-8708\(73\)90011-X](https://doi.org/10.1016/0001-8708(73)90011-X).
- [33] N. LINIAL, E. LONDON, AND Y. RABINOVICH, *The geometry of graphs and some of its algorithmic applications*, Combinatorica, 15 (1995), pp. 215–245, <https://doi.org/10.1007/BF01200757>.
- [34] F. LUST-PIQUARD, *Inégalités de Khintchine dans C_p ($1 < p < \infty$)*, C. R. Acad. Sci. Paris Sér. I Math., 303 (1986), pp. 289–292.
- [35] H. L. MONTGOMERY, *The pair correlation of zeros of the zeta function*, in Analytic Number Theory (Proc. Sympos. Pure Math., Vol. XXIV, St. Louis Univ., St. Louis, Mo., 1972), American Mathematical Society, Providence, RI, 1973, pp. 181–193.
- [36] J. NEEMAN, B. SHI, AND R. WARD, *Concentration inequalities for sums of Markov-dependent random matrices*, Inf. Inference, 13 (2024), iaae032, <https://doi.org/10.1093/imaiai/iaae032>.
- [37] R. I. OLIVEIRA, *Concentration of the Adjacency Matrix and of the Laplacian in Random Graphs with Independent Edges*, preprint, <https://arxiv.org/abs/0911.0600>, 2010.
- [38] L. PASTUR AND M. SHCHERBINA, *Eigenvalue Distribution of Large Random Matrices*, American Mathematical Society, Providence, RI, 2011, <https://doi.org/10.1090/surv/171>.
- [39] D. PAULIN, L. MACKEY, AND J. A. TROPP, *Efron-Stein inequalities for random matrices*, Ann. Probab., 44 (2016), pp. 3431–3473, <https://doi.org/10.1214/15-AOP1054>.
- [40] M. RUDELSON, *Random vectors in the isotropic position*, J. Funct. Anal., 164 (1999), pp. 60–72, <https://doi.org/10.1006/jfan.1998.3384>.
- [41] D. A. SPIELMAN AND N. SRIVASTAVA, *Graph sparsification by effective resistances*, SIAM J. Comput., 40 (2011), pp. 1913–1926, <https://doi.org/10.1137/080734029>.
- [42] D. A. SPIELMAN AND S.-H. TENG, *Nearly-linear time algorithms for graph partitioning, graph sparsification, and solving linear systems*, in Proceedings of the 36th Annual ACM Symposium on Theory of Computing (STOC), ACM Press, New York, 2004, pp. 81–90, <https://doi.org/10.1145/1007352.1007372>.
- [43] H. R. STAHL, *Proof of the BMV conjecture*, Acta Math., 211 (2013), pp. 255–290, <https://doi.org/10.1007/s11511-013-0104-z>.

- [44] S.-H. TENG, *The Laplacian paradigm: Emerging algorithms for massive graphs*, in Theory and Applications of Models of Computation, Springer, Berlin, 2010, pp. 2–14, https://doi.org/10.1007/978-3-642-13562-0_2.
- [45] N. TOMCZAK-JAEGERMANN, *The moduli of smoothness and convexity and the Rademacher averages of trace classes S_p ($1 \leq p < \infty$)*, Studia Math., 50 (1974), pp. 163–182, <http://eudml.org/doc/217886>.
- [46] J. A. TROPP, *Freedman's inequality for matrix martingales*, Electron. Commun. Probab., 16 (2011), pp. 262–270, <https://doi.org/10.1214/ECP.v16-1624>.
- [47] J. A. TROPP, *User-friendly tail bounds for sums of random matrices*, Found. Comput. Math., 12 (2012), pp. 389–434, <https://doi.org/10.1007/s10208-011-9099-z>.
- [48] J. A. TROPP, *An introduction to matrix concentration inequalities*, Found. Trends Mach. Learn., 8 (2015), pp. 1–230, <https://doi.org/10.1561/22000000048>.
- [49] J. A. TROPP, *The expected norm of a sum of independent random matrices: An elementary approach*, in High Dimensional Probability VII, Springer, Cham, 2016, pp. 173–202, https://doi.org/10.1007/978-3-319-40519-3_8.
- [50] J. A. TROPP, *Second-order matrix concentration inequalities*, Appl. Comput. Harmon. Anal., 44 (2018), pp. 700–736, <https://doi.org/10.1016/j.acha.2016.07.005>.
- [51] J. A. TROPP, *Matrix Concentration & Computational Linear Algebra*, CMS Lecture Notes 2019-01, Caltech, Pasadena, CA, 2019, <https://doi.org/10.7907/nbx6-vm05>.
- [52] J. A. TROPP, *Comparison Theorems for the Minimum Eigenvalue of a Random Positive-Semidefinite Matrix*, preprint, <https://doi.org/10.48550/arXiv.2501.16578>, 2025.
- [53] J. A. TROPP, *Universality Laws for Random Matrices via Exchangeable Pairs*, manuscript, 2025.
- [54] R. VERSHYNIN, *High-Dimensional Probability: An Introduction with Applications in Data Science*, Cambridge University Press, Cambridge, 2018, <https://doi.org/10.1017/9781108231596>.
- [55] D. V. VOICULESCU, K. J. DYKEMA, AND A. NICA, *Free Random Variables*, American Mathematical Society, Providence, RI, 1992, <https://doi.org/10.1090/crmn/001>.
- [56] M. J. WAINWRIGHT, *High-Dimensional Statistics: A Non-Asymptotic Viewpoint*, Cambridge University Press, Cambridge, 2019, <https://doi.org/10.1017/9781108627771>.
- [57] E. P. WIGNER, *Characteristic vectors of bordered matrices with infinite dimensions*, Ann. of Math. (2), 62 (1955), pp. 548–564, <https://doi.org/10.2307/1970079>.
- [58] J. WISHART, *The generalised product moment distribution in samples from a normal multivariate population*, Biometrika, 20A (1928), pp. 32–52, <https://doi.org/10.1093/BIOMET/20A.1-2.32>.