

**USER-FRIENDLY TAIL BOUNDS
FOR MATRIX MARTINGALES**

JOEL A. TROPP

Technical Report No. 2011-01
January 2011

APPLIED & COMPUTATIONAL MATHEMATICS
CALIFORNIA INSTITUTE OF TECHNOLOGY
mail code 217-50 · pasadena, ca 91125

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE JAN 2011		2. REPORT TYPE		3. DATES COVERED 00-00-2011 to 00-00-2011	
4. TITLE AND SUBTITLE User-Friendly Tail Bounds for Matrix Martingales				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) California Institute of Technology, Applied and Computational Mathematics, Pasadena, CA, 91125				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT This report presents probability inequalities for sums of adapted sequences of random self-adjoint matrices. The results frame simple, easily verifiable hypotheses on the summands, and they yield strong conclusions about the large-deviation behavior of the maximum eigenvalue of the sum. The methods also specialize to sums of independent random matrices.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 14	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

USER-FRIENDLY TAIL BOUNDS FOR MATRIX MARTINGALES

JOEL A. TROPP

ABSTRACT. This report presents probability inequalities for sums of adapted sequences of random, self-adjoint matrices. The results frame simple, easily verifiable hypotheses on the summands, and they yield strong conclusions about the large-deviation behavior of the maximum eigenvalue of the sum. The methods also specialize to sums of independent random matrices.

1. MAIN RESULTS

This technical report is a companion to two other works, the papers “User-friendly tail bounds for sums of random matrices” [Tro10c] and “Freedman’s inequality for matrix martingales” [Tro10a]. Since this report is intended as a supplement, we have removed most of the background discussion, citations to related work, and auxiliary commentary that places the research in a wider context. We recommend that the reader peruse the original papers before studying this report.

The paper [Tro10a] describes a martingale technique that leads to an extension of Freedman’s inequality in the matrix setting, which is similar to the result [Oli10a, Thm. 1.2]. The purpose of this work is to show how the arguments from [Tro10a] allow us to establish the matrix probability inequalities for sums of *independent* random matrices that appear in [Tro10c]. The discussion here also contains some new probability inequalities for sums of adapted sequences of random matrices; we have removed these results from the other two papers because they are somewhat specialized.

1.1. Roadmap. The rest of the report is organized as follows. The balance of §1 provides an overview of the main results for sums of independent random matrices. Section 2 contains the main technical ingredients for the proof. Sections 3–5 complete the proofs of the matrix probability inequalities for adapted sequences. Appendix A provides an overview of the background material that we require.

1.2. Rademacher and Gaussian Series. Let $\|\cdot\|$ denote the usual norm for operators on a Hilbert space, which returns the largest singular value of its argument, and let $\lambda_{\max}$ denote the algebraically largest eigenvalue of a self-adjoint matrix. The extreme eigenvalues of a Rademacher series with self-adjoint matrix coefficients exhibit normal concentration.

Theorem 1.1 (Matrix Rademacher and Gaussian Series). *Consider a finite sequence $\{\mathbf{A}_k\}$ of fixed self-adjoint matrices with dimension d , and let $\{\varepsilon_k\}$ be a finite sequence of independent Rademacher variables. Compute the norm of the sum of squared coefficient matrices:*

$$\sigma^2 := \left\| \sum_k \mathbf{A}_k^2 \right\|. \tag{1.1}$$

Date: 25 April 2010. Revised on 15 June 2010, 10 August 2010, 14 November 2010, and 16 January 2011.

Key words and phrases. Discrete-time martingale, large deviation, probability inequality, random matrix, sum of independent random variables.

2010 *Mathematics Subject Classification.* Primary: 60B20. Secondary: 60F10, 60G50, 60G42.

JAT is with Applied and Computational Mathematics, MC 305-16, California Inst. Technology, Pasadena, CA 91125. E-mail: jtropp@acm.caltech.edu. Research supported by ONR award N00014-08-1-0883, DARPA award N66001-08-1-2065, and AFOSR award FA9550-09-1-0643.

For all $t \geq 0$,

$$\mathbb{P} \left\{ \lambda_{\max} \left(\sum_k \varepsilon_k \mathbf{A}_k \right) \geq t \right\} \leq d \cdot e^{-t^2/2\sigma^2}. \quad (1.2)$$

In particular,

$$\mathbb{P} \left\{ \left\| \sum_k \varepsilon_k \mathbf{A}_k \right\| \geq t \right\} \leq 2d \cdot e^{-t^2/2\sigma^2}. \quad (1.3)$$

The same bounds hold when we replace $\{\varepsilon_k\}$ by a finite sequence of independent standard normal random variables.

See [Tro10c, §4] for a detailed discussion of Theorem 1.1, which indicates that it is essentially sharp. We present the proof in §5.

1.3. Sums of Random Semidefinite Matrices. Chernoff bounds describe the upper and lower tails of a sum of nonnegative random variables. In the matrix case, the analogous results concern a sum of positive-semidefinite random matrices. The matrix Chernoff bound shows that the extreme eigenvalues of this sum exhibit the same binomial-type behavior as in the scalar setting.

Theorem 1.2 (Matrix Chernoff). *Consider a finite sequence $\{\mathbf{X}_k\}$ of independent, random, positive-semidefinite matrices with dimension d . Suppose that*

$$\lambda_{\max}(\mathbf{X}_k) \leq R \quad \text{almost surely.}$$

Compute the eigenvalues of the sum of the expectations:

$$\mu_{\min} := \lambda_{\min} \left(\sum_k \mathbb{E} \mathbf{X}_k \right) \quad \text{and} \quad \mu_{\max} := \lambda_{\max} \left(\sum_k \mathbb{E} \mathbf{X}_k \right).$$

Then

$$\begin{aligned} \mathbb{P} \left\{ \lambda_{\min} \left(\sum_k \mathbf{X}_k \right) \leq (1 - \delta) \mu_{\min} \right\} &\leq d \cdot \left[\frac{e^{-\delta}}{(1 - \delta)^{1-\delta}} \right]^{\mu_{\min}/R} \quad \text{for } \delta \in [0, 1), \text{ and} \\ \mathbb{P} \left\{ \lambda_{\max} \left(\sum_k \mathbf{X}_k \right) \geq (1 + \delta) \mu_{\max} \right\} &\leq d \cdot \left[\frac{e^{\delta}}{(1 + \delta)^{1+\delta}} \right]^{\mu_{\max}/R} \quad \text{for } \delta \geq 0. \end{aligned}$$

We establish Theorem 1.2 in §3, where it emerges as a consequence of Theorem 3.1, a Chernoff inequality for sums of adapted sequences of positive-semidefinite matrices.

1.4. Adding Variance Information. In the scalar case, a well-known inequality of Bernstein shows that the sum exhibits normal concentration near its mean with variance controlled by the variance of the sum. On the other hand, the tail of the sum decays subexponentially on a scale determined by a uniform upper bound for the summands. Sums of independent random matrices exhibit the same type of behavior, where the normal concentration depends on a matrix generalization of the variance and the tails are controlled by a uniform bound for the largest eigenvalue of each summand.

Theorem 1.3 (Matrix Bernstein). *Consider a finite sequence $\{\mathbf{X}_k\}$ of independent, random, self-adjoint matrices with dimension d . Suppose that*

$$\mathbb{E} \mathbf{X}_k = \mathbf{0} \quad \text{and} \quad \lambda_{\max}(\mathbf{X}_k) \leq R \quad \text{almost surely.}$$

Compute the norm of the total variance:

$$\sigma^2 := \left\| \sum_k \mathbb{E} (\mathbf{X}_k^2) \right\|.$$

For all $t \geq 0$,

$$\mathbb{P} \left\{ \lambda_{\max} \left(\sum_k \mathbf{X}_k \right) \geq t \right\} \leq d \cdot \exp \left(\frac{-t^2/2}{\sigma^2 + Rt/3} \right).$$

The matrix Bernstein inequality, Theorem 1.3, follows from a more detailed result, which provides stronger Poisson-type decay for the tail. In §4, we derive these results from a martingale result.

1.5. **Miscellaneous Results.** The methods in this paper deliver a number of other results:

- All of the results described in the front matter follow from more general bounds for large deviations of matrix martingales. See §3–5 for the full story.
- All the inequalities we have mentioned, with exception of the matrix Chernoff bounds, have variants that hold for rectangular matrices. The extensions follow immediately from the self-adjoint case by applying an elegant device from operator theory, called the self-adjoint dilation of a matrix [Pau86]. See [Tro10c, §4.2] for additional details.

2. TAIL BOUNDS VIA MARTINGALE METHODS

This section contains the main part of the argument, which parallels Freedman’s argument for producing large deviation bounds for scalar martingales [Fre75]. The material here duplicates the note [Tro10a].

2.1. **Matrix Moments and Cumulants.** Consider a random s.a. matrix $\mathbf{X}$ that has moments of all orders. By analogy with the classical definitions for scalar random variables, we construct the matrix *moment generating function* (mgf) and *cumulant generating function* (cgf).

$$\mathbf{M}_{\mathbf{X}}(\theta) := \mathbb{E} e^{\theta \mathbf{X}} \quad \text{and} \quad \mathbf{\Xi}_{\mathbf{X}}(\theta) := \log \mathbb{E} e^{\theta \mathbf{X}} \quad \text{for } \theta \in \mathbb{R}. \quad (2.1)$$

The mgf has a formal power series expansion that displays the raw moments of the random matrix:

$$\mathbf{M}_{\mathbf{X}}(\theta) = \mathbf{I} + \sum_{j=1}^{\infty} \frac{\theta^j}{j!} \cdot \mathbb{E}(\mathbf{X}^j).$$

In the scalar setting, the cgf can be interpreted as an *exponential mean*, a weighted average of a random variable that emphasizes large (positive) deviations. The matrix cgf admits a similar intuition, and we treat it as a measure of the variability of a random matrix.

2.2. **The Large Deviation Supermartingale.** In this section, we extend Freedman’s martingale techniques [Fre75] to the matrix setting. The matrix cgf and Lieb’s result, Theorem A.1, play a central role in this development.

We begin with a filtration $\{\mathcal{F}_k : k = 0, 1, 2, \dots\}$ of a master probability space, and we write $\mathbb{E}_k$ for the conditional expectation with respect to $\mathcal{F}_k$. Consider an adapted random process $\{\mathbf{X}_k : k = 1, 2, 3, \dots\}$ and a previsible random process $\{\mathbf{V}_k : k = 1, 2, 3, \dots\}$ whose values are s.a. matrices with dimension d . Suppose that the two processes are related through a conditional cgf bound of the form

$$\log \mathbb{E}_{k-1} e^{\theta \mathbf{X}_k} \preceq g(\theta) \cdot \mathbf{V}_k \quad \text{almost surely for } \theta > 0. \quad (2.2)$$

The function $g : (0, \infty) \rightarrow [0, \infty]$, and—for simplicity—we do not allow this function to depend on the index k . It is convenient to define the partial sums of the original process and the partial sums of the conditional cgf bounds:

$$\begin{aligned} \mathbf{Y}_0 &:= \mathbf{0} \quad \text{and} \quad \mathbf{Y}_k := \sum_{j=1}^k \mathbf{X}_j. \\ \mathbf{W}_0 &:= \mathbf{0} \quad \text{and} \quad \mathbf{W}_k := \sum_{j=1}^k \mathbf{V}_j. \end{aligned}$$

In almost all our examples, $\{\mathbf{V}_k\}$ is a sequence of psd matrices, and so $\{\mathbf{W}_k\}$ increases with respect to the semidefinite order. The random matrix $\mathbf{W}_k$ can be viewed as a measure of the total variability of the process $\{\mathbf{X}_k\}$ up to time k .

To continue, we fix the function g and a positive number θ . Define a real-valued function with two s.a. matrix arguments:

$$G_{\theta}(\mathbf{Y}, \mathbf{W}) := \text{tr} \exp(\theta \mathbf{Y} - g(\theta) \cdot \mathbf{W}).$$

We use the function G_θ to construct a real-valued random process.

$$S_k := S_k(\theta) = G_\theta(\mathbf{Y}_k, \mathbf{W}_k) \quad \text{for } k = 0, 1, 2, \dots \quad (2.3)$$

This process is an evolving measure of the discrepancy between the partial sum process $\{\mathbf{Y}_k\}$ and the cumulant sum process $\{\mathbf{W}_k\}$. The following lemma describes the key properties of this random sequence. In particular, the average discrepancy decreases with time. The proof relies on Lieb's result, Theorem A.1.

Lemma 2.1. *For each fixed $\theta > 0$, the random process $\{S_k(\theta) : k = 0, 1, 2, \dots\}$ defined in (2.3) is a positive supermartingale whose initial value $S_0 = d$.*

Proof. It is easily seen that S_k is positive because the exponential of a self-adjoint matrix is pd, and the trace of a pd matrix is positive. We obtain the initial value from a short calculation:

$$S_0 = \text{tr} \exp(\theta \mathbf{Y}_0 - \mathbf{W}_0(\theta)) = \text{tr} \exp(\mathbf{0}_d) = \text{tr} \mathbf{I}_d = d.$$

To prove that the process is a supermartingale, we ascend a short chain of inequalities.

$$\begin{aligned} \mathbb{E}_{k-1} S_k &= \mathbb{E}_{k-1} \text{tr} \exp\left(\theta \mathbf{Y}_{k-1} - g(\theta) \cdot \mathbf{W}_k + \log e^{\theta \mathbf{X}_k}\right) \\ &\leq \text{tr} \exp\left(\theta \mathbf{Y}_{k-1} - g(\theta) \cdot \mathbf{W}_k + \log \mathbb{E}_{k-1} e^{\theta \mathbf{X}_k}\right) \\ &\leq \text{tr} \exp(\theta \mathbf{Y}_{k-1} - g(\theta) \cdot \mathbf{W}_k + g(\theta) \cdot \mathbf{V}_k) \\ &= \text{tr} \exp(\theta \mathbf{Y}_{k-1} - g(\theta) \cdot \mathbf{W}_{k-1}) \\ &= S_{k-1}. \end{aligned}$$

In the first step, we remove the term $\mathbf{X}_k$ from the partial sum $\mathbf{Y}_k$ and rewrite it using the definition (A.7) of the matrix logarithm. Next, we invoke Lieb's Theorem, conditional on $\mathcal{F}_{k-1}$, to verify the concavity of the function

$$\mathbf{A} \mapsto \text{tr} \exp(\theta \mathbf{Y}_{k-1} - g(\theta) \cdot \mathbf{W}_k + \log(\mathbf{A})).$$

We apply Jensen's inequality (A.9) to draw the conditional expectation inside the function. This act is legal because $\mathbf{Y}_{k-1}$ and $\mathbf{W}_k$ are both measurable with respect to $\mathcal{F}_{k-1}$. The second inequality depends on the assumption (2.2) together with the fact (A.6) that the trace of the matrix exponential is monotone. The final step recalls that $\{\mathbf{W}_k\}$ is the sequence of partial sums of $\{\mathbf{V}_k\}$. $\square$

Finally, we present a simple inequality for the function G_θ that holds when we have control on the eigenvalues of its arguments.

Lemma 2.2. *Suppose that $\lambda_{\max}(\mathbf{Y}) \geq y$ and that $\lambda_{\max}(\mathbf{W}) \leq w$. For each $\theta > 0$,*

$$G_\theta(\mathbf{Y}, \mathbf{W}) \geq e^{\theta y - g(\theta)w}.$$

Proof. Recall that $g(\theta) \geq 0$. The bound results from a straightforward calculation:

$$G_\theta(\mathbf{Y}, \mathbf{W}) = \text{tr} e^{\theta \mathbf{Y} - g(\theta) \cdot \mathbf{W}} \geq \text{tr} e^{\theta \mathbf{Y} - g(\theta)w \mathbf{I}} \geq \lambda_{\max}\left(e^{\theta \mathbf{Y} - g(\theta)w \mathbf{I}}\right) = e^{\theta \lambda_{\max}(\mathbf{Y}) - g(\theta)w} \geq e^{\theta y - g(\theta)w}.$$

The first inequality depends on the fact that $\mathbf{W} \preceq w \mathbf{I}$ and the monotonicity (A.6) of the trace exponential. The second inequality relies on the property (A.1) that the trace of a psd matrix is at least as large as its maximum eigenvalue. The third identity follows from the spectral mapping theorem and elementary properties of the maximum eigenvalue map. $\square$

2.3. The Main Result. Our key theorem provides a bound on the probability that the partial sum of a matrix-valued random process is large.

Theorem 2.3. *Consider an adapted sequence $\{\mathbf{X}_k\}$ and a previsible sequence $\{\mathbf{V}_k\}$ of self-adjoint matrices with dimension d . Assume these sequences satisfy the relations*

$$\log \mathbb{E}_{k-1} e^{\theta \mathbf{X}_k} \preceq g(\theta) \cdot \mathbf{V}_k \quad \text{almost surely for each } \theta > 0,$$

where $g : (0, \infty) \rightarrow [0, \infty]$. Define the partial sums

$$\mathbf{Y}_k := \sum_{j=1}^k \mathbf{X}_j \quad \text{and} \quad \mathbf{W}_k := \sum_{j=1}^k \mathbf{V}_j.$$

For all $t, w \in \mathbb{R}$,

$$\mathbb{P} \{ \exists k : \lambda_{\max}(\mathbf{Y}_k) \geq t \quad \text{and} \quad \lambda_{\max}(\mathbf{W}_k) \leq w \} \leq d \cdot \inf_{\theta > 0} e^{-\theta t + g(\theta)w}.$$

In particular, the cumulant bound holds when

$$\mathbb{E}_{k-1} e^{\theta \mathbf{X}_k} \preceq e^{g(\theta) \cdot \mathbf{V}_k} \quad \text{almost surely for each } \theta > 0.$$

Proof. First, note that the cgf hypothesis holds when

$$\mathbb{E}_{k-1} e^{\theta \mathbf{X}_k} \preceq e^{g(\theta) \cdot \mathbf{V}_k}$$

because of the operator monotonicity (A.8) of the logarithm.

The strategy for the main argument is identical with the stopping-time technique used by Freedman [Fre75]. Fix a positive parameter θ , which we will optimize later. Following the discussion in Section 2.2, we introduce the random process $S_k = G_\theta(\mathbf{Y}_k, \mathbf{W}_k)$. Lemma 2.1 implies that $\{S_k\}$ is a positive supermartingale with initial value d . Let us emphasize that these simple properties of the auxiliary random process distill all the essential information from the hypotheses of the theorem.

Define a stopping time κ by finding the first time instant k when the maximum eigenvalue of the partial sum process $\{\mathbf{Y}_k\}$ reaches the level t even though the sum of cumulant bounds has maximum eigenvalue no larger than w .

$$\kappa := \inf \{ k \geq 0 : \lambda_{\max}(\mathbf{Y}_k) \geq t \quad \text{and} \quad \lambda_{\max}(\mathbf{W}_k) \leq w \}.$$

When the infimum is empty, the stopping time $\kappa = \infty$. Consider a system of exceptional events:

$$E_k := \{ \lambda_{\max}(\mathbf{Y}_k) \geq t \quad \text{and} \quad \lambda_{\max}(\mathbf{W}_k) \leq w \} \quad \text{for } k = 0, 1, 2, \dots$$

Construct the event $E := \bigcup_{k=0}^{\infty} E_k$ that one or more of these exceptional situations takes place. The intuition behind this definition is that the partial sum $\mathbf{Y}_k$ is typically not large unless the process $\{\mathbf{X}_k\}$ has varied substantially, a situation that the bound on $\mathbf{W}_k$ disallows. As a result, the event E is rather unlikely.

We are prepared to estimate the probability of the exceptional event. First, note that $\kappa < \infty$ on the event E . Therefore, Lemma 2.2 provides a conditional lower bound for the process $\{S_k\}$ at the stopping time κ :

$$S_\kappa = G_\theta(\mathbf{Y}_\kappa, \mathbf{W}_\kappa) \geq e^{\theta t - g(\theta)w} \quad \text{on the event } E.$$

Since $\mathbb{E} S_k \leq d$ for each (finite) index k ,

$$\begin{aligned} d &\geq \sum_{k=1}^{\infty} \mathbb{E}[S_k | \kappa = k] \cdot \mathbb{P} \{ \kappa = k \} = \mathbb{E}[S_\kappa | \kappa < \infty] \geq \int_{\{\kappa < \infty\}} S_\kappa d\mathbb{P} \\ &\geq \int_E S_\kappa d\mathbb{P} \geq \mathbb{P}(E) \cdot \inf_E S_\kappa \geq \mathbb{P}(E) \cdot e^{\theta t - g(\theta)w}. \end{aligned}$$

We require the fact that S_κ is positive to justify these inequalities. Rearrange the relation to obtain

$$\mathbb{P}(E) \leq d \cdot e^{-\theta t + g(\theta)w}.$$

Minimize the right-hand side with respect to θ to complete the main part of the argument. $\square$

We often prefer to use a corollary of Theorem 2.3 that describes the sum of a finite process. This focus allows us to avoid distracting details about the convergence of infinite series.

Corollary 2.4. *Suppose the hypotheses of Theorem 2.3 are in force, and suppose the random processes are finite in length. Define*

$$\mathbf{Y} := \sum_k \mathbf{X}_k \quad \text{and} \quad \mathbf{W} := \sum_k \mathbf{V}_k.$$

For all $t, w \in \mathbb{R}$,

$$\mathbb{P} \{ \lambda_{\max}(\mathbf{Y}) \geq t \quad \text{and} \quad \lambda_{\max}(\mathbf{W}) \leq w \} \leq d \cdot \inf_{\theta > 0} e^{-\theta t + g(\theta)w}.$$

3. SUMS OF RANDOM SEMIDEFINITE MATRICES

In this section, we establish Chernoff inequalities for the sum of an adapted sequence of random psd matrices. This result extends the Chernoff bounds for independent random matrices, Theorem 1.2, that we presented in §1.3.

Theorem 3.1 (Matrix Chernoff: Adapted Sequences). *Consider a finite adapted sequence $\{\mathbf{X}_k\}$ of positive-semidefinite matrices with dimension d , and suppose that*

$$\lambda_{\max}(\mathbf{X}_k) \leq R \quad \text{almost surely.}$$

Define the finite series

$$\mathbf{Y} := \sum_k \mathbf{X}_k \quad \text{and} \quad \mathbf{W} := \sum_k \mathbb{E}_{k-1} \mathbf{X}_k.$$

For all $\mu \geq 0$,

$$\begin{aligned} \mathbb{P} \{ \lambda_{\min}(\mathbf{Y}) \leq (1 - \delta)\mu \quad \text{and} \quad \lambda_{\min}(\mathbf{W}) \geq \mu \} &\leq d \cdot \left[\frac{e^{-\delta}}{(1 - \delta)^{1-\delta}} \right]^{\mu/R} \quad \text{for } \delta \in [0, 1), \text{ and} \\ \mathbb{P} \{ \lambda_{\max}(\mathbf{Y}) \geq (1 + \delta)\mu \quad \text{and} \quad \lambda_{\max}(\mathbf{W}) \leq \mu \} &\leq d \cdot \left[\frac{e^{\delta}}{(1 + \delta)^{1+\delta}} \right]^{\mu/R} \quad \text{for } \delta \geq 0. \end{aligned}$$

The Chernoff bound for independent random matrices, Theorem 1.2, follows as an immediate corollary.

Proof of Theorem 1.2 from Theorem 3.1. In this case, we assume that $\{\mathbf{X}_k\}$ is an independent sequence of psd matrices. Then the matrix $\mathbf{W}$ is not random, so we can define the numbers

$$\mu_{\min} := \lambda_{\min}(\mathbf{W}) \quad \text{and} \quad \mu_{\max} := \lambda_{\max}(\mathbf{W}).$$

As a consequence, we can replace μ with $\mu_{\min}$ or $\mu_{\max}$, as appropriate, and remove the part of the event involving $\mathbf{W}$ from both probabilities in Theorem 3.1. $\square$

3.1. Proofs. We begin with a semidefinite bound for the mgf of a random psd matrix. This argument transfers a linear upper bound for the scalar exponential to the matrix case.

Lemma 3.2 (Chernoff mgf). *Suppose that $\mathbf{X}$ is a random psd matrix that satisfies $\lambda_{\max}(\mathbf{X}) \leq 1$. Then*

$$\mathbb{E} e^{\theta \mathbf{X}} \preceq \exp \left((e^{\theta} - 1)(\mathbb{E} \mathbf{X}) \right) \quad \text{for } \theta \in \mathbb{R}.$$

Proof. Consider the function $f(x) = e^{\theta x}$. Since f is convex, its graph lies below the chord connecting two points. In particular,

$$f(x) \leq f(0) + [f(1) - f(0)] \cdot x \quad \text{for } x \in [0, 1].$$

More explicitly,

$$e^{\theta x} \leq 1 + (e^{\theta} - 1) \cdot x \quad \text{for } x \in [0, 1].$$

Since the eigenvalues of $\mathbf{X}$ lie in the interval $[0, 1]$, the transfer rule (A.3) implies that

$$\mathbf{e}^{\theta \mathbf{X}} \preceq \mathbf{I} + (\mathbf{e}^\theta - 1)\mathbf{X}.$$

Expectation respects the semidefinite order, so

$$\mathbb{E} \mathbf{e}^{\theta \mathbf{X}} \preceq \mathbf{I} + (\mathbf{e}^\theta - 1)(\mathbb{E} \mathbf{X}) \preceq \exp \left((\mathbf{e}^\theta - 1)(\mathbb{E} \mathbf{X}) \right),$$

where the second relation is (A.4). $\square$

We prove the upper Chernoff bound first, since the argument is slightly easier.

Proof of Theorem 3.1, Upper Bound. By homogeneity, we may assume that $\lambda_{\max}(\mathbf{X}_k) \leq 1$; the general case follows by re-scaling. An application of Lemma 3.2 demonstrates that

$$\mathbb{E}_{k-1} \mathbf{e}^{\theta \mathbf{X}_k} \preceq \mathbf{e}^{g(\theta) \cdot \mathbb{E}_{k-1} \mathbf{X}_k} \quad \text{where } g(\theta) = \mathbf{e}^\theta - 1 \text{ for } \theta > 0.$$

Corollary 2.4 provides that

$$\mathbb{P} \left\{ \lambda_{\max} \left(\sum_k \mathbf{X}_k \right) \geq (1 + \delta)\mu \quad \text{and} \quad \lambda_{\max} \left(\sum_k \mathbb{E}_{k-1} \mathbf{X}_k \right) \leq \mu \right\} \leq d \cdot \inf_{\theta > 0} \mathbf{e}^{-\theta(1+\delta)\mu + g(\theta)\mu}.$$

The infimum is achieved when $\theta = \log(1 + \delta)$. Substitute and simplify to complete the proof. $\square$

The lower Chernoff bound follows from a similar argument.

Proof of Theorem 3.1, Lower Bound. As before, we may assume that $\lambda_{\max}(\mathbf{X}_k) \leq 1$. This time, we intend to apply Corollary 2.4 to the sequence $\{-\mathbf{X}_k\}$. Lemma 3.2 demonstrates that

$$\mathbb{E}_{k-1} \mathbf{e}^{(-\theta) \mathbf{X}_k} \preceq \mathbf{e}^{g(\theta) \cdot \mathbb{E}_{k-1} (-\mathbf{X}_k)} \quad \text{where } g(\theta) = 1 - \mathbf{e}^{-\theta} \text{ for } \theta > 0.$$

Corollary 2.4 delivers

$$\mathbb{P} \left\{ \lambda_{\max} \left(-\sum_k \mathbf{X}_k \right) \geq -(1 - \delta)\mu \quad \text{and} \quad \lambda_{\max} \left(-\sum_k \mathbb{E}_{k-1} \mathbf{X}_k \right) \leq -\mu \right\} \leq d \cdot \inf_{\theta > 0} \mathbf{e}^{(\theta(1-\delta) - g(\theta))\mu}.$$

Since $\lambda_{\max}(-\mathbf{A}) = -\lambda_{\min}(\mathbf{A})$ for each s.a. matrix $\mathbf{A}$, we can draw the negation out of the eigenvalue maps and reverse the sense of the inequalities inside the probability. Finally, we observe that the infimum occurs when $\theta = -\log(1 - \delta)$. $\square$

4. INCORPORATING VARIANCE INFORMATION

In this section, we establish a variant of the Freedman inequality for martingales [Fre75, Thm. (1.6)]. This inequality demonstrates that a sum of random matrices has normal concentration around its mean and Poisson-type decay in the tails.

Theorem 4.1 (Matrix Bennett: Adapted Sequences). *Consider a finite adapted sequence $\{\mathbf{X}_k\}$ of self-adjoint matrices with dimension d that satisfy the relations*

$$\mathbb{E}_{k-1} \mathbf{X}_k = \mathbf{0} \quad \text{and} \quad \lambda_{\max}(\mathbf{X}_k) \leq R \quad \text{almost surely.}$$

Define the finite series

$$\mathbf{Y} := \sum_k \mathbf{X}_k \quad \text{and} \quad \mathbf{W} := \sum_k \mathbb{E}_{k-1} (\mathbf{X}_k^2).$$

For all $t \geq 0$ and $\sigma^2 > 0$,

$$\mathbb{P} \left\{ \lambda_{\max}(\mathbf{Y}) \geq t \quad \text{and} \quad \lambda_{\max}(\mathbf{W}) \leq \sigma^2 \right\} \leq d \cdot \exp \left\{ -\frac{\sigma^2}{R^2} \cdot h \left(\frac{Rt}{\sigma^2} \right) \right\}.$$

The function $h(u) := (1 + u) \log(1 + u) - u$ for $u \geq 0$.

We obtain a Freedman-type inequality for matrix martingales when we simplify the right-hand side of the probability bound in Theorem 4.1.

Corollary 4.2 (Matrix Freedman). *Under the hypotheses of Theorem 4.1,*

$$\mathbb{P} \left\{ \lambda_{\max}(\mathbf{Y}) \geq t \quad \text{and} \quad \lambda_{\max}(\mathbf{W}) \leq \sigma^2 \right\} \leq d \cdot \exp \left(\frac{-t^2/2}{\sigma^2 + Rt/3} \right).$$

Proof. This corollary is a direct consequence of Theorem 4.1 and the numerical inequality

$$h(u) = (1+u) \log(1+u) - u \geq \frac{u^2/2}{1+u/3} \quad \text{for } u \geq 0,$$

which can be obtained by comparing derivatives. $\square$

The Bernstein inequality, Theorem 1.3, for sums of independent random matrices follows directly from the Freedman inequality, Corollary 4.2.

Proof of Theorem 1.3 from Corollary 4.2. Indeed, when $\{\mathbf{X}_k\}$ is an independent family of random matrices, the matrix $\mathbf{W}$ is deterministic. Therefore, if the bound $\sigma^2 \geq \|\mathbf{W}\|$ holds, then it holds almost surely. As a result, we can remove the condition on $\mathbf{W}$ from the probability bound in the theorem. We can derive a matrix Bennett inequality from Theorem 4.1 in precisely the same manner. $\square$

The proof of Theorem 4.1 appears below. Remark 4.4 shows that we can obtain the same results if we are provided with a set of bounds on the moments of the summands.

4.1. Proofs. The first lemma shows how to bound the mgf of a zero-mean random matrix using an almost-sure bound for its largest eigenvalue. We learned this argument from Yao-Liang Yu.

Lemma 4.3 (Bennett mgf). *Suppose that $\mathbf{X}$ is a random s.a. matrix that satisfies*

$$\mathbb{E} \mathbf{X} = \mathbf{0} \quad \text{and} \quad \lambda_{\max}(\mathbf{X}) \leq 1 \quad \text{almost surely.}$$

Then

$$\mathbb{E} e^{\theta \mathbf{X}} \preceq \exp \left((e^\theta - \theta - 1) \cdot \mathbb{E}(\mathbf{X}^2) \right) \quad \text{for } \theta > 0.$$

Proof. Fix the parameter $\theta > 0$, and define a continuous function f on the real line:

$$f(x) = \frac{e^{\theta x} - \theta x - 1}{x^2} \quad \text{for } x \neq 0 \quad \text{and} \quad f(0) = \frac{\theta^2}{2}.$$

An exercise in differential calculus verifies that f is nonnegative and increasing. The matrix $\mathbf{X}$ has a (random) eigenvalue decomposition $\mathbf{X} = \mathbf{Q} \mathbf{\Lambda} \mathbf{Q}^*$ where $\mathbf{\Lambda} \preceq \mathbf{I}$ almost surely. We see that

$$f(\mathbf{X}) = \mathbf{Q} f(\mathbf{\Lambda}) \mathbf{Q}^* \preceq \mathbf{Q} \cdot f(\mathbf{I}) \cdot \mathbf{Q}^* = f(1) \cdot \mathbf{I}.$$

Expanding the matrix exponential and invoking the conjugation rule (A.2), we discover that

$$e^{\theta \mathbf{X}} = \mathbf{I} + \theta \mathbf{X} + \mathbf{X} f(\mathbf{X}) \mathbf{X} \preceq \mathbf{I} + \theta \mathbf{X} + f(1) \cdot \mathbf{X}^2.$$

To complete the proof, we take the expectation of this semidefinite relation.

$$\mathbb{E} e^{\theta \mathbf{X}} \preceq \mathbf{I} + f(1) \cdot \mathbb{E}(\mathbf{X}^2) \preceq \exp(f(1) \cdot \mathbb{E}(\mathbf{X}^2))$$

The final step follows from (A.4). $\square$

We are ready to establish the Bennett inequality for adapted sequences of random matrices.

Proof of Theorem 4.1. We assume that $R = 1$; the general result follows by re-scaling since $\mathbf{Y}$ is 1-homogeneous and $\mathbf{W}$ is 2-homogeneous. Invoke Lemma 4.3 to see that

$$\mathbb{E}_{k-1} e^{\theta \mathbf{X}_k} \preceq \exp(g(\theta) \cdot \mathbb{E}_{k-1}(\mathbf{X}_k^2)) \quad \text{where } g(\theta) = e^\theta - \theta - 1.$$

Corollary 2.4 implies that

$$\mathbb{P} \left\{ \lambda_{\max} \left(\sum_k \mathbf{X}_k \right) \geq t \quad \text{and} \quad \lambda_{\max} \left(\sum_k \mathbb{E}_{k-1}(\mathbf{X}_k^2) \right) \leq \sigma^2 \right\} \leq d \cdot \inf_{\theta > 0} e^{\theta t - g(\theta) \sigma^2}.$$

The infimum is achieved when $\theta = \log(1 + t/\sigma^2)$. $\square$

Remark 4.4. We can also establish the Bennett mgf bound under appropriate hypotheses on the growth of the moments of $\mathbf{X}$. This argument proceeds by estimating each term in the Taylor series of the matrix exponential.

Suppose that $\mathbf{X}$ is a random s.a. matrix with $\mathbb{E} \mathbf{X} = \mathbf{0}$, and assume the moment growth bounds

$$\mathbb{E}(\mathbf{X}^j) \preceq R^{j-2} \cdot \mathbf{A}^2 \quad \text{for } j = 2, 3, 4, \dots$$

We demonstrate that

$$\mathbb{E} e^{\theta \mathbf{X}} \preceq \exp \left(\frac{e^{\theta R} - \theta R - 1}{R^2} \cdot \mathbf{A}^2 \right) \quad \text{for } \theta > 0.$$

Indeed, the growth condition for the moments yields the bound

$$\begin{aligned} \mathbb{E} e^{\theta \mathbf{X}} &= \mathbf{I} + \theta \cdot \mathbb{E} \mathbf{X} + \sum_{j=2}^{\infty} \frac{\theta^j \mathbb{E}(\mathbf{X}^j)}{j!} \preceq \mathbf{I} + \frac{1}{R^2} \sum_{j=2}^{\infty} \frac{(\theta R)^j}{j!} \cdot \mathbf{A}^2 \\ &= \mathbf{I} + \frac{e^{\theta R} - \theta R - 1}{R^2} \cdot \mathbf{A}^2 \preceq \exp \left(\frac{e^{\theta R} - \theta R - 1}{R^2} \cdot \mathbf{A}^2 \right). \end{aligned}$$

As usual, the last relation follows from (A.4).

5. RADEMACHER AND GAUSSIAN SERIES

This section establishes normal concentration for Rademacher and Gaussian series with matrix coefficients. The first step is to verify the bounds for the mgf of a fixed matrix modulated by a Rademacher variable or a Gaussian variable; see also [Oli10b, Lem. 2].

Lemma 5.1 (Rademacher and Gaussian mgfs). *Suppose that $\mathbf{A}$ is an s.a. matrix. Let ε be a Rademacher random variable, and let γ be a standard normal random variable. Then*

$$\mathbb{E} e^{\varepsilon \mathbf{A}} \preceq e^{\theta^2 \mathbf{A}^2/2} \quad \text{and} \quad \mathbb{E} e^{\gamma \mathbf{A}} = e^{\theta^2 \mathbf{A}^2/2} \quad \text{for } \theta \in \mathbb{R}.$$

Proof. By absorbing θ into $\mathbf{A}$, we may assume $\theta = 1$ in each case. We begin with the Rademacher mgf. By direct calculation,

$$\mathbb{E} e^{\varepsilon \mathbf{A}} = \cosh(\mathbf{A}) \preceq e^{\mathbf{A}^2/2},$$

where the second relation is (A.5).

Recall that the moments of a standard normal variable are

$$\mathbb{E}(\gamma^{2j+1}) = 0 \quad \text{and} \quad \mathbb{E}(\gamma^{2j}) = \frac{(2j)!}{j! 2^j} \quad \text{for } j = 0, 1, 2, \dots$$

Therefore,

$$\mathbb{E} e^{\gamma \mathbf{A}} = \mathbf{I} + \sum_{j=1}^{\infty} \frac{\mathbb{E}(\gamma^{2j}) \mathbf{A}^{2j}}{(2j)!} = \mathbf{I} + \sum_{j=1}^{\infty} \frac{(\mathbf{A}^2/2)^j}{j!} = e^{\mathbf{A}^2/2}.$$

The first identity holds because the odd terms in the series vanish. $\square$

We immediately obtain the bound for Rademacher and Gaussian series.

Proof of Theorem 1.1. Let $\{\xi_k\}$ be a finite sequence of independent Rademacher variables or independent standard normal variables. Invoke Lemma 5.1 to obtain

$$\mathbb{E} e^{\xi_k \theta \mathbf{A}_k} \preceq e^{\theta^2 \mathbf{A}_k^2/2}.$$

By assumption, $\lambda_{\max}(\sum_k \mathbf{A}_k^2) \leq \sigma^2$ almost surely. Therefore, Corollary 2.3 yields

$$\mathbb{P} \left\{ \lambda_{\max} \left(\sum_k \xi_k \mathbf{A}_k \right) \geq t \right\} \leq d \cdot \inf_{\theta > 0} e^{-\theta t + \theta^2 \sigma^2/2}.$$

The infimum is attained at $\theta = t/\sigma^2$. $\square$

APPENDIX A. MATHEMATICAL BACKGROUND

This section provides a short introduction to the background material we use in our proofs. Section A.1 discusses matrix theory, and Section A.2 reviews some relevant ideas from probability.

A.1. Matrix Theory. Most of these results can be located in Bhatia's books on matrix analysis [Bha97, Bha07]. The works of Horn and Johnson [HJ85, HJ94] also serve as good general references. Higham's book [Hig08] is an excellent source for information about matrix functions.

A.1.1. Conventions. A *matrix* is a finite, two-dimensional array of complex numbers. *In this paper, all matrices are square unless otherwise noted.* We add the qualification *rectangular* when we need to refer to a general array, which may be square or nonsquare. Many parts of the discussion do not depend on the size of a matrix, so we specify dimensions only when it matters. In particular, we usually do not state the size of a matrix when it is determined by the context.

A.1.2. Basic Matrices. We write $\mathbf{0}$ for the zero matrix and $\mathbf{I}$ for the identity matrix. Occasionally, we add a subscript to specify the dimension, e.g., $\mathbf{I}_d$ is the $d \times d$ identity.

A matrix that satisfies $\mathbf{Q}\mathbf{Q}^* = \mathbf{I} = \mathbf{Q}^*\mathbf{Q}$ is called *unitary*. We reserve the symbol $\mathbf{Q}$ for a unitary matrix. The symbol $*$ denotes the conjugate transpose.

A square matrix that satisfies $\mathbf{A} = \mathbf{A}^*$ is called *self-adjoint* (briefly, *s.a.*). We adopt Parlett's convention that letters symmetric around the vertical axis ($\mathbf{A}, \mathbf{H}, \dots, \mathbf{Y}$) represent s.a. matrices unless otherwise noted.

A.1.3. The Semidefinite Order. An s.a. matrix $\mathbf{A}$ with nonnegative eigenvalues is called *positive semidefinite* (briefly, *psd*). When the eigenvalues are strictly positive, we say the matrix is *positive definite* (briefly, *pd*). An easy consequence of the definition is that

$$\lambda_{\max}(\mathbf{A}) \leq \text{tr } \mathbf{A} \quad \text{when } \mathbf{A} \text{ is psd} \quad (\text{A.1})$$

because the trace is the sum of the eigenvalues.

The set of all psd matrices with fixed dimension forms a closed, convex cone. Therefore, we may define the *semidefinite partial order* on s.a. matrices of the same size by the rule

$$\mathbf{A} \preceq \mathbf{H} \iff \mathbf{H} - \mathbf{A} \text{ is psd.}$$

In particular, we may write $\mathbf{A} \succeq \mathbf{0}$ to indicate that $\mathbf{A}$ is psd and $\mathbf{A} \succ \mathbf{0}$ to indicate that $\mathbf{A}$ is pd. For a diagonal matrix, $\mathbf{A} \succeq \mathbf{0}$ means that each entry of $\mathbf{A}$ is nonnegative.

The semidefinite order is preserved by conjugation:

$$\mathbf{A} \preceq \mathbf{H} \implies \mathbf{B}^* \mathbf{A} \mathbf{B} \preceq \mathbf{B}^* \mathbf{H} \mathbf{B} \quad \text{for each matrix } \mathbf{B}. \quad (\text{A.2})$$

We refer to (A.2) as the *conjugation rule*.

A.1.4. Matrix Functions. Let us describe the most direct method for lifting functions on the reals to functions on s.a. matrices. Consider a function $f : \mathbb{R} \rightarrow \mathbb{R}$. First, extend f to a map on diagonal matrices by applying the function to each diagonal entry:

$$(f(\mathbf{A}))_{jj} := f(\mathbf{A}_{jj}) \quad \text{for each index } j.$$

We extend f to all s.a. matrices by way of the eigenvalue decomposition. If $\mathbf{A} = \mathbf{Q}\mathbf{\Lambda}\mathbf{Q}^*$, then

$$f(\mathbf{A}) = f(\mathbf{Q}\mathbf{\Lambda}\mathbf{Q}^*) := \mathbf{Q}f(\mathbf{\Lambda})\mathbf{Q}^*.$$

The *spectral mapping theorem* states that each eigenvalue of $f(\mathbf{A})$ has the form $f(\lambda)$, where λ is an eigenvalue of $\mathbf{A}$. This point is obvious from our definition.

Inequalities for real functions extend to semidefinite relationships for matrix functions:

$$f(a) \leq g(a) \quad \text{for } a \in I \implies f(\mathbf{A}) \preceq g(\mathbf{A}) \quad \text{when the eigenvalues of } \mathbf{A} \text{ lie in } I. \quad (\text{A.3})$$

Indeed, let us decompose $\mathbf{A} = \mathbf{Q}\mathbf{\Lambda}\mathbf{Q}^*$. It is immediate that $f(\mathbf{\Lambda}) \preceq g(\mathbf{\Lambda})$. Conjugate by $\mathbf{Q}$, as justified by (A.2), and invoke the definition of a matrix function. We sometimes refer to (A.3) as the *transfer rule*.

When a real function has a convergent power series expansion, we can also define an s.a. matrix function via the same power series expansion:

$$f(a) = c_0 + \sum_{j=1}^{\infty} c_j a^j \implies f(\mathbf{A}) := c_0 \mathbf{I} + \sum_{j=1}^{\infty} c_j \mathbf{A}^j.$$

In this case, the two definitions of a matrix function coincide.

Beware: One must never take for granted that a standard property of a real function generalizes to the associated matrix function.

A.1.5. The Matrix Exponential. We may define the matrix exponential of an s.a. matrix $\mathbf{A}$ via the power series

$$\exp(\mathbf{A}) := e^{\mathbf{A}} = \mathbf{I} + \sum_{j=1}^{\infty} \frac{\mathbf{A}^j}{j!}.$$

The exponential of an s.a. matrix is always pd because of the spectral mapping theorem.

On account of the transfer rule (A.3), the matrix exponential satisfies some simple semidefinite relations that we collect here. Since $1 + a \leq e^a$ for real a , we have

$$\mathbf{I} + \mathbf{A} \preceq e^{\mathbf{A}} \quad \text{for each s.a. matrix } \mathbf{A}. \quad (\text{A.4})$$

By comparing Taylor series, one verifies that $\cosh(a) \leq e^{a^2/2}$ for real a . Therefore,

$$\cosh(\mathbf{A}) \preceq e^{\mathbf{A}^2/2} \quad \text{for each s.a. matrix } \mathbf{A}. \quad (\text{A.5})$$

We often work with the trace of the matrix exponential

$$\text{tr exp} : \mathbf{A} \longmapsto \text{tr } e^{\mathbf{A}}.$$

The trace exponential is monotone with respect to the semidefinite order:

$$\mathbf{A} \preceq \mathbf{H} \implies \text{tr } e^{\mathbf{A}} \leq \text{tr } e^{\mathbf{H}}. \quad (\text{A.6})$$

See [Pet94, Sec. 2] for a short proof of this fact.

A.1.6. The Matrix Logarithm. The matrix logarithm is defined as the functional inverse of the matrix exponential:

$$\log(e^{\mathbf{A}}) := \mathbf{A} \quad \text{for each s.a. matrix } \mathbf{A}. \quad (\text{A.7})$$

This formula determines the logarithm on the pd cone, which is adequate for our purposes. The matrix logarithm is monotone with respect to the semidefinite order.

$$\mathbf{0} \prec \mathbf{A} \preceq \mathbf{H} \implies \log(\mathbf{A}) \preceq \log(\mathbf{H}). \quad (\text{A.8})$$

A.1.7. A Theorem of Lieb. The central tool in this paper is a deep theorem of Lieb from his seminal 1973 work on convex trace functions [Lie73, Thm. 6]. Epstein provides an alternative proof of this bound in [Eps73, Sec. II], and Ruskai offers a simplified account of Epstein's argument in [Rus02, Rus05]. For another approach that is based on the joint convexity of quantum relative entropy [Lin74, Lem. 2], see the recent note [Tro10b].

Theorem A.1 (Lieb). *Fix a self-adjoint matrix $\mathbf{H}$. The function*

$$\mathbf{A} \longmapsto \text{tr exp}(\mathbf{H} + \log(\mathbf{A}))$$

is concave on the positive-definite cone.

A.2. Probability. We continue with some material from probability, focusing on connections with matrices. Rogers and Williams [RW00] is our main source for information about martingales.

A.2.1. Conventions. We prefer to avoid abstraction and unnecessary technical detail, so we frame the standing assumption that all random variables are sufficiently regular that we are justified in computing expectations, interchanging limits, and so forth. Furthermore, we often state that a random variable satisfies some relation and omit the qualification “almost surely.” We reserve the letters $\mathbf{V}, \mathbf{W}, \mathbf{X}, \mathbf{Y}$ for random s.a. matrices.

A.2.2. Adapted Sequences. Let $(\Omega, \mathcal{F}, \mathbb{P})$ be a master probability space. Consider a filtration $\{\mathcal{F}_k\}$ contained in the master sigma algebra:

$$\mathcal{F}_0 \subset \mathcal{F}_1 \subset \mathcal{F}_2 \subset \cdots \subset \mathcal{F}_\infty \subset \mathcal{F}.$$

Given such a filtration, we define the conditional expectation

$$\mathbb{E}_k[\cdot] := \mathbb{E}[\cdot \mid \mathcal{F}_k].$$

We say that a sequence $\{\mathbf{X}_k\}$ of random matrices is *adapted* to the filtration when each $\mathbf{X}_k$ is measurable with respect to $\mathcal{F}_k$. Loosely speaking, an adapted sequence is one where the present depends only upon the past.

We say that a sequence $\{\mathbf{V}_k\}$ of random matrices is *previsible* when each $\mathbf{V}_k$ is measurable with respect to $\mathcal{F}_{k-1}$. In particular, the sequence $\{\mathbb{E}_{k-1} \mathbf{X}_k\}$ of conditional expectations of an adapted sequence $\{\mathbf{X}_k\}$ is previsible.

A *stopping time* is a random variable $\kappa : \Omega \rightarrow \{0, 1, 2, \dots, \infty\}$ that satisfies

$$\{\kappa \leq k\} \subset \mathcal{F}_k \quad \text{for } k = 0, 1, 2, \dots, \infty.$$

In words, we can determine if the stopping time has arrived from past experience.

A.2.3. Matrix Martingales. We say that an adapted sequence $\{\mathbf{Y}_k : k = 0, 1, 2, \dots\}$ of s.a. matrices is a *matrix martingale* when

$$\mathbb{E}_{k-1} \mathbf{Y}_k = \mathbf{Y}_{k-1} \quad \text{for } k = 1, 2, 3, \dots$$

We also impose an L_1 boundedness criterion:

$$\mathbb{E} \|\mathbf{Y}_k\| < \infty \quad \text{for } k = 1, 2, 3, \dots$$

Since all norms on a finite-dimensional space are equivalent, this condition is the same as the requirement that each coordinate of each matrix $\mathbf{Y}_k$ is integrable. It follows that we obtain a scalar martingale if we track any fixed coordinate of the sequence $\{\mathbf{Y}_k\}$.

Given a matrix martingale $\{\mathbf{Y}_k\}$, we construct the *difference sequence*

$$\mathbf{X}_k := \mathbf{Y}_k - \mathbf{Y}_{k-1} \quad \text{for } k = 1, 2, 3, \dots$$

Observe that the difference sequence is conditionally zero mean: $\mathbb{E}_{k-1} \mathbf{X}_k = \mathbf{0}$. Alternatively, we may begin with an adapted sequence $\{\mathbf{X}_k\}$ of conditionally zero-mean random matrices and then form the partial sum process

$$\mathbf{Y}_0 := \mathbf{0} \quad \text{and} \quad \mathbf{Y}_k := \sum_{j=1}^k \mathbf{X}_j.$$

It is easy to verify that $\{\mathbf{Y}_k\}$ is a martingale, provided that the integrability requirement holds.

A.2.4. Inequalities for Expectation. Jensen’s inequality describes how averaging interacts with convexity. Let $\mathbf{Z}$ be a random matrix, and let f be a real-valued function on matrices. Then

$$\mathbb{E} f(\mathbf{Z}) \leq f(\mathbb{E} \mathbf{Z}) \quad \text{when } f \text{ is concave.} \tag{A.9}$$

Since the expectation of a random matrix can be viewed as a convex combination and the psd cone is convex, expectation preserves the semidefinite order:

$$\mathbf{X} \preceq \mathbf{Y} \quad \text{almost surely} \implies \mathbb{E} \mathbf{X} \preceq \mathbb{E} \mathbf{Y}.$$

Finally, let us emphasize that each of these bounds holds when we replace the expectation $\mathbb{E}$ by the conditional expectation $\mathbb{E}_k$.

ACKNOWLEDGMENTS

I would like to thank Vern Paulsen and Bernhard Bodmann for some helpful conversations connected with this project. Klas Markström and David Gross provided some references to related work. Roberto Oliveira introduced me to Freedman's inequality and encouraged me to apply the methods in the paper [Tro10c] to this problem. It was Oliveira's elegant work [Oli10b] on matrix probability inequalities that spurred me to pursue this project in the first place. Finally, I would like to thank Yao-Liang Yu, who pointed out an inconsistency in the proof of Theorem 2.3 and who proposed the argument in Lemma 4.3. Richard Chen and Alex Gittens have also helped me root out typographic errors.

REFERENCES

- [Bha97] R. Bhatia. *Matrix Analysis*. Number 169 in Graduate Texts in Mathematics. Springer, Berlin, 1997.
- [Bha07] R. Bhatia. *Positive Definite Matrices*. Princeton Univ. Press, Princeton, NJ, 2007.
- [Eps73] H. Epstein. Remarks on two theorems of E. Lieb. *Comm. Math. Phys.*, 31:317–325, 1973.
- [Fre75] D. A. Freedman. On tail probabilities for martingales. *Ann. Probab.*, 3(1):100–118, Feb. 1975.
- [Hig08] N. J. Higham. *Functions of Matrices: Theory and Computation*. Society for Industrial and Applied Mathematics, Philadelphia, PA, 2008.
- [HJ85] R. A. Horn and C. R. Johnson. *Matrix Analysis*. Cambridge Univ. Press, Cambridge, 1985.
- [HJ94] R. A. Horn and C. R. Johnson. *Topics in Matrix Analysis*. Cambridge Univ. Press, Cambridge, 1994.
- [Lie73] E. H. Lieb. Convex trace functions and the Wigner–Yanase–Dyson conjecture. *Adv. Math.*, 11:267–288, 1973.
- [Lin74] G. Lindblad. Expectations and entropy inequalities for finite quantum systems. *Comm. Math. Phys.*, 39:111–119, 1974.
- [Oli10a] R. I. Oliveira. Concentration of the adjacency matrix and of the Laplacian in random graphs with independent edges. Available at [arXiv:0911.0600](#), Feb. 2010.
- [Oli10b] R. I. Oliveira. Sums of random Hermitian matrices and an inequality by Rudelson. *Elect. Comm. Probab.*, 15:203–212, 2010.
- [Pau86] V. I. Paulsen. *Completely Bounded Maps and Dilations*. Number 146 in Pitman Research Notes in Mathematics. Longman Scientific & Technical, New York, NY, 1986.
- [Pet94] D. Petz. A survey of certain trace inequalities. In *Functional analysis and operator theory*, volume 30 of *Banach Center Publications*, pages 287–298, Warsaw, 1994. Polish Acad. Sci.
- [Rus02] M. B. Ruskai. Inequalities for quantum entropy: A review with conditions for equality. *J. Math. Phys.*, 43(9):4358–4375, Sep. 2002.
- [Rus05] M. B. Ruskai. Erratum: Inequalities for quantum entropy: A review with conditions for equality [*J. Math. Phys.* 43, 4358 (2002)]. *J. Math. Phys.*, 46(1):0199101, 2005.
- [RW00] L. C. G. Rogers and D. Williams. *Diffusions, Markov Processes, and Martingales. Volume I: Foundations*. Cambridge Univ. Press, Cambridge, 2nd edition, 2000.
- [Tro10a] J. A. Tropp. Freedman's inequality for matrix martingales. Available at [arXiv](#)., June 2010.
- [Tro10b] J. A. Tropp. From the joint convexity of quantum relative entropy to a concavity theorem of Lieb. Available at [arXiv:1101.1070](#), Dec. 2010.
- [Tro10c] J. A. Tropp. User-friendly tail bounds for sums of random matrices. Available at [arXiv:1004.4389](#), Apr. 2010.