

Randomized and quantum approximate matrix multiplication

Simon Apers¹, Arjan Cornelissen², and Samson Wang³

¹Université Paris Cité, CNRS, IRIF, Paris, France

²Simons Institute, UC Berkeley, California, USA

³Institute for Quantum Information and Matter, Caltech, California, USA

October 10, 2025

Abstract

The complexity of matrix multiplication is a central topic in computer science. While the focus has traditionally been on exact algorithms, a long line of literature also considers randomized algorithms, which return an approximate solution in faster time. In this work, we adopt a unifying perspective that frames these randomized algorithms in terms of mean estimation. Using it, we first give refined analyses of classical algorithms based on random walks by Cohen-Lewis ('99), and based on sketching by Sarlós ('06) and Drineas-Kannan-Mahoney ('06). We then propose an improvement on Cohen-Lewis that yields a single classical algorithm that is faster than all the other approaches, if we assume no use of (exact) fast matrix multiplication as a subroutine. Second, we demonstrate a quantum speedup on top of these algorithms by using the recent quantum multivariate mean estimation algorithm by Cornelissen-Hamoudi-Jerbi ('22).

1 Introduction

Over the last 6 decades, the complexity of multiplying two n -by- n matrices has gone down from the naïve $O(n^3)$ upper bound to $O(n^\omega)$ with $\omega < 2.372$ the matrix multiplication coefficient – see [Alm+25] for the latest improvement. Whether this coefficient can be brought down to its trivial optimum $\omega = 2$ is a major open question in computer science. In the meantime, however, the complexity can be reduced by considering algorithms for *approximate* matrix multiplication.

Definition 1.1 (Approximate matrix multiplication). *Given a description of matrices $\mathbf{A}$ and $\mathbf{B}$, return a matrix $\mathbf{C}$ which satisfies $\|\mathbf{C} - \mathbf{AB}\|_* \leq \varepsilon$ for some chosen norm $\|\cdot\|_*$.*

In our work we specifically consider $\|\cdot\|_* = \|\cdot\|_{\max}$ (the “max-norm problem”) and $\|\cdot\|_* = \|\cdot\|_F$ (the “Frobenius norm problem”).

Approximate matrix multiplication is a setting that allows to exploit *randomized* algorithms, and this seems in contrast to exact matrix multiplication where the best known algorithms are deterministic. In particular, there are well-known Monte Carlo algorithms for approximate matrix multiplication that build on random walks (or “path integral Monte Carlo”) [CL99] and sketching [Sar06; DKM06]. These algorithms can ultimately be viewed from the perspective of *multivariate mean estimation*, i.e., they cleverly define a random variable $\mathbf{X}$, taking values in $\mathbb{R}^d$ with $d = n^2$, whose mean $\boldsymbol{\mu}$ contains the entries of the matrix product $\mathbf{AB}$. Using e.g. the vector Bernstein inequality, one can argue that with high

*All authors contributed equally and are listed in alphabetical order.

probability, the empirical mean $\tilde{\boldsymbol{\mu}}$ of L samples of $\mathbf{X}$, with covariance matrix $\boldsymbol{\Sigma} := \mathbb{E}[(\mathbf{X} - \mathbb{E}[\mathbf{X}])(\mathbf{X} - \mathbb{E}[\mathbf{X}])^T]$, satisfies

$$\text{multivariate mean estimation: } \begin{cases} \|\tilde{\boldsymbol{\mu}} - \boldsymbol{\mu}\|_{\infty} = \mathcal{O}\left(\sqrt{\max_i [\boldsymbol{\Sigma}]_{ii}/L}\right) \\ \|\tilde{\boldsymbol{\mu}} - \boldsymbol{\mu}\|_2 = \mathcal{O}\left(\sqrt{\text{Tr}[\boldsymbol{\Sigma}]/L}\right) \end{cases}. \quad (1)$$

The fact that these randomized algorithms can be phrased as a mean estimation task means that they are amenable to a quantum speedup, since we can make use of *quantum multivariate mean estimation* from [CHJ22]: given quantum access to $\tilde{O}(L)$ samples of $\mathbf{X}$, we can return an estimator $\tilde{\boldsymbol{\mu}}$ that with high probability satisfies

$$\text{quantum multivariate mean estimation: } \begin{cases} \|\tilde{\boldsymbol{\mu}} - \boldsymbol{\mu}\|_{\infty} = \mathcal{O}\left(\sqrt{\text{Tr}[\boldsymbol{\Sigma}]/L}\right) \\ \|\tilde{\boldsymbol{\mu}} - \boldsymbol{\mu}\|_2 = \mathcal{O}\left(\sqrt{d \text{Tr}[\boldsymbol{\Sigma}]/L}\right) \end{cases}. \quad (2)$$

For the ℓ_2 -norm, for instance, this yields a speedup over classical mean estimation when the number of quantum samples $L \gg d = n^2$.

In this work, we study classical and quantum algorithms for approximate matrix multiplication utilizing mean estimation. We make the following contributions:

- We refine the seminal algorithm of Cohen & Lewis [CL99]. While the original algorithm is incomparable to the main other approach based on matrix sketching [Sar06; DKM06], our refinement yields a classical algorithm with better time complexity than prior classical art (if no fast matrix multiplication is allowed).
- We give a unified analysis of sketching algorithms for matrix multiplication [Sar06; DKM06] based on mean estimation. Using this, we construct quantum analogues of all above algorithms. We find a quantum speedup in certain parameter regimes, and there is an unconditional speedup over prior quantum art.
- We extend our sketching-based analysis to give classical and quantum matrix multiplication algorithms for more than two matrices.

We display a simplified account of our results in Table 1 (a more intricate comparison of time complexities can be found in Section 5). We express the runtimes using the element-wise matrix norm: for $\mathbf{M} \in \mathbb{R}^{n \times m}$ and $p, q \in [1, \infty]$, we write

$$\|\mathbf{M}\|_{p,q} = \left(\sum_{j=1}^n \left(\sum_{k=1}^m |[\mathbf{M}]_{jk}|^p \right)^{\frac{q}{p}} \right)^{\frac{1}{q}}.$$

Specifically, $\|\cdot\|_{\max} = \|\cdot\|_{\infty, \infty}$ and $\|\cdot\|_F = \|\cdot\|_{2,2}$ are the max-norm and Frobenius norm, respectively (see Section 2.2 for more details). In order to facilitate navigation of the table, we also present the partial ordering of all norms that appear in Figure 1.

To our knowledge the best previous quantum algorithm is that of Shao [Sha18], which operates via a quantum subroutine to perform inner product estimation on each of the n^2 inner products that determine a matrix product.⁽¹⁾ The data model Shao requires is efficient access to the unitaries which encode the rows and columns of $\mathbf{A}$ and $\mathbf{B}$ in the amplitudes of quantum states – this is achievable in $O(n^2)$ preprocessing time and $\text{polylog}(n)$ time per query with a QROM [Pra14] (we define and discuss quantum memory models in Section 2.3). We note that Shao also presents other algorithms which depend on the condition number of the matrix which we do not display in our table. Whilst these approaches

⁽¹⁾See also the algorithm [Ber+24] which can also be adapted to give an algorithm of similar complexity, though it is specialized to outputting a quantum state.

		Runtime ($\tilde{O}(\cdot)$)		Data model	Multiple matrices?
		$\ \cdot\ _{\max}$	$\ \cdot\ _F$		
Classical	Cohen-Lewis [CL99]	$\frac{\ \overline{\mathbf{A}\mathbf{B}}\ _{1,2}^2}{\varepsilon^2}$	$n \frac{\ \overline{\mathbf{A}\mathbf{B}}\ _{1,2}^2}{\varepsilon^2}$	RAM	✓
	Sarlós [Sar06]	$n^2 \frac{\ \mathbf{A}\ _{2,\infty}^2 \ \mathbf{B}^T\ _{2,\infty}^2}{\varepsilon^2}$	$n^2 \frac{\ \mathbf{A}\ _{2,2}^2 \ \mathbf{B}\ _{2,2}^2}{\varepsilon^2}$	circuit	✓
	Drineas-Kannan-Mahoney [DKM06]	$n^2 \frac{\ \mathbf{A}^T\ _{\infty,2}^2 \ \mathbf{B}\ _{\infty,2}^2}{\varepsilon^2}$	$n^2 \frac{\ \mathbf{A}\ _{2,2}^2 \ \mathbf{B}\ _{2,2}^2}{\varepsilon^2}$	circuit	
	This work (random walk)	$\frac{\ \overline{\mathbf{A}\mathbf{B}}\ _{\infty,\infty} \ \overline{\mathbf{A}\mathbf{B}}\ _{1,1}}{\varepsilon^2}$	$\frac{\ \overline{\mathbf{A}\mathbf{B}}\ _{1,1}^2}{\varepsilon^2}$	RAM	✓
Quantum	Shao [Sha18]	$\frac{\ \mathbf{A}\ _{2,1} \ \mathbf{B}^T\ _{2,1}}{\varepsilon}$	$n \frac{\ \mathbf{A}\ _{2,1} \ \mathbf{B}^T\ _{2,1}}{\varepsilon}$	QROM	
	This work (quantum tug-of-war sketch)	$n^2 \frac{\ \mathbf{A}\ _{2,2} \ \mathbf{B}\ _{2,2}}{\varepsilon}$		quantum circuit	✓
	This work (quantum random walk)	$\frac{\ \overline{\mathbf{A}\mathbf{B}}\ _{1,1}}{\varepsilon}$	$n \frac{\ \overline{\mathbf{A}\mathbf{B}}\ _{1,1}}{\varepsilon}$	QRAM	✓

Table 1: **Time complexities of approximate matrix multiplication (without use of fast matrix multiplication).** For simplicity of presentation, we assume square matrices $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$, and a matrix multiplication exponent of $\omega = 3$. We use an overline over any matrix to denote the matrix of element-wise absolute values. The element-wise norms we use are defined in Section 2.2, and we present their partial ordering in Figure 1. If within a block (i.e., “Classical” or “Quantum”) and a particular column there is an algorithm which is unconditionally faster than the others, we have shaded its complexity. Each of the algorithms utilizes $O(n^2)$ preprocessing time which we omit from the table for simplicity. We discuss the different (classical and quantum) data models in Section 2.3. We present a more detailed table of complexities in Table 5. The cell left open would have a complexity that exceeds $O(n^3)$, and hence is not very useful.

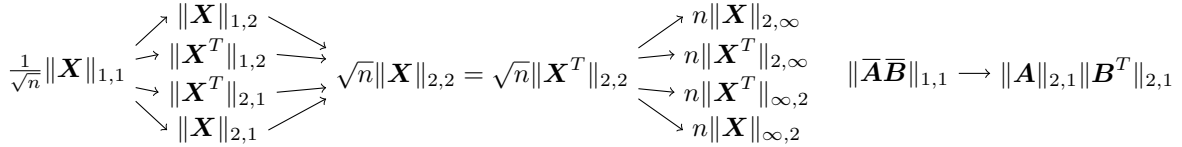


Figure 1: Partial ordering of matrix norms appearing in Table 1. Here, $\mathbf{X}, \mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$, and $\overline{\mathbf{A}}$ represents the matrix containing all element-wise absolute values of $\mathbf{A}$. “ $\rightarrow$ ” denotes “ $\leq$ ”. If there is no arrow in between two quantities then they are incomparable.

could be advantageous in certain settings, we focus on comparing algorithms for generic matrices which are efficient without additional assumptions.

We remark that our algorithms with the most refined time complexities inherit a particular property of the algorithm of Cohen & Lewis [CL99] — that the complexity depends on the norm of a matrix product, rather than a product of matrix norms. This can be especially advantageous for non-negative matrices or matrices with (approximate) sparsity patterns, and this advantage can be amplified in generalizations to multiple matrix products. As an example, when $\mathbf{A}$ and $\mathbf{B}$ are stochastic matrices, all Frobenius norm algorithms are slower than their counterparts based on [CL99] by a factor up to $\Omega(n)$.

1.1 Random walk algorithms

While the Cohen-Lewis algorithm [CL99] works for arbitrary matrices, its idea is best illustrated for the multiplication of stochastic matrices $\mathbf{A}_1, \dots, \mathbf{A}_k$, i.e., every matrix $\mathbf{A}_j \in \mathbb{R}^{n \times n}$ is entry-wise non-negative, and all row-sums equal 1. The matrix product $\mathbf{A}_1, \dots, \mathbf{A}_k$ can then be interpreted as the probability transition matrix of a random walk on a directed graph, displayed in Figure 2. Simulating this random walk samples from the distribution $\mathbf{d}^T \mathbf{A}_1 \cdots \mathbf{A}_k$, for an arbitrary initial distribution $\mathbf{d} \in \mathbb{R}^n$. The original idea from Cohen & Lewis is to reconstruct the matrix product $\mathbf{A}_1 \cdots \mathbf{A}_k$ by learning it row-by-row, i.e., learning all the probability distributions by starting from $\mathbf{d} = \mathbf{e}_j$, with j running from 1 to n .

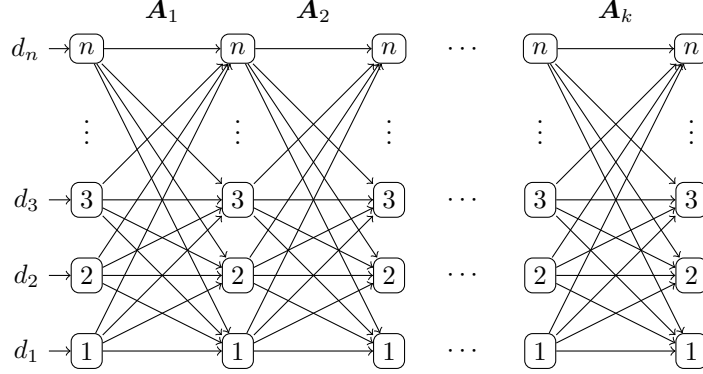


Figure 2: A pictorial representation of the matrix product $\mathbf{A}_1 \cdots \mathbf{A}_k$ of stochastic matrices as a random walk on a directed graph. The random walk samples a vertex on the left side of the graph with probability distribution $\mathbf{d}$, and then transitions through the graph, with the probability transition matrix of every layer equal to $\mathbf{A}_j$. After k steps, the probability distribution on the final layer of vertices is $\mathbf{d}^T \mathbf{A}_1 \cdots \mathbf{A}_k$.

We subtly modify this approach, by analyzing what happens when we start with an arbitrary initial distribution $\mathbf{d}$. We denote the resulting sampled path of vertices $(j_0, \dots, j_k)$, and define the corresponding random variable $\mathbf{X}(j_0, \dots, j_k) = \mathbf{e}_{j_0} \mathbf{e}_{j_k}^T$. We then prove in Lemma 3.3 that this can be used as an estimator for the matrix product, as

$$\mathbb{E}[\mathbf{X}] = \text{diag}(\mathbf{d}) \cdot \mathbf{A}_1 \cdots \mathbf{A}_k, \quad \text{and} \quad \text{Tr}[\Sigma] \leq 1,$$

where we denote $\text{diag}(\mathbf{d})$ as the diagonal matrix who shares diagonal entries with $\mathbf{d}$.

Next, we adapt the above approach to work for matrices with negative entries as well, i.e., we explain how to compute $\mathbf{A}_1 \cdots \mathbf{A}_k$, where $\bar{\mathbf{A}}_1, \dots, \bar{\mathbf{A}}_k$ are all stochastic matrices. The idea here, already observed by Cohen & Lewis, is to label all the edges of the directed graph with the signs of the matrix entries. For a sampled path $(j_0, \dots, j_k)$, we then multiply all the signs together to obtain $s(j_0, \dots, j_k) \in \{\pm 1\}$, and we modify the random variable to be $\mathbf{X}(j_0, \dots, j_k) = s(j_0, \dots, j_k) \mathbf{e}_{j_0} \mathbf{e}_{j_k}^T$. Crucially, this does not impact the upper bound on the trace of the covariance matrix, since we can still show that $\text{Tr}[\Sigma] \leq 1$.

Finally, we adapt the construction to any matrix product, by observing that for any sequence of matrices $\mathbf{A}_1, \dots, \mathbf{A}_k$, we can construct a *stochastic matrix product decomposition*. That is, we can find a vector $\mathbf{d}$ and a sequence of matrices $\mathbf{A}'_1, \dots, \mathbf{A}'_k$, such that $\bar{\mathbf{A}}'_1, \dots, \bar{\mathbf{A}}'_k$ are all stochastic, and such that $\mathbf{A}_1 \cdots \mathbf{A}_k = \text{diag}(\mathbf{d}) \mathbf{A}'_1 \cdots \mathbf{A}'_k$ (see Theorem 3.1). Moreover, we prove that we can perform this stochastic matrix product decomposition as a preprocessing routine in time linear in the size of the input. We can now interpret the vector $\mathbf{d}$ as the probability distribution $\mathbf{d}/\|\mathbf{d}\|_1$, and absorb the resulting factor $\|\mathbf{d}\|_1$ in the random variable, i.e., $\mathbf{X}(j_0, \dots, j_k) = \|\mathbf{d}\|_1 \cdot s(j_0, \dots, j_k) \cdot \mathbf{e}_{j_0} \mathbf{e}_{j_k}^T$. We then find that

$$\mathbb{E}[\mathbf{X}] = \mathbf{A}_1 \cdots \mathbf{A}_k, \quad \max_{ij} [\Sigma]_{ii} \leq \|\mathbf{A}_1 \cdots \mathbf{A}_k\|_{\infty, \infty} \cdot \|\mathbf{A}_1 \cdots \mathbf{A}_k\|_{1, 1}, \quad \text{and} \quad \text{Tr}[\Sigma] \leq \|\mathbf{A}_1 \cdots \mathbf{A}_k\|_{1, 1}^2.$$

Plugging these expressions into Eqs. (1), (2) results in the complexities in Table 1.

To actually implement these algorithms, we require efficiently addressable read/write-memory in both the classical and quantum settings. Putting all these ideas together gives the following result.

Theorem 1.2 (Simplified version of Theorems 3.4 & 3.5). *Consider a matrix product of k matrices $\mathbf{A}_1 \cdots \mathbf{A}_k$ and the problem in Definition 1.1. We use an overline over any matrix to denote the matrix of element-wise absolute values. We give classical algorithms in the RAM model, with run-times*

$$\tilde{O}\left(k \cdot \frac{\|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{\infty, \infty} \cdot \|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,1}}{\varepsilon^2}\right), \quad \text{and} \quad \tilde{O}\left(k \cdot \frac{\|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,1}^2}{\varepsilon^2}\right),$$

that solve the max-norm and the Frobenius-norm problems, respectively, with high probability. Similarly, we give quantum algorithms in the QRAM-model, with run-times

$$\tilde{O}\left(k \cdot \frac{\|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,1}}{\varepsilon}\right), \quad \text{and} \quad \tilde{O}\left(k \cdot \sqrt{nm} \frac{\|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,1}}{\varepsilon}\right),$$

that solve the max-norm and the Frobenius-norm problems, respectively, with high probability. Here, $\mathbf{A}_1$ has n rows, and $\mathbf{A}_k$ has m columns. Both algorithms require classical preprocessing time linear in the size of the input up to polylogarithmic factors.

The quantum complexity for Frobenius norm approximation which we demonstrate above simply follows by considering standard norm conversion between the max-norm and the Frobenius norm, and rescaling the error parameter.

We remark that aside from being unconditionally faster than other algorithms in certain norm regimes (and the absence of fast matrix multiplication), the algorithms based on random walks are particularly well-suited to stochastic matrices. In particular, their complexities grow linearly in the number of matrices, rather than geometrically as is the case for all other approaches.

1.2 Sketching algorithms

The key idea of sketching algorithms as used for matrix multiplication is to sample vectors $\mathbf{s} \in \mathbb{R}^n$ and to compute pairs of vectors $\mathbf{A}\mathbf{s}$ and $\mathbf{s}^T \mathbf{B}$, which together form “sketches” of $\mathbf{A}$ and $\mathbf{B}$, respectively. With judicious choice of $\mathbf{s}$, we can ensure first that the outer product $\mathbf{A}\mathbf{s}\mathbf{s}^T \mathbf{B}$ forms an unbiased estimator of $\mathbf{A}\mathbf{B}$, and second that the second-order moments of this estimator are controlled. Sarlós [Sar06] and Drineas et al. [DKM06] propose different constructions for $\mathbf{s}$, and we build on these for our generalization to multiple matrices, as well as for our quantum algorithms. We will refer to their approaches as the tug-of-war sketch and column-sample sketch, respectively.

The tug-of-war sketch [AMS96] (also known as the AMS sketch) populates entries of $\mathbf{s}$ with a 4-wise independent hash function $h : [n] \rightarrow \{-1, 1\}$ the resulting sketched rows of $\mathbf{A}$ and sketched columns of $\mathbf{B}$ are “pulled” in either the positive or negative directions with equal probability, hence the name “tug-of-war sketch”. The column-sample sketch samples a column of $\mathbf{A}$ (and correspondingly, a row of $\mathbf{B}$) according to a some probability distribution of choice. Thus, $\mathbf{s}$ takes the form of a 1-sparse vector weighted by the chosen probability distribution. Utilizing these sketches we give our results on sketching-based matrix multiplication algorithms.

Theorem 1.3 (Simplified version of Theorems 4.10 and 4.11). *Consider a matrix product of k square matrices $\mathbf{A}_1 \cdots \mathbf{A}_k \in \mathbb{R}^{n \times n}$ and the problem in Definition 1.1. With classical preprocessing time linear in the size of the input up to polylogarithmic factors, we give:*

$$\text{a classical algorithm for the Frobenius norm problem with run-time} \quad \tilde{O}\left(n^2 \cdot \left(3^k \frac{\|\mathbf{A}_1\|_{2,2}^2 \cdots \|\mathbf{A}_k\|_{2,2}^2}{\varepsilon^2}\right)^{\omega-2}\right),$$

$$\text{a quantum algorithm for the max-norm problem with run-time} \quad \tilde{O}\left(n^2 \cdot 3^k \frac{\|\mathbf{A}_1\|_{2,2} \cdots \|\mathbf{A}_k\|_{2,2}}{\varepsilon}\right),$$

with high probability, where ω is the matrix multiplication exponent. Both algorithms do not require efficiently-addressable memory, and can run in the circuit model.

The complexity of our classical algorithm we state above reduces to that of [Sar06] and [DKM06] in the case of multiplication of two matrices. To the authors' knowledge this property is not shared by previous analyses of sketching algorithms for multiple matrices – they either only naturally extend to 3 matrices [DKM06, Appendix A1] or incur large (super-quadratic) runtime in the inverse precision ε^{-1} [Sar06, Appendix B2].

The classical algorithm can be sped up by use of fast matrix multiplication, as can all the sketching-based classical algorithms we consider in this manuscript. Interestingly, for any matrix multiplication exponent $\omega < 2.5$, this yields an algorithm with time complexity scaling sublinearly with ε^{-1} , which has better dependence on the precision even over quantum algorithms. The origin of the advantage using fast matrix multiplication can be viewed as follows. Conventionally, we might wish to estimate $\mathbb{E}[\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}] = \mathbf{A}\mathbf{B}$ by taking a sample mean of c independently drawn samples $\{\mathbf{s}_j\}_j$. Naturally, the time complexity is linear in c . However, due to the structure of the outer product, such a sample mean is also exactly equal to $\mathbf{A}\mathbf{S}^T\mathbf{S}\mathbf{B}$ where the rows of the matrix $\mathbf{S} \in \mathbb{R}^{c \times n}$ are given by $\{\mathbf{s}_j/\sqrt{c}\}_j$. The matrix $\mathbf{A}\mathbf{S}^T\mathbf{S}\mathbf{B}$ can now be evaluated with fast matrix multiplication with time complexity proportional to $c^{\omega-2}$. This provides a way to return the sample mean sublinearly (after preprocessing time proportional to nc to construct $\mathbf{S}$).

Our quantum algorithm as quoted above is based on the tug-of-war sketch utilized by Sarlós. We can also give a quantum algorithm with similar complexity based on the column-sample sketch of Drineas et al., however, unlike the algorithms based on the tug-of-war sketch both classical and quantum approaches can only accommodate multiplication of two matrices. We note that the quantum algorithm we state has unconditionally worse time complexity than the one we give in Theorem 1.2. However, we record it as it uses a weaker assumption on the type of data structure required, which we discuss in Section 2.3 (and it can handle multiple matrices which the algorithm of [Sha18] cannot).

1.3 Discussion

Our quantum algorithms for max-norm approximation depend on the trace of the covariance matrix, which is the same quantity that characterizes approximation in Frobenius norm. We note that in classical mean estimation, there is a different quantity that controls max-norm approximations: the maximum entry-wise variance — as clear examples, note the disparity in the complexity comparison of Table 1. It is an interesting open question if the quantum multivariate mean estimation of [CHJ22] in max-norm can be refined to depend on the same second-moment quantity that we see for classical mean estimation.

In the presentation of Table 1 we have made the simplifying assumption that the matrix multiplication exponent is $\omega = 3$. In reality, one can choose $\omega < 3$ by exploiting fast matrix multiplication, in which case there are interesting asymptotic conclusions about the runtime of sketching algorithms — they can outperform the quantum approaches we present in certain regimes. We remark that the most costly step of our quantum algorithms essentially is matrix-vector multiplication in coherent arithmetic (rather than rectangular matrix multiplication in classical sketching algorithms), and we leave it as an open question whether or not fast matrix multiplication can be exploited in our quantum algorithms. Meanwhile, in comparing our classical algorithm to those which can be sped up with fast matrix multiplication, we do stress that we may be more interested in practical considerations for immediate application. From this perspective, we recall that the Strassen algorithm [Str69] with $\omega \approx 2.8$ is commonly implemented. However, refinements based around the work of Coppersmith and Winograd beyond $\omega < 2.38$ [CW90] are considered “galactic”, and not practically implementable [Ili89].

Finally, our work considers generic unstructured matrices. However, other classical approaches to approximate matrix multiplication are particularly useful in exploiting sparse matrices [Pag13]. We leave open the question of whether we can exploit sparsity in the quantum setting as well.

2 Preliminaries

2.1 Notation

We denote matrices and vectors in bold type $\mathbf{A} \in \mathbb{C}^{N \times N}$ and $\mathbf{a} \in \mathbb{C}^N$. We denote their entries as $[\mathbf{A}]_{ij}$ and a_j respectively for all $i, j \in [N]$. Denote $[\mathbf{A}]_{k\bullet} \in \mathbb{C}^{1 \times N}$ as the k^{th} row matrix of $\mathbf{A}$ for all $k \in [N]$, and denote $[\mathbf{A}]_{\bullet j} \in \mathbb{C}^{N \times 1}$ as the j^{th} column matrix of $\mathbf{A}$ for all $j \in [N]$. For any matrix, we write $\bar{\mathbf{A}}$ as the matrix with element-wise absolute values $[\bar{\mathbf{A}}]_{ij} = |[\mathbf{A}]_{ij}|$. We denote $\mathbf{I}$ as the identity matrix and we denote $\mathbf{1}_m$ as the all-ones vector of dimension m .

The notation $\tilde{O}(\cdot)$ hides polylogarithmic factors, that is we denote $\tilde{O}(g(n)) = O(g(n) \log^k(g(n)))$ for some constant k .

2.2 Linear algebra and probability theory results

For vectors $\mathbf{a} \in \mathbb{R}^n$, we denote the usual ℓ_p -norm as $\|\mathbf{a}\|_p = (\sum_i a_i^p)^{1/p}$. We also make use of some element-wise matrix norms:

Definition 2.1 (Element-wise matrix norms). *Let $\mathbf{M} \in \mathbb{R}^{n \times m}$, and let $p, q \in [1, \infty]$. We write*

$$\|\mathbf{M}\|_{p,q} = \left(\sum_{j=1}^n \left(\sum_{k=1}^m |[\mathbf{M}]_{jk}|^p \right)^{\frac{q}{p}} \right)^{\frac{1}{q}}.$$

We note $\|\cdot\|_{\max} = \|\cdot\|_{\infty, \infty}$ and $\|\cdot\|_F = \|\cdot\|_{2,2}$ are the max-norm and Frobenius norm, respectively.

We can think of this norm intuitively as follows: first we take the vector of ℓ_p -norms of all the rows of the matrix, and then we take the ℓ_q -norm of the resulting vector. If we want to swap the roles of rows and columns, we can consider the norm of the transpose instead. An important subtlety is that $\|\mathbf{A}\|_{p,q}$ is in general not equal to $\|\mathbf{A}^T\|_{q,p}$. On the other hand, if $p = q$, then the norm for any given matrix is the same as the norm for its transpose.

Using this interpretation of the element-wise matrix norms, we also observe that Hölder's inequality holds.

Lemma 2.2 (Hölder's inequality for element-wise matrix norms). *Let $\mathbf{M} \in \mathbb{R}^{n \times m}$. Let $1 \leq p' \leq p \leq \infty$, and $1 \leq q' \leq q \leq \infty$. Then,*

$$\|\mathbf{M}\|_{p,q} \leq \|\mathbf{M}\|_{p',q'} \leq n^{\frac{1}{q} - \frac{1}{q'}} m^{\frac{1}{p} - \frac{1}{p'}} \|\mathbf{M}\|_{p,q}.$$

As a preliminary for classical mean estimation, we recall the following inequalities. They are slight reformulations of the concentration inequalities known as the vector and matrix Bernstein inequalities, see e.g. [Gro11; Tro15].

Lemma 2.3 (Vector and matrix Bernstein inequality). *Let $\tilde{\mathbf{a}}^{(L)} = \sum_{i=1}^L \mathbf{a}_i / L$ be an empirical average over L independent and identically distributed d -dimensional random variables $\{\mathbf{a}_i\}_i$ which satisfy $\mathbb{E}[\mathbf{a}_i] = \mathbf{a}$. Then, with probability at least $(1 - \delta)$ we have*

$$\|\tilde{\mathbf{a}}^{(L)} - \mathbf{a}\|_{\infty} = O\left(\sqrt{\frac{\max_{i \in [d]} [\boldsymbol{\Sigma}]_{ii} \log(d/\delta)}{L}}\right), \quad \text{and} \quad \|\tilde{\mathbf{a}}^{(L)} - \mathbf{a}\|_2 = O\left(\sqrt{\frac{\text{Tr}[\boldsymbol{\Sigma}] \log(1/\delta)}{L}}\right),$$

where $\boldsymbol{\Sigma} = \mathbb{E}[(\mathbf{a}_i - \mathbf{a})(\mathbf{a}_i - \mathbf{a})^T]$ is the covariance matrix.

2.3 Quantum computational model and memory

In both our classical and quantum algorithms, we assume access to different levels of memory architecture (see “Data model” in Table 1). Specifically, some algorithms call an architecture that allows for efficient

reading and writing operations – given a register containing a memory address, we assume to be able to swap the bit at that location in memory with one from another register. In the quantum setting, this means that we assume to have access to a quantum random-access gate (QRAM gate), that for all data $\mathbf{x} \in \{0, 1\}^d$, $j \in [d]$ and $b \in \{0, 1\}$ acts as

$$\text{QRAM} : |x_1, \dots, x_d\rangle |j\rangle |b\rangle \mapsto |x_1, \dots, x_{j-1}, b, x_{j+1}, \dots, x_d\rangle |j\rangle |x_j\rangle.$$

We remark that this is a stronger assumption than having access to a quantum read-only memory (QROM),⁽²⁾ where one assumes to have access to a gate

$$\text{QROM} : |x_1, \dots, x_d\rangle |j\rangle |b\rangle \mapsto |x_1, \dots, x_d\rangle |j\rangle |b \oplus x_j\rangle.$$

Having access to efficiently-addressable write operations is standard in the classical literature, but its quantum counterpart has been somewhat contentious, particularly as implementing an $O(d)$ -sized QRAM/QROM fault-tolerantly may require $O(d)$ (classical) effort [JR23; Dal+25]. We stress here that we only require our quantum algorithm to be able to perform the same operations as its corresponding classical counterpart, i.e., if our original classical algorithms do not actually use efficiently-addressable writing operations, then neither do our corresponding quantum algorithms. In this work we focus on theoretical comparison between the classical and quantum computational models.

In other settings, we do not need to assume any efficient read or write operations. Instead, we operate in a *circuit model*. Here, our classical (or quantum) algorithm can be written wholly in terms of a fixed classical (or quantum) circuit acting on input data [Sip96]. When we consider this setting, it is sufficient to consider the input data in the form of a binary encoding $|x_1\rangle \otimes \dots \otimes |x_d\rangle$.

So far in the above exposition we have used a simplifying picture of binary data. In both the classical and quantum models, we make the assumption that we can represent real numbers exactly in memory, and that we can retrieve them and perform basic arithmetic operations (that is, addition, multiplication, subtraction and division) on them in constant time. This is a standard assumption in the classical literature and we make the same assumption quantumly.

2.4 Classical subroutines

First, we recall how we can sample from a given distribution of n outcomes, with probabilities written in memory. The most efficient method to achieve this is known as the alias method, and we recall a result by Vose.

Theorem 2.4 (Alias method [Vos91]). *Given a probability distribution over n elements, with the probabilities written in memory, we can sample from it using $O(n)$ preprocessing time in the RAM-model, and $O(1)$ time per sample performing only read operations.*

The fact that we only perform read operations per sample is important, because it means that we can port this algorithm to the QROM-setting. Indeed, we can implement Vose’s sampling algorithm using a quantum read-only memory of size $O(n)$, which we initialize using classical RAM only, and then access with read-only operations by a quantum algorithm.

2.5 Quantum subroutines

We will make use of a multivariate quantum mean estimation routine from [CHJ22] (and later refined in [Tan25]), which we restate here for convenience.

⁽²⁾We also remark here that the many different nomenclatures for read-only and read-write operations on quantum memory. Instead of the acronyms QROM/QRAM [Cor25], one can find QRAM/QRAG [All+23], QCRAM/full-QRAM [Ape20], QACM/QAQM [NS20], QRACM/QRAQM [Kup13] and more. We prefer the nomenclature QROM/QRAM, to mimic the nomenclature ROM/RAM for the corresponding classical counterparts.

Theorem 2.5 (Quantum mean estimation [CHJ22, Theorem 3.4]). *Let $d \in \mathbb{N}$, $\delta \in (0, 1)$, $(\Omega, \mathbb{P})$ be a probability space, and $\mathbf{X} : \Omega \rightarrow \mathbb{R}^d$ a random variable. Suppose we have access to the following two operations:*

$$U_{\mathbb{P}} : |0\rangle \mapsto \sum_{\omega \in \Omega} \sqrt{\mathbb{P}(\omega)} |\omega\rangle, \quad \text{and} \quad O_{\mathbf{X}} : |\mathbf{Y}\rangle |\omega\rangle |0\rangle \mapsto |\mathbf{Y}\rangle |\omega\rangle |\mathbf{Y}^T \mathbf{X}(\omega)\rangle,$$

where $\omega \in \Omega$ and $\mathbf{Y} \in \mathbb{R}^d$. Let $\boldsymbol{\mu}$ and $\boldsymbol{\Sigma}$ be the expectation and covariance matrix of X , and let $N \geq \log(d/\delta)$. Then, we can produce a vector $\tilde{\boldsymbol{\mu}} \in \mathbb{R}^d$, such that $\|\tilde{\boldsymbol{\mu}} - \boldsymbol{\mu}\|_{\infty} \leq \sqrt{\text{Tr}[\boldsymbol{\Sigma}] \log(d/\delta)/N}$, with probability at least $(1 - \delta)$, with $\tilde{O}(N)$ queries to $U_{\mathbb{P}}$, $O_{\mathbf{X}}$ and an inner product routine between the random variable $X(\omega)$ and an arbitrary vector $\mathbf{Y} \in \mathbb{R}^d$, and $\tilde{O}(d)$ additional elementary gates.

The time complexity is not explicitly analyzed in [CHJ22], so we discuss this now here. The three key steps of the algorithm consist of (1) preparation of a uniform superposition of a grid of points $\mathbf{Y} \in G \subset (-1/2, 1/2)^d$; (2) computing the inner product of $\mathbf{Y}$ with a data vector; (3) applying the inverse quantum Fourier transform.

Step (1) costs $\tilde{O}(d)$ without any need of quantum memory assumptions, due to efficient state preparation algorithms for uniform superpositions. The inverse quantum Fourier transform is also efficient. We can thus focus on step (2) as the non-generic cost. Explicitly, we require an operation that computes the inner product, i.e., for all $\mathbf{Y} \in \mathbb{R}^d$,

$$O'_{\mathbf{X}} : |\mathbf{Y}\rangle |\omega\rangle |0\rangle \mapsto |\mathbf{Y}\rangle |\omega\rangle |\mathbf{Y}^T \mathbf{X}(\omega)\rangle.$$

Since we assume to be able to do arithmetic operations in unit time, this leads to an additional $O(d)$ overhead per sample in the circuit model, bringing the total time complexity to $\tilde{O}(dN)$. A similar observation was made in [AG23]. However, what we will use in this work is that if we can implement this inner product more efficiently, say with time complexity η , then the resulting total time complexity becomes $\tilde{O}(\eta N + d)$.

3 The Cohen-Lewis algorithm

In this section, we describe our first classical and quantum algorithms for implementing approximate matrix multiplication. The idea of this algorithm is based on the seminal work by Cohen and Lewis [CL99].

Suppose that we are given k matrices, $\mathbf{A}_1, \dots, \mathbf{A}_k$, such that $\mathbf{A}_j \in \mathbb{R}^{n_{j-1}, n_j}$, for all $j \in [k]$. Then, $\mathbf{A}_1 \cdots \mathbf{A}_k$ is well-defined, and our goal is to output an approximation of this matrix product. We write $n = n_0$ and $m = n_k$, so that the resulting matrix product is of dimension $n \times m$. We assume that full descriptions of the matrices are available to us in memory.

3.1 Stochastic matrix product decomposition

We start by showing that we can assume all matrices have absolute row-sum 1, without loss of generality. To that end, we show that we can classically perform a stochastic matrix product decomposition as a preprocessing step, using time and memory that is similar to the size of the matrices themselves.

We start by formally introducing the decomposition.

Theorem 3.1 (Stochastic matrix product decomposition). *Let $k \in \mathbb{N}$, $n_j \in \mathbb{N}$ for all $\ell \in [k]$, and let $\mathbf{A}_{\ell} \in \mathbb{R}^{n_{\ell-1} \times n_{\ell}}$ for all $\ell \in [k]$. We recursively define diagonal matrices $\mathbf{D}_0, \dots, \mathbf{D}_k$, such that*

$$\mathbf{D}_k = I, \quad \text{and} \quad \forall \ell \in [k], i \in [n_{\ell-1}], \quad [\mathbf{D}_{\ell-1}]_{ii} = \sum_{j=1}^{n_{\ell}} [\bar{\mathbf{A}}_{\ell}]_{ij} \cdot [\mathbf{D}_{\ell}]_{jj} = \|[\bar{\mathbf{A}}_{\ell} \mathbf{D}_{\ell}]_{i\bullet}\|_1.$$

Next, for all $\ell \in [k]$, we let $\mathbf{A}'_{\ell} = \mathbf{D}_{\ell-1}^{-1} \mathbf{A}_{\ell} \mathbf{D}_{\ell}$. Then, $\mathbf{A}_1 \cdots \mathbf{A}_k = \mathbf{D}_0 \mathbf{A}'_1 \cdots \mathbf{A}'_k$, all $(\mathbf{A}'_j)$'s have all

row-sums equal to 1, and for all $\ell \in [k]$ and $i \in [n_{\ell-1}]$,

$$[\mathbf{D}_{\ell-1}]_{ii} = \sum_{j=1}^{n_\ell} [\overline{\mathbf{A}}_\ell \cdots \overline{\mathbf{A}}_k]_{ij} = [\overline{\mathbf{A}}_\ell \cdots \overline{\mathbf{A}}_k \mathbf{1}_{n_k}]_i.$$

Moreover, we can compute all $\mathbf{D}_j$'s and $(\mathbf{A}'_j)$'s classically in time linear in the input size.

Proof. It is clear by plugging in the definitions of $\mathbf{A}'_\ell$ that for all $\ell \in [k]$,

$$\mathbf{D}_{\ell-1} \mathbf{A}'_\ell \cdots \mathbf{A}'_k = \mathbf{A}_\ell \cdots \mathbf{A}_k \mathbf{D}_k = \mathbf{A}_\ell \cdots \mathbf{A}_k.$$

Moreover, for every $\ell \in [k]$ and $i \in [n_{\ell-1}]$, we have

$$\sum_{j=1}^{n_\ell} |[\mathbf{A}'_\ell]_{ij}| = \frac{1}{[\mathbf{D}_{\ell-1}]_{ii}} \cdot \sum_{j=1}^{n_\ell} |[\mathbf{A}_\ell]_{ij}| \cdot [\mathbf{D}_\ell]_{jj} = 1.$$

The characterization of $[\mathbf{D}_{\ell-1}]_{ii}$ follows by induction. Finally, it is clear from the recursive definition how we compute the $\mathbf{D}_\ell$'s and $(\mathbf{A}'_\ell)$'s in linear time. $\square$

Using this decomposition, we can also compute some element-wise norms of the matrix product efficiently. Specifically, for $p = 1$ and $q \in [1, \infty]$, we show that the element-wise matrix norm can be computed in time linear in the input size.

Lemma 3.2. *Let $k \in \mathbb{N}$, $n_j \in \mathbb{N}$ for all $\ell \in [k]_0$, and let $\mathbf{A}_\ell \in \mathbb{R}^{n_{\ell-1} \times n_\ell}$ for all $\ell \in [k]$. Let $q \in [1, \infty]$. Let $\mathbf{D}_0$ be as in Theorem 3.1, and let $\mathbf{d}_0$ be the vector containing the diagonal of $\mathbf{D}_0$. Then,*

$$\|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,q} = \|\mathbf{d}_0\|_q,$$

and as such it can be computed in time linear in the total size of the input.

Proof. The result follows directly from the characterization of $(\mathbf{D}_0)_{ii}$ in Theorem 3.1. $\square$

We now state the crucial observation that motivates the stochastic matrix product decomposition. It is a variant on path-integral Monte Carlo, and it mimics the approach taken by Cohen and Lewis [CL99], but differs in one crucial aspect: Cohen and Lewis estimate every row of the product matrix in a separate subroutine. Here, we sample over the rows as well, leading to a random variable with expectation exactly equal to the matrix product we are after. This is the core of the improvement we present in this work.

Lemma 3.3. *Let $k \in \mathbb{N}$, $n_j \in \mathbb{N}$ for all $\ell \in [k]_0$, and let $\mathbf{A}_\ell \in \mathbb{R}^{n_{\ell-1} \times n_\ell}$ for all $\ell \in [k]$. Let $\mathbf{D}_0, \mathbf{A}'_1, \dots, \mathbf{A}'_k$ be the stochastic matrix product decomposition from Theorem 3.1. Let $\mathbf{d}_0$ be the vector containing the diagonal entries of $\mathbf{D}_0$. Then, we define the probability distribution p by sampling $j_0 \sim \mathbf{q}$, for a probability distribution $\mathbf{q}$, and then $j_\ell \sim (\overline{\mathbf{A}}'_\ell)_{j_{\ell-1}, \bullet}$, for all $\ell \in [k]$. Finally, we define the random variable*

$$\mathbf{X}(j_0, \dots, j_k) = \frac{d_{j_0}}{q_{j_0}} \cdot \left(\prod_{\ell=1}^k \text{sign}([\mathbf{A}'_\ell]_{j_{\ell-1}, j_\ell}) \right) \mathbf{e}_{j_0} \mathbf{e}_{j_k}^T,$$

where $\text{sign}(x)$ is 1 if $x \geq 0$, and -1 otherwise. Then,

$$\mathbf{A}_1 \cdots \mathbf{A}_k = \mathbb{E}_{(j_0, \dots, j_k) \sim p} [\mathbf{X}], \quad \text{and} \quad \text{Var}_{(j_0, \dots, j_k) \sim p} [\mathbf{X}_{ij}] \leq \frac{(d_0)_i}{q_i} \cdot (\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k)_{ij}.$$

Hence, if we take $\mathbf{q} = \mathbf{d}_0 / \|\mathbf{d}_0\|_1$, then we obtain

$$\max_{i \in [n_0], j \in [n_k]} \text{Var}[\mathbf{X}_{ij}] \leq \|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{\infty, \infty} \cdot \|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,1}, \quad \text{and} \quad \text{Tr}[\Sigma] \leq \|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,1}^2,$$

where $\Sigma = \text{Cov}(\mathbf{X})$. If we instead take $q_i \propto \|(\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k)_{i,\bullet}\|_1 \cdot \|(\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k)_{i,\bullet}\|_\infty$, then

$$\max_{i \in [n_0], j \in [n_k]} \text{Var}[\mathbf{X}_{ij}] \leq \sum_{i=1}^{n_0} \|(\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k)_{i,\bullet}\|_1 \cdot \|(\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k)_{i,\bullet}\|_\infty,$$

and if we take $\mathbf{q} = \mathbf{d}_0^2 / \|\mathbf{d}_0\|_2^2$, then we obtain

$$\max_{i \in [n_0], j \in [n_k]} \text{Var}[\mathbf{X}_{ij}] \leq \|\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k\|_{1,2}^2.$$

Proof. Unwrapping the expectation yields

$$\begin{aligned} & \mathbb{E}_{(j_0, \dots, j_k) \sim p} [\mathbf{X}(j_0, \dots, j_k)] \\ &= \frac{(\mathbf{d}_0)_{j_0}}{\mathbf{q}_{j_0}} \cdot \sum_{j_0=1}^{n_0} \mathbf{q}_{j_0} \cdot \sum_{j_1=1}^{n_1} |\mathbf{A}'_1|_{j_0, j_1} \text{sign}((\mathbf{A}'_1)_{j_0, j_1}) \cdots \sum_{j_k=1}^{n_k} |\mathbf{A}'_k|_{j_{k-1}, j_k} \text{sign}((\mathbf{A}'_k)_{j_{k-1}, j_k}) \mathbf{e}_{j_0} \mathbf{e}_{j_k}^T = \\ &= \sum_{j_0=1}^{n_0} \mathbf{e}_{j_0} (\mathbf{d}_0)_{j_0} \cdot \sum_{j_1=1}^{n_1} (\mathbf{A}'_1)_{j_0, j_1} \cdots \sum_{j_k=1}^{n_k} (\mathbf{A}'_k)_{j_{k-1}, j_k} \mathbf{e}_{j_k}^T \\ &= \sum_{j_0=1}^{n_0} \mathbf{e}_{j_0} \mathbf{e}_{j_0}^T \mathbf{D}_0 \mathbf{e}_{j_0} \cdot \sum_{j_1=1}^{n_1} \mathbf{e}_{j_0}^T \mathbf{A}'_1 \mathbf{e}_{j_1} \cdots \sum_{j_k=1}^{n_k} \mathbf{e}_{j_{k-1}}^T \mathbf{A}'_k \mathbf{e}_{j_k} \mathbf{e}_{j_k}^T \\ &= \sum_{j'_0=1}^{n_0} \mathbf{e}_{j'_0} \mathbf{e}_{j'_0}^T \mathbf{D}_0 \sum_{j_0=1}^{n_0} \mathbf{e}_{j_0} \mathbf{e}_{j_0}^T \mathbf{A}'_1 \sum_{j_1=1}^{n_1} \mathbf{e}_{j_1} \mathbf{e}_{j_1}^T \cdots \mathbf{A}'_k \sum_{j_k=1}^{n_k} \mathbf{e}_{j_k} \mathbf{e}_{j_k}^T = \mathbf{D}_0 \mathbf{A}'_1 \cdots \mathbf{A}'_k = \mathbf{A}_1 \cdots \mathbf{A}_k. \end{aligned}$$

Next, for the variance of the (i, j) -th entry of $\mathbf{X}$, we observe that

$$\begin{aligned} \text{Var}_{(j_0, \dots, j_k) \sim p} [\mathbf{X}(j_0, \dots, j_k)_{ij}] &\leq \mathbb{E}_{(j_0, \dots, j_k) \sim p} [\mathbf{X}(j_0, \dots, j_k)_{ij}^2] = \left(\frac{(\mathbf{d}_0)_i}{\mathbf{q}_i} \right)^2 \cdot \mathbb{P}_{(j_0, \dots, j_k) \sim p} [j_0 = i \wedge j_k = j] \\ &= \frac{(\mathbf{d}_0)_i}{\mathbf{q}_i} \cdot (\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k)_{ij}. \end{aligned}$$

The final expressions follow directly. $\square$

From the above analysis, we observe that when we choose $\mathbf{q} \propto \mathbf{d}_0$, we minimize the trace of the covariance matrix, which means that this choice yields the best algorithm for the Frobenius-norm problem. On the other hand, by choosing $\mathbf{q} \propto \mathbf{d}_0^2$, we recover the complexity from [CL99] for the max-norm problem. Both choices for $\mathbf{q}$ yield incomparable complexities for the max-norm problem.

It is tempting to simply choose the optimal $\mathbf{q}$ for the max-norm problem, which is when $q_i \propto \|(\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k)_{i,\bullet}\|_\infty \cdot \|(\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k)_{i,\bullet}\|_1$. However, the problem is that it is not clear how to compute this choice for $\mathbf{q}$ efficiently in linear time as a preprocessing step. Hence, we are left with an algorithm with incomparable complexity to the max-norm algorithm given by Cohen and Lewis' approach, for instance if we take $\mathbf{q} = \mathbf{d}_0 / \|\mathbf{d}_0\|_1$ (see next section for further discussion). If we additionally have good upper bounds on $\|(\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k)_{i,\bullet}\|_\infty$ a priori, for all $i \in [n]$, then we could indeed further improve the algorithm.

3.2 Classical algorithm

The idea for the classical algorithm is to compute the expectation from Lemma 3.3 using Monte Carlo sampling. It remains to figure out the number of samples required, and the cost to perform sampling, which we analyze in the following theorem.

Theorem 3.4 (Classical random walk algorithm). *Let $k \in \mathbb{N}$, $n_j \in \mathbb{N}$ for all $\ell \in [k]_0$, and let $\mathbf{A}_\ell \in \mathbb{R}^{n_{\ell-1} \times n_\ell}$ for all $\ell \in [k]$. Let $\varepsilon > 0$ and $\delta \in (0, 1)$, and let $n_0 = n$ and $n_k = m$. Then, with classical*

preprocessing time linear in the size of the input up to polylogarithmic factors, we give a classical algorithm in the RAM model with run-time

$$\tilde{O}\left(k \cdot \frac{\|\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k\|_{\infty, \infty} \cdot \|\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k\|_{1,1}}{\varepsilon^2} \cdot \log\left(\frac{d}{\delta}\right)\right) \quad \text{and} \quad \tilde{O}\left(k \cdot \frac{\|\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k\|_{1,1}^2}{\varepsilon^2} \cdot \log\left(\frac{1}{\delta}\right)\right),$$

that computes an approximation $\tilde{\mathbf{A}}$ of $\mathbf{A}_1 \cdots \mathbf{A}_k$ with probability at least $(1 - \delta)$, satisfying the precision bounds $\|\tilde{\mathbf{A}} - \mathbf{A}_1 \cdots \mathbf{A}_k\|_{\max} \leq \varepsilon$ and $\|\tilde{\mathbf{A}} - \mathbf{A}_1 \cdots \mathbf{A}_k\|_F \leq \varepsilon$, respectively.

Proof. We take the average of N realizations of the random variable from Lemma 3.3. Using the bound on the trace of the covariance matrix, and the vector-Bernstein inequality Lemma 2.3, we obtain that with probability at least $(1 - \delta)$,

$$\|\tilde{\mathbf{A}} - \mathbf{A}_1 \cdots \mathbf{A}_k\|_F = \mathcal{O}\left(\|\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k\|_{1,1} \cdot \sqrt{\frac{\log(1/\delta)}{N}}\right),$$

from which we deduce that we can take $N = \Theta(\|\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k\|_{1,1}^2 / \varepsilon^2 \cdot \log(1/\delta))$ to ensure that the Frobenius error is at most ε . A similar analysis yields the corresponding complexity for the max-norm case.

It remains to prove that we can generate each sample in time $O(k)$. To that end, observe from Theorem 3.1 that we can compute the stochastic matrix product decomposition in time linear in the size of the input, and from Lemma 3.2 that we can compute the choice for N in time linear in the size of the input as well. Next, we observe that we can use Theorem 2.4 to sample from $\mathbf{d}_0 / \|\mathbf{d}_0\|_1$ and from any of the rows of the matrices $\mathbf{A}'_\ell$ in polylogarithmic time, with total preprocessing time again linear in the size of the input. Finally, since we need to sample from k probability distributions, and multiply k signs to obtain one realization of the random variable, the cost per sample is indeed $O(k)$. $\square$

The resulting complexity for the Frobenius norm unconditionally improves over the result in [CL99]. Indeed, by comparing the expressions in Table 1, we can use Hölder's inequality to conclude that

$$\|\mathbf{A}_1 \cdots \mathbf{A}_k\|_{1,1}^2 \leq n \|\mathbf{A}_1 \cdots \mathbf{A}_k\|_{1,2}^2.$$

The situation is a bit more complicated for the max-norm problem, since the bound from [CL99] and our complexity are incomparable. Indeed, if the resulting matrix product $\mathbf{A}$ is the $n \times n$ all-ones matrix, then the Cohen-Lewis bound is $\|\mathbf{A}\|_{1,2}^2 = n^3$, whereas our bound is $\|\mathbf{A}\|_{\infty, \infty} \cdot \|\mathbf{A}\|_{1,1} = n^2$. On the other hand, if the resulting matrix product $\mathbf{A}$ is $\text{diag}(\sqrt{n}, 1, \dots, 1)$, then the Cohen-Lewis bound is $2n$, whereas our bound is $n^{3/2}$.

3.3 Quantum algorithm

In the quantum setting, we wish to speed-up the classical algorithm using the multivariate quantum mean estimation routine [CHJ22]. We describe our result in the following theorem, in the proof of which we compute the costs of implementing the input routines for this procedure.

Theorem 3.5 (Quantum random walk algorithm). *Let $k \in \mathbb{N}$, $n_j \in \mathbb{N}$ for all $\ell \in [k]_0$, and let $\mathbf{A}_\ell \in \mathbb{R}^{n_{\ell-1} \times n_\ell}$ for all $\ell \in [k]$. Let $\varepsilon > 0$ and $\delta \in (0, 1)$, and let $n_0 = n$ and $n_k = m$. Then, with classical preprocessing in the RAM model with time linear in the size of the input up to polylogarithmic factors, there exists a quantum algorithm in the QRAM model with run-time*

$$\tilde{O}\left(k \cdot \frac{\|\bar{\mathbf{A}}_1 \cdots \bar{\mathbf{A}}_k\|_{1,1}}{\varepsilon} \cdot \log\left(\frac{nm}{\delta}\right)\right),$$

that computes an approximation $\tilde{\mathbf{A}}$ satisfying $\|\tilde{\mathbf{A}} - \mathbf{A}_1 \cdots \mathbf{A}_k\|_{\max} < \varepsilon$, with probability at least $(1 - \delta)$.

Proof. We take N as the number of q-samples in the quantum multivariate mean estimation algorithm. Then, combining Lemma 3.3 and Theorem 2.5, the resulting estimate $\tilde{\mathbf{A}}$ satisfies with probability at least

$(1 - \delta)$,

$$\|\tilde{\mathbf{A}} - \mathbf{A}_1 \cdots \mathbf{A}_k\|_{\max} = \tilde{O}\left(\|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,1} \cdot \frac{\log(nm/\delta)}{N}\right),$$

and so it suffices to choose $N = \Theta(\|\overline{\mathbf{A}}_1 \cdots \overline{\mathbf{A}}_k\|_{1,1} \log(nm/\delta)/\varepsilon)$ to ensure that the error is at most ε .

It remains to prove that we can implement the input oracles in time $\tilde{O}(k)$. To that end, observe that coherent sampling from p can be done by implementing the classical sampling procedure reversibly, as described in [Mon15], with the same time complexity $\tilde{O}(k)$. Thus, total cost per sample for implementing $U_{\mathbb{P}}$ is indeed $\tilde{O}(k)$.

Finally we compute the cost of implementing the oracle computing the random variable. To that end, observe that it must implement the following operation:

$$O_{\mathbf{X}} : |j_0, \dots, j_k\rangle |0\rangle \mapsto |j_0, \dots, j_k\rangle \left| \prod_{\ell=1}^k \text{sign}([\mathbf{A}'_{\ell}]_{j_{\ell-1}, j_{\ell}}) \cdot \mathbf{e}_{j_0} \mathbf{e}_{j_k}^T \right\rangle.$$

This essentially boils down to k look-ups of the signs of the entries in the input, and then $O(k)$ multiplications of signs. This can be implemented in $\tilde{O}(k)$ time with the aid of QROM access to preprocessed data. We complete the proof by observing that we can also implement the inner product routine in the same way if we have access to QRAM, i.e., we can store the matrix $|\mathbf{Y}\rangle$ in QRAM and index into it efficiently to spend at most $\tilde{O}(k)$ time to implement

$$O'_{\mathbf{X}} : |\mathbf{Y}\rangle |j_0, \dots, j_k\rangle |0\rangle \mapsto |\mathbf{Y}\rangle |j_0, \dots, j_k\rangle \left| \prod_{\ell=1}^k \text{sign}([\mathbf{A}'_{\ell}]_{j_{\ell-1}, j_{\ell}}) Y_{j_0, j_k} \right\rangle. \quad \square$$

This approach gives us a quantum algorithm that estimates the matrix product in max-norm, whereas the corresponding classical algorithm provides an approximation in terms of Frobenius norm. We can convert between the two using norm conversion, since $\|\cdot\|_F \leq \sqrt{nm} \|\cdot\|_{\max}$. The resulting expressions can be found in Table 1 and Table 5.

4 Sketching-based algorithms

4.1 Generic analysis

In this section we provide a unifying picture for matrix multiplication algorithms via sketching, which encapsulates the algorithms of [DKM06] and [Sar06], and also enables us to construct classical algorithms for multiple matrices as well as quantum algorithms. For simplicity, in this part of the exposition let us assume both $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$, though the discussion also generalizes to rectangular matrices. Our goal is to construct a matrix-valued random variable $\mathbf{S} \in \mathbb{R}^{c \times n}$ (with $c < n$) such that for all $\mathbf{A}, \mathbf{B}$ we have

$$\mathbb{E}[\mathbf{A} \mathbf{S}^T \mathbf{S} \mathbf{B}] = \mathbf{A} \mathbf{B}, \quad (3)$$

which in particular implies that $\mathbb{E}[\mathbf{S}^T \mathbf{S}] = \mathbf{I}$. The matrices $\mathbf{A} \mathbf{S}^T$ and $\mathbf{S} \mathbf{B}$ may be thought of as sketches of $\mathbf{A}$ and $\mathbf{B}$ respectively. Specifically, judicious choices of $\mathbf{S}$ can enable efficient approximation of $\mathbf{A} \mathbf{B}$, as the complexity of evaluating $\mathbf{A} \mathbf{S}^T \mathbf{S} \mathbf{B}$ is $O(n^2 c^{\omega-2})$, compared to $O(n^{\omega})$ for the exact matrix product (here ω can either be $\omega = 3$, if we disallow fast matrix multiplication, or else it is the matrix multiplication exponent $\omega < 2.37$). Supposing the rows of $\mathbf{S}$ are chosen independently, and we denote them as the column vectors $\{\hat{\mathbf{s}}_j\}_{j=1}^c$, then we can also write

$$\mathbb{E}[\mathbf{A} \mathbf{s} \mathbf{s}^T \mathbf{B}] = \mathbf{A} \mathbf{B}, \quad (4)$$

for a vector-valued random variable $\mathbf{s} := \hat{\mathbf{s}} \sqrt{c}$, and we may view the matrix multiplication as a mean over outer products, each costing $O(n^2)$ times to evaluate. In this perspective, we can view one instance of $\mathbf{A} \mathbf{S}^T \mathbf{S} \mathbf{B}$ as a *sample mean* of $\mathbf{A} \mathbf{s} \mathbf{s}^T \mathbf{B}$ over c samples. For the sake of analysis, we use this picture

and find it useful. However, we stress that the quantum and classical algorithm algorithmically adopt different approaches. The classical algorithms use Eq. (3), explicitly computing $\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{S}\mathbf{B}$, and thus can exploit fast matrix multiplication; whereas the quantum algorithms use Eq. (4) and explicitly implement a mean estimation of outer products.

Now that the problem is framed in terms of mean estimation, we can determine the convergence speed by investigating properties of the covariance matrix so that we may use Lemma 2.3 and Theorem 2.5 to analyze the classical and quantum algorithm respectively.

Lemma 4.1 (Covariance matrix for sketching algorithms). *Consider the matrix-valued random variable $\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}$. The covariance matrix Σ satisfies*

$$\begin{aligned}\max_i [\Sigma]_{ii} &= \max_{ij} \text{Var} [\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}]_{ij} = \max_{ij} \sum_{k\ell qr} [\mathbf{A}]_{ik} [\mathbf{B}]_{\ell j} [\mathbf{A}]_{iq} [\mathbf{B}]_{rj} \text{Cov} [s_k s_\ell, s_q s_r], \\ \text{Tr}[\Sigma] &= \mathbb{E} [\|\mathbf{A}\mathbf{B} - \mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{S}\mathbf{B}\|_F^2] = \sum_{ij} \sum_{k\ell qr} [\mathbf{A}]_{ik} [\mathbf{B}]_{\ell j} [\mathbf{A}]_{iq} [\mathbf{B}]_{rj} \text{Cov} [s_k s_\ell, s_q s_r].\end{aligned}$$

Proof. In order to define the covariance matrix in terms of a vector-valued random variable, denote the super index $\alpha = (i, j)$. Recall that the diagonal entries of the covariance matrix are variances, that is

$$[\Sigma]_{\alpha\alpha} = \text{Var} [\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}]_{\alpha} = \text{Var} [\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}]_{ij}$$

From this the first expression for $\max_{\alpha} [\Sigma]_{\alpha\alpha}$ follows, and the first expression for $\text{Tr}[\Sigma]$ follows by noting that

$$\begin{aligned}\text{Tr}[\Sigma] &= \sum_{ij} \text{Var} [\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}]_{ij} \\ &= \sum_{ij} \mathbb{E} [(\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B} - \mathbf{A}\mathbf{B})_{ij}^2] \\ &= \mathbb{E} \left[\sum_{ij} (\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B} - \mathbf{A}\mathbf{B})_{ij}^2 \right] \\ &= \mathbb{E} [\|\mathbf{A}\mathbf{B} - \mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}\|_F^2].\end{aligned}$$

The final expressions can be gleaned by using the standard formula for the variance of a sum of random variables

$$\begin{aligned}\text{Var} [\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}]_{ij} &= \text{Var} \left[\sum_{k\ell qr} [\mathbf{A}]_{ik} [\mathbf{s}\mathbf{s}^T]_{k\ell} [\mathbf{B}]_{\ell j} [\mathbf{A}]_{iq} [\mathbf{s}\mathbf{s}^T]_{qr} [\mathbf{B}]_{rj} \right] \\ &= \sum_{k\ell qr} [\mathbf{A}]_{ik} [\mathbf{B}]_{\ell j} [\mathbf{A}]_{iq} [\mathbf{B}]_{rj} \text{Cov} [[\mathbf{s}\mathbf{s}^T]_{k\ell}, [\mathbf{s}\mathbf{s}^T]_{qr}],\end{aligned}$$

where to obtain the final expression we observe that $[\mathbf{s}\mathbf{s}^T]_{k\ell} = s_k s_\ell$. $\square$

Lemma 4.1 shows that the convergence of both classical and quantum algorithms is controlled simply by covariances of elements of $\mathbf{s}$. In what follows we present two choices for $\mathbf{s}$, and their implications for both classical and quantum algorithms.

4.2 Classical algorithms

Let us now consider a general rectangular matrix product $\mathbf{A}\mathbf{B}$ for $\mathbf{A} \in \mathbb{R}^{n \times q}$, $\mathbf{B} \in \mathbb{R}^{q \times m}$. In classical algorithms we will sample one matrix $\mathbf{S} \in \mathbb{R}^{c \times q}$ with rows constructed independently from some distribution, and explicitly compute $\mathbf{A}\mathbf{S}^T\mathbf{S}\mathbf{B}$ (recall that in order to analyze convergence, we will view this as taking a sample mean over c samples of $\mathbf{A}\mathbf{s}\mathbf{s}^T\mathbf{B}$). This costs $O((mn + nq + mq)c^{\omega-2})$ runtime. We recall

that the approximation error in Frobenius norm and max-norm are controlled by $\text{Tr}[\Sigma]$ and $\max_i[\Sigma]_{ii}$ respectively.

4.2.1 Tug-of-war sketch

In the tug-of-war sketch [AMS96] (also known as AMS sketch), $\mathbf{S}$ is chosen to be the dense matrix with entries $[\mathbf{S}]_{ij} = h_i(j)/\sqrt{c}$ with each $h_i : [n] \rightarrow \{-1, 1\}$ drawn uniformly at random.⁽³⁾ In the mean estimation picture, our sketching vectors $\mathbf{s}$ thus satisfy $s_j = h(j)$.

Lemma 4.2 (Tug-of-war sketch covariance matrix). *For the tug-of-war sketch, the covariance matrix Σ satisfies*

$$\text{Tr}[\Sigma] = \|\mathbf{A}\|_F^2 \|\mathbf{B}\|_F^2 + \|\mathbf{AB}\|_F^2 - 2 \sum_{k=1}^n \|\mathbf{A}_{\bullet k}\|_2^2 \|\mathbf{B}_{k\bullet}\|_2^2,$$

and

$$\max_i [\Sigma]_{ii} = \|\mathbf{A}\|_{2,\infty}^2 \|\mathbf{B}^T\|_{2,\infty}^2 + \|\mathbf{AB}\|_{\max}^2 - 2 \min_{ij} \sum_k [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{kj}^2.$$

Proof. We can first verify that

$$\mathbb{E}[[\mathbf{s}\mathbf{s}^T]_{ij}] = \mathbb{E}[s_i s_j] = \mathbb{E}[h(i)h(j)] = \delta_{ij} \implies \mathbb{E}[\mathbf{s}\mathbf{s}^T] = \mathbf{I}.$$

The variance of diagonal terms satisfy

$$\begin{aligned} \text{Var}[s_i^2] &= \mathbb{E}[s_i^4] - 1 \\ &= \mathbb{E}[h(i)^4] - 1 \\ &= 0. \end{aligned}$$

The variance of off-diagonal terms ($i \neq j$) satisfy

$$\begin{aligned} \text{Var}[s_i s_j] &= \mathbb{E}[(s_i s_j)^2] \\ &= \mathbb{E}[h(i)^2 h(j)^2] \\ &= 1. \end{aligned}$$

Thus, we can write $\text{Var}[s_i s_j] = (1 - \delta_{ij})$. Finally, the covariance (for $\{i, j\} \neq \{k, \ell\}$) can be evaluated as

$$\begin{aligned} \text{Cov}[s_i s_j, s_k s_\ell] &= \mathbb{E}[s_i s_j s_k s_\ell] - \mathbb{E}[s_i s_j] \mathbb{E}[s_k s_\ell] \\ &= \mathbb{E}[h(i)h(j)h(k)h(\ell)] - \delta_{ij}\delta_{k\ell} \\ &= \delta_{ij}\delta_{k\ell} - \delta_{ij}\delta_{k\ell} \\ &= 0. \end{aligned}$$

This means that overall, we have

$$\text{Cov}[s_i s_j, s_k s_\ell] = (\delta_{ik}\delta_{j\ell} + \delta_{i\ell}\delta_{jk})(1 - \delta_{ij}). \quad (5)$$

We can substitute this into the expressions in Lemma 4.1 to give

$$\begin{aligned} \text{Tr}[\Sigma] &= \sum_{ij} \left(\sum_{k\ell} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{\ell j}^2 + [\mathbf{AB}]_{ij}^2 - 2 \sum_{k\ell} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{kj}^2 \right) \\ &= \|\mathbf{A}\|_F^2 \|\mathbf{B}\|_F^2 + \|\mathbf{AB}\|_F^2 - 2 \sum_{ijk} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{kj}^2, \end{aligned}$$

⁽³⁾In fact, drawing entries randomly from a 4-wise independent hash family suffices [Sar06], though we will not exploit this.

and

$$\max_i [\Sigma]_{ii} = \max_{ij} \left(\sum_{k\ell} [A]_{ik}^2 [B]_{\ell j}^2 + [AB]_{ij}^2 - 2 \sum_{k\ell} [A]_{ik}^2 [B]_{k\ell}^2 \right) \quad (6)$$

$$= \|A\|_{2,\infty}^2 \|B^T\|_{2,\infty}^2 + \|AB\|_{\max}^2 - 2 \min_{ij} \sum_k [A]_{ik}^2 [B]_{kj}^2, \quad (7)$$

as required. $\square$

With this, we can write down complexities for approximate matrix multiplication using the tug-of-war sketch. This was detailed for Frobenius norm approximation in [Sar06], and we use the above to write the result for max-norm approximation also.

Theorem 4.3 (Classical tug-of-war algorithm, 2 matrices). *Let $A \in \mathbb{R}^{n \times q}$, $B \in \mathbb{R}^{q \times m}$. By sampling a single tug-of-war matrix of appropriate dimension, we can return a $C \in \mathbb{R}^{n \times m}$ in the circuit model with success probability at least $(1 - \delta)$, which satisfies*

$$\|C - AB\|_{\max} \leq \varepsilon \quad \text{in time} \quad O \left((mn + nq + qm) \left(\frac{\|A\|_{2,\infty}^2 \|B^T\|_{2,\infty}^2}{\varepsilon^2} \right)^{\omega-2} \log \left(\frac{1}{\delta} \right) \right), \quad (8)$$

$$\|C - AB\|_F \leq \varepsilon \quad \text{in time} \quad O \left((mn + nq + qm) \left(\frac{\|A\|_F^2 \|B\|_F^2}{\varepsilon^2} \right)^{\omega-2} \log \left(\frac{1}{\delta} \right) \right). \quad (9)$$

Proof. We can directly use the matrix Bernstein inequalities of Lemma 2.3, with expressions for covariance properties given by Lemma 4.2. The max-norm algorithm follows by choice of $c = L = 2\|A\|_{2,\infty}^2 \|B^T\|_{2,\infty}^2 \log(nm/\delta)/\varepsilon^2$ for the max-norm algorithm and $c = L = 2\|A\|_F^2 \|B\|_F^2 \log(1/\delta)/\varepsilon^2$ for the Frobenius norm algorithm. Finally, as discussed at the top of this section, the time complexity of computing $AS^T SB$ classically is $O((mn + nq + mq)c^{\omega-2})$. $\square$

4.2.2 Column-sample sketch

In the algorithm of Drineas, Kannan and Mahoney, for a matrix product of $A \in \mathbb{R}^{n \times q}$ with $B \in \mathbb{R}^{q \times m}$, $S \in \mathbb{R}^{c \times q}$ is chosen to be a row 1-sparse matrix with entries populated as $[S]_{i\bullet} = e_j^T / \sqrt{cp_j}$ with probability p_j , where we denote e_j as the unit column vector with non-zero j th entry, and for some probability distribution $\{p_k\}_k$ [DKM06]. This is their general framework, and in the end $\{p_k\}_k$ is judiciously chosen to minimize the convergence time for their algorithm in Frobenius norm approximation – note, however, that the minimizing distribution for the Frobenius norm does not optimize convergence in other norms, and indeed we find this not to be the case. In the mean estimation picture, we can write sketching vectors chosen as $s = e_j / \sqrt{p_j} \in \mathbb{R}^q$ with probability p_j .

Lemma 4.4 (Column-sample sketch covariance). *For the column-sample sketch as described above we have $\text{Cov}[s_i^2, s_j^2] = \frac{1}{p_i} \delta_{ij} - 1$ and $\text{Cov}[s_i s_j, s_k u_\ell] = 0$ for and $i \neq j$ or $k \neq \ell$. This implies that*

$$\begin{aligned} \text{Tr}[\Sigma] &= \sum_{i=1}^n \frac{1}{p_i} \|[A]_{\bullet,i}\|_2^2 \|[B]_{i,\bullet}\|_2^2 - \|AB\|_F^2, \\ \max_i [\Sigma]_{ii} &= \max_{ij} \left(\sum_k \frac{1}{p_k} [A]_{ik}^2 [B]_{kj}^2 - (AB)_{ij}^2 \right). \end{aligned}$$

Proof. We can first verify that

$$\mathbb{E}[ss^T] = \sum_{j=1}^n p_j \frac{e_j e_j^T}{p_j} = I.$$

We note that $\mathbf{s}\mathbf{s}^T$ is always diagonal, and thus $\text{Var}[s_i s_j] = \text{Var}[(\mathbf{s}\mathbf{s}^T)_{ij}] = 0$ for all $i \neq j$, and similarly for the covariance. The covariances of the diagonal terms are

$$\begin{aligned} \text{Cov}[s_i^2, s_j^2] &= \mathbb{E}[s_i^2 s_j^2] - \mathbb{E}[s_i^2] \mathbb{E}[s_j^2] \\ &= \sum_{k=1}^q p_k \frac{[\mathbf{e}_k \mathbf{e}_k^T]_{ii}}{p_k} \frac{[\mathbf{e}_k \mathbf{e}_k^T]_{jj}}{p_k} - \sum_{k=1}^q p_k \frac{[\mathbf{e}_k \mathbf{e}_k^T]_{ii}}{p_k} \sum_{\ell=1}^q p_\ell \frac{[\mathbf{e}_\ell \mathbf{e}_\ell^T]_{jj}}{p_\ell} \\ &= \sum_{k=1}^q \frac{1}{p_k} \delta_{ik} \delta_{jk} - \sum_{k=1}^q \delta_{ik} \sum_{\ell=1}^q \delta_{\ell j} \\ &= \frac{1}{p_i} \delta_{ij} - 1. \end{aligned}$$

We can now substitute this result into the expressions in Lemma 4.1 and write

$$\begin{aligned} \text{Tr}[\Sigma] &= \sum_{ij} \left(\sum_k \frac{1}{p_k} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{kj}^2 - \sum_{k\ell} [\mathbf{A}]_{ik} [\mathbf{B}]_{kj} [\mathbf{A}]_{i\ell} [\mathbf{B}]_{\ell j} \right) \\ &= \sum_{k=1}^n \frac{1}{p_k} \|\mathbf{A}_{\bullet, k}\|_2^2 \|\mathbf{B}_{k, \bullet}\|_2^2 - \|\mathbf{AB}\|_F^2, \end{aligned}$$

and

$$\begin{aligned} \max_i [\Sigma]_{ii} &= \max_{ij} \left(\sum_k \frac{1}{p_k} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{kj}^2 - \sum_{k\ell} [\mathbf{A}]_{ik} [\mathbf{B}]_{kj} [\mathbf{A}]_{i\ell} [\mathbf{B}]_{\ell j} \right) \\ &= \max_{ij} \left(\sum_k \frac{1}{p_k} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{kj}^2 - (\mathbf{AB})_{ij}^2 \right). \end{aligned}$$

which completes the proof. $\square$

The probability distribution can now be optimized – note that the optimal probability distribution for the Frobenius norm (as found in [DKM06]) may differ from the optimal distribution for the max-norm, and indeed we find this to be the case.

Theorem 4.5 (Classical column-sampling algorithm, 2 matrices). *Let $\mathbf{A} \in \mathbb{R}^{n \times q}$, $\mathbf{B} \in \mathbb{R}^{q \times m}$. By sampling a single column-sampling matrix of appropriate dimension, we can return a matrix $\mathbf{C} \in \mathbb{R}^{n \times m}$, with success probability at least $(1 - \delta)$, which satisfies*

$$\|\mathbf{C} - \mathbf{AB}\|_{\max} \leq \varepsilon \quad \text{in time} \quad O\left(nm \left(\frac{\|\mathbf{A}^T\|_{\infty, 2}^2 \|\mathbf{B}\|_{\infty, 2}^2}{\varepsilon^2}\right)^{\omega-2} \log\left(\frac{1}{\delta}\right)\right), \quad (10)$$

$$\|\mathbf{C} - \mathbf{AB}\|_F \leq \varepsilon \quad \text{in time} \quad O\left(nm \left(\frac{\|\mathbf{A}\|_F^2 \|\mathbf{B}\|_F^2}{\varepsilon^2}\right)^{\omega-2} \log\left(\frac{1}{\delta}\right)\right), \quad (11)$$

where ω is the matrix multiplication exponent.

Proof. Let us consider the Frobenius-norm algorithm first. Drineas et al. show that the minimizing probability distribution for $\text{Tr}[\Sigma]$ is the one with $p_k \propto \|\mathbf{A}_{\bullet, k}\|_2 \|\mathbf{B}_{k, \bullet}\|_2$, which gives

$$\text{Tr}[\Sigma] = \left(\sum_i \|\mathbf{A}_{\bullet, i}\|_2 \|\mathbf{B}_{i, \bullet}\|_2 \right)^2 - \|\mathbf{AB}\|_F^2.$$

The Frobenius norm algorithm follows by inspecting the matrix Bernstein inequalities (see Lemma 2.3) and taking $c = L = \left(\sum_i \|\mathbf{A}_{\bullet, i}\|_2 \|\mathbf{B}_{i, \bullet}\|_2 \right)^2 \leq \|\mathbf{A}\|_F^2 \|\mathbf{B}\|_F^2 / \varepsilon^2$. For the max-norm algorithm we

can write

$$\max_i [\Sigma]_{ii} \leq \sum_k \max_{ij} [A]_{ik}^2 [B]_{kj}^2 / p_k - \min_{ij} (AB)_{ij}^2.$$

By the method of Lagrange multipliers, or the Cauchy-Schwarz inequality, one can show that the minimizing probability distribution for this quantity is the one with probabilities $p_k \propto \max_{ij} [A]_{ik} [B]_{kj}$. For this choice of probability distribution we then have

$$\begin{aligned} \max_i [\Sigma]_{ii} &\leq \left(\sum_k \max_{ij} |[A]_{ik} [B]_{kj}| \right)^2 - \min_{ij} (AB)_{ij}^2 \\ &= \left(\sum_k \max_i |[A]_{ik}| \max_j |[B]_{kj}| \right)^2 - \min_{ij} (AB)_{ij}^2 \\ &\leq \left(\sum_k \max_i |[A]_{ik}|^2 \right) \left(\sum_\ell \max_j |[B]_{\ell j}|^2 \right) - \min_{ij} (AB)_{ij}^2 \\ &= \|A^T\|_{\infty,2}^2 \|B\|_{\infty,2}^2 - \min_{ij} (AB)_{ij}^2. \end{aligned} \tag{12}$$

where the second inequality is due to Cauchy-Schwarz. The max-norm algorithm follows by choice of $c = L = \|A\|_{\infty,2}^2 \|B^T\|_{\infty,2}^2 / \varepsilon^2$ and inspecting the max-norm matrix Bernstein inequality of Lemma 2.3.

Finally, we note that we can bring the complexity down with respect to dependencies on the matrix dimension, by observing that the computational cost of evaluating $AS^T SB$ is $O(nmc^{\omega-2})$ due to the sparsity of S (first evaluate AS^T and SB , taking time $O(nc)$ and $O(mc)$ respectively, and then evaluate their product using fast matrix multiplication), and there is no dependence on the inner dimension. $\square$

4.3 Quantum algorithm

In this section we quantize the tug-of-war and column-sample matrix sketches. Here we coherently construct $Ass^T B$ for given $s \in \mathbb{R}^q$ from the tug-of-war sketch (see Section 4.2.1) and use quantum multivariate mean estimation to estimate $\mathbb{E}[Ass^T B]$. As the algorithm is effectively computing matrix-vector products, it is unclear how to exploit fast matrix multiplication in this setting, unlike in the classical algorithms discussed in the previous section which can bunch together samples. Both algorithms have the same complexity for square matrices, but the quantum algorithm based on the column-sample sketch uses a QROM as the sketching vectors are constructed from column and row norms of the matrices. When considering rectangular matrices, ability to use QROM yields an advantage for the column-sample sketch over the tug-of-war sketch for matrix products with large inner dimension (A “wide”, B “tall”), as we can exploit the sparsity of the sketching vectors.

Theorem 4.6 (Quantum tug-of-war algorithm, 2 matrices). *Let $A \in \mathbb{R}^{n \times q}$, $B \in \mathbb{R}^{q \times m}$. We have a quantum algorithm which works in the circuit model and returns a matrix $C \in \mathbb{R}^{n \times m}$, with success probability at least $(1 - \delta)$, which satisfies*

$$\|C - AB\|_{\max} \leq \varepsilon \quad \text{in time} \quad \tilde{O}\left((nm + mq + nq) \left(\frac{\|A\|_{2,2} \|B\|_{2,2}}{\varepsilon}\right) \log\left(\frac{1}{\delta}\right)\right), \tag{13}$$

Proof. In order to perform quantum multivariate mean estimation (as outlined in Section 2.5), we use the following oracles. The action of the binary oracle may be written as

$$O_X : |s\rangle |0\rangle \mapsto |s\rangle |Ass^T B\rangle,$$

which can be obtained in $\tilde{O}(nm + mq + nq)$ time, given QROM access to A and B . Note that we will use the first register in superposition, but we do not need fast coherent read or write for the entries of A and B – thus with a binary encoding of both we can also achieve the same runtime in the circuit model,

without QROM. For the probability oracle, we relax the 4-wise independence assumption, and allow our event space to include all signed bitstrings $\Omega = \{-1, +1\}^n$. Here, the probability oracle has action:

$$U_{\mathbb{P}} : |0\rangle \xrightarrow{H^{\otimes n}} \frac{1}{\sqrt{2^n}} \sum_{2s' - \mathbf{1}_n \in \Omega} |s'\rangle .$$

which is instantiated simply by the Hadamard operation in constant time. Finally, we can write the action of the inner product oracle as

$$O_{Y, X} : |Y\rangle |s\rangle |0\rangle \mapsto |Y\rangle |s\rangle |\text{Tr}[Y^T A s s^T B]\rangle ,$$

for chosen instances $Y \in (-\frac{1}{2}, \frac{1}{2})^{n \times m}$ from the mean estimation algorithm of Theorem 2.5. We remark that the Hilbert-Schmidt (trace) inner product we denote above is entirely equivalent to the vector inner product, should Y and $A s s^T B$ be unfurled into column vectors. This inner product can be instantiated as

$$\begin{aligned} |Y\rangle |s\rangle |0\rangle |0\rangle &\xrightarrow{O_X} |Y\rangle |s\rangle |A s s^T B\rangle |0\rangle \\ &\longrightarrow |Y\rangle |s\rangle |A s s^T B\rangle |\text{Tr}[Y^T A s s^T B]\rangle \xrightarrow{O_X^\dagger} |Y\rangle |s\rangle |0\rangle |\text{Tr}[Y^T A s s^T B]\rangle , \end{aligned}$$

where the second step costs $O(nm)$. The quantum multivariate mean estimation routine of Theorem 2.5 can thus be instantiated with complexity $\tilde{O}((nm + nq + mq)N)$, attaining max-norm error $\sqrt{\text{Tr}[\Sigma]} \log(nm/\delta)/N$ with success probability at least $(1 - \delta)$, where $\text{Tr}[\Sigma]$ is given in Lemma 4.2. The theorem statement is satisfied by picking $N = \sqrt{2} \|A\|_F \|B\|_F \log(nm/\delta)/\varepsilon$. $\square$

Theorem 4.7 (Quantum column-sample algorithm, 2 matrices). *Let $A \in \mathbb{R}^{n \times q}$, $B \in \mathbb{R}^{q \times m}$. We have a quantum algorithm in the QROM model which returns a matrix $C \in \mathbb{R}^{n \times m}$, with success probability at least $(1 - \delta)$, which satisfies*

$$\|C - AB\|_{\max} \leq \varepsilon \quad \text{in time} \quad \tilde{O}\left(nm \left(\frac{\|A\|_{2,2} \|B\|_{2,2}}{\varepsilon}\right) \log\left(\frac{1}{\delta}\right)\right) . \quad (14)$$

Proof. The analysis follows in the same way as in the proof of Theorem 4.6, now with the column-sample sketch and equivalent bound for $\text{Tr}[\Sigma]$ given in Lemma 4.4. However, now s is constructed from a probability distribution based on column (row) norms of A (B). Thus, we require a QROM for those norms to instantiate the binary oracle $O_X : |s\rangle |0\rangle \mapsto |s\rangle |A s s^T B\rangle$. Compared to Theorem 4.6, we can improve the runtime to only depend on outer dimensions n, m by exploiting the sparsity of s in the column-sample sketch: As only costs $\tilde{O}(n)$ time to evaluate and $s^T B$ only costs $\tilde{O}(m)$ time to evaluate given QROM access to A and B , and their outer product costs $\tilde{O}(nm)$ time. The rest of the analysis follows as in the proof of Theorem 4.6, and the quantum multivariate mean estimation routine of Theorem 2.5 can thus be instantiated with complexity $\tilde{O}((nm)N)$ with the stated error guarantee attained using the same choice of N . $\square$

4.4 Multiple matrices

Now let's look at the general case of the product of k matrices which we label as $A_1 \cdots A_k$, with $A_\ell \in \mathbb{R}^{n_\ell \times n_{\ell+1}}$ for all $\ell \in [k]$. We approximate this by inserting $s_\ell s_\ell^T$ in between each matrix, that is we evaluate $(\prod_{\ell=1}^{k-1} A_\ell s_\ell s_\ell^T) A_k$. This product costs $O(k \max_{i \neq j} n_i n_j)$ to evaluate classically (or quantumly, using coherent arithmetic). As for 2 matrices, we will estimate the mean of this quantity, and can repeat the analysis as before, with both quantum and classical algorithms analyzed through the lens of mean estimation, if we can bound the relevant covariances.

Lemma 4.8 (Trace of covariance, sketched multiple matrix product). *Consider the matrix-valued random variable $(\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{s}_\ell \mathbf{s}_\ell^T) \mathbf{A}_k$ with each $\mathbf{s}_\ell \in \mathbb{R}^{n_\ell+1}$ drawn independently. We have*

$$\begin{aligned} \text{Tr}[\Sigma] &= \sum_{\substack{r_0=u_0, \\ q_k=t_k}} \sum_{\substack{(q_1, r_1), \dots, (q_{k-1}, r_{k-1}) \\ (t_1, u_1), \dots, (t_{k-1}, u_{k-1})}} \prod_{\ell=1}^k [\mathbf{A}_\ell]_{r_{\ell-1} q_\ell} [\mathbf{A}_\ell]_{u_{\ell-1} t_\ell} \cdot \\ &\quad \cdot \left(\prod_{\ell=1}^{k-1} \left(\text{Cov} \left[s_{q_\ell}^{(\ell)} s_{r_\ell}^{(\ell)}, s_{t_\ell}^{(\ell)} s_{u_\ell}^{(\ell)} \right] + \delta_{q_\ell r_\ell} \delta_{t_\ell u_\ell} \right) - \prod_{\ell=1}^{k-1} \delta_{q_\ell r_\ell} \delta_{t_\ell u_\ell} \right), \end{aligned} \quad (15)$$

where $s_j^{(\ell)}$ denotes the j th entry of $\mathbf{s}^{(\ell)}$.

Proof. We will use the fact that for a product of independent random variables we have that

$$\begin{aligned} \text{Var} \left[\prod_i^k X_i \right] &= \prod_i^k (\text{Var}[X_i] + \mathbb{E}[X_i]^2) - \prod_i^k \mathbb{E}[X_i]^2, \\ \text{Cov} \left[\prod_i^k X_i, \prod_i^k Y_i \right] &= \prod_i^k (\text{Cov}[X_i, Y_i] + \mathbb{E}[X_i] \mathbb{E}[Y_i]) - \prod_i^k \mathbb{E}[X_i] \mathbb{E}[Y_i] \end{aligned}$$

Using this, we can write

$$\begin{aligned} \text{Tr}[\Sigma] &= \sum_{ij} \text{Var} \left[\left(\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{s}_\ell \mathbf{s}_\ell^T \mathbf{A}_k \right)_{ij} \right] \\ &= \sum_{ij} \text{Var} \left[\sum_{(q_1, r_1), \dots, (q_{k-1}, r_{k-1})} \prod_{\ell=1}^{k-1} [\mathbf{A}_\ell]_{r_{\ell-1} q_\ell} s_{q_\ell}^{(\ell)} s_{r_\ell}^{(\ell)} [\mathbf{A}_k]_{r_{k-1} j} \right] \\ &= \sum_{\substack{r_0=u_0, \\ q_k=t_k}} \sum_{\substack{(q_1, r_1), \dots, (q_{k-1}, r_{k-1}) \\ (t_1, u_1), \dots, (t_{k-1}, u_{k-1})}} \prod_{\ell=1}^k [\mathbf{A}_\ell]_{r_{\ell-1} q_\ell} [\mathbf{A}_\ell]_{u_{\ell-1} t_\ell} \text{Cov} \left[\prod_{\ell=1}^{k-1} s_{q_\ell}^{(\ell)} s_{r_\ell}^{(\ell)}, \prod_{\ell'=1}^{k-1} s_{t_{\ell'}}^{(\ell')} s_{u_{\ell'}}^{(\ell')} \right] \\ &= \sum_{\substack{r_0=u_0, \\ q_k=t_k}} \sum_{\substack{(q_1, r_1), \dots, (q_{k-1}, r_{k-1}) \\ (t_1, u_1), \dots, (t_{k-1}, u_{k-1})}} \prod_{\ell=1}^k [\mathbf{A}_\ell]_{r_{\ell-1} q_\ell} [\mathbf{A}_\ell]_{u_{\ell-1} t_\ell} \cdot \\ &\quad \cdot \left(\prod_{\ell=1}^{k-1} \left(\text{Cov} \left[s_{q_\ell}^{(\ell)} s_{r_\ell}^{(\ell)}, s_{t_\ell}^{(\ell)} s_{u_\ell}^{(\ell)} \right] + \delta_{q_\ell r_\ell} \delta_{t_\ell u_\ell} \right) - \prod_{\ell=1}^{k-1} \delta_{q_\ell r_\ell} \delta_{t_\ell u_\ell} \right), \end{aligned}$$

where in the third line we have assigned the indices $r_0 = u_0 = i, q_k = t_k = j$. $\square$

Lemma 4.9 (Tug-of-war variances, multiple matrices). *Adopt the setting of Lemma 4.8, using tug-of-war sketch matrices. We have*

$$\text{Tr}[\Sigma] \leq 3^k \|\mathbf{A}_1\|_{2,2}^2 \cdots \|\mathbf{A}_k\|_{2,2}^2.$$

Proof. From Lemma 4.2 recall that for the tug-of-war sketch we have $\text{Cov} [s_i s_j, s_k s_\ell] = (\delta_{ik} \delta_{jl} + \delta_{il} \delta_{jk})(1 - \delta_{ij})$. As these correspond only to variances, they are always positive (if we upper bound all coefficients by their positive counterparts) and we can write $\text{Cov} [s_i s_j, s_k s_\ell] \leq (\delta_{ik} \delta_{jl} + \delta_{il} \delta_{jk})$. This

can be substituted into the expression for $\text{Tr}[\Sigma]$ in Lemma 4.8, giving

$$\begin{aligned}
\text{Tr}[\Sigma] &\leq \sum_{\substack{r_0=u_0, \\ q_k=t_k}} \sum_{\substack{(q_1, r_1), \dots, (q_{k-1}, r_{k-1}) \\ (t_1, u_1), \dots, (t_{k-1}, u_{k-1})}} \prod_{\ell=1}^k [\overline{\mathbf{A}_\ell}]_{r_{\ell-1}q_\ell} [\overline{\mathbf{A}_\ell}]_{u_{\ell-1}t_\ell} \cdot \prod_{\ell'=1}^{k-1} (\delta_{q_{\ell'}t_{\ell'}} \delta_{r_{\ell'}u_{\ell'}} + \delta_{q_{\ell'}u_{\ell'}} \delta_{r_{\ell'}t_{\ell'}} + \delta_{q_{\ell'}r_{\ell'}} \delta_{t_{\ell'}u_{\ell'}}) \\
&= \sum_{r_0, q_k} \sum_{(q_1, r_1), \dots, (q_{k-1}, r_{k-1})} ([\overline{\mathbf{A}_1}]_{r_0q_1}^2 + [\overline{\mathbf{A}_1}]_{r_0q_1} [\overline{\mathbf{A}_1}]_{r_0r_1} + [\overline{\mathbf{A}_1}]_{r_0q_1} [\overline{\mathbf{A}_1}]_{r_0r_1}) \cdot \\
&\quad \cdot \prod_{\ell=2}^{k-1} ([\overline{\mathbf{A}_\ell}]_{r_{\ell-1}q_\ell}^2 + [\overline{\mathbf{A}_\ell}]_{r_{\ell-1}q_\ell} [\overline{\mathbf{A}_\ell}]_{q_{\ell-1}r_\ell} + [\overline{\mathbf{A}_\ell}]_{q_{\ell-1}q_\ell} [\overline{\mathbf{A}_\ell}]_{r_{\ell-1}r_\ell}) \cdot \\
&\quad \cdot ([\overline{\mathbf{A}_k}]_{r_{k-1}q_k}^2 + [\overline{\mathbf{A}_k}]_{r_{\ell-1}q_k} [\overline{\mathbf{A}_k}]_{q_{k-1}q_k} + [\overline{\mathbf{A}_k}]_{q_{k-1}q_k} [\overline{\mathbf{A}_k}]_{r_{\ell-1}q_k}) \\
&\leq \sum_{r_0, q_k} \sum_{(q_1, r_1), \dots, (q_{k-1}, r_{k-1})} ([\overline{\mathbf{A}_1}]_{r_0q_1}^2 + [\overline{\mathbf{A}_1}]_{r_0q_1} [\overline{\mathbf{A}_1}]_{r_0r_1} + [\overline{\mathbf{A}_1}]_{r_0q_1} [\overline{\mathbf{A}_1}]_{r_0r_1}) \cdot \\
&\quad \cdot \prod_{\ell=2}^{k-1} ([\overline{\mathbf{A}_\ell}]_{r_{\ell-1}q_\ell}^2 + [\overline{\mathbf{A}_\ell}]_{r_{\ell-1}q_\ell} [\overline{\mathbf{A}_\ell}]_{q_{\ell-1}r_\ell} + [\overline{\mathbf{A}_\ell}]_{q_{\ell-1}q_\ell} [\overline{\mathbf{A}_\ell}]_{r_{\ell-1}r_\ell}) \cdot \\
&\quad \cdot ([\overline{\mathbf{A}_k}]_{r_{k-1}q_k}^2 + [\overline{\mathbf{A}_k}]_{r_{\ell-1}q_k} [\overline{\mathbf{A}_k}]_{q_{k-1}q_k} + [\overline{\mathbf{A}_k}]_{q_{k-1}q_k} [\overline{\mathbf{A}_k}]_{r_{\ell-1}q_k}) \\
&\leq \sum_{r_0, q_k} \sum_{\substack{(q_1, r_1), \dots, (q_{k-1}, r_{k-1}) \\ (q'_1, r'_1), \dots, (q'_{k-1}, r'_{k-1})}} ([\overline{\mathbf{A}_1}]_{r_0q'_1}^2 + [\overline{\mathbf{A}_1}]_{r_0q'_1} [\overline{\mathbf{A}_1}]_{r_0r'_1} + [\overline{\mathbf{A}_1}]_{r_0q'_1} [\overline{\mathbf{A}_1}]_{r_0r'_1}) \cdot \\
&\quad \cdot \prod_{\ell=2}^{k-1} ([\overline{\mathbf{A}_\ell}]_{r_{\ell-1}q'_\ell}^2 + [\overline{\mathbf{A}_\ell}]_{r_{\ell-1}q'_\ell} [\overline{\mathbf{A}_\ell}]_{q_{\ell-1}r'_\ell} + [\overline{\mathbf{A}_\ell}]_{q_{\ell-1}q'_\ell} [\overline{\mathbf{A}_\ell}]_{r_{\ell-1}r'_\ell}) \cdot \\
&\quad \cdot ([\overline{\mathbf{A}_k}]_{r_{k-1}q_k}^2 + [\overline{\mathbf{A}_k}]_{r_{\ell-1}q_k} [\overline{\mathbf{A}_k}]_{q_{k-1}q_k} + [\overline{\mathbf{A}_k}]_{q_{k-1}q_k} [\overline{\mathbf{A}_k}]_{r_{\ell-1}q_k})
\end{aligned}$$

where in the second line we have done a relabeling in the third term in each bracket $u_\ell \rightarrow r_\ell$, and in the final line we have decoupled sums using the fact that all terms are positive. We note that the expression in the final line is a sum of products of squared Frobenius norms of different groupings of matrices from $\mathbf{A}_1$ to $\mathbf{A}_k$. Using the submultiplicativity of Frobenius norms, and by counting terms, we have

$$\text{Tr}[\Sigma] \leq 3^k \|\mathbf{A}_1\|_F \cdots \|\mathbf{A}_k\|_F,$$

as required. $\square$

Theorem 4.10 (Sketching algorithms for multiple matrices). *Consider a matrix product of k matrices $\mathbf{A}_1 \cdots \mathbf{A}_k$, with $\mathbf{A}_\ell \in \mathbb{R}^{n_{\ell-1} \times n_\ell}$ for all $\ell \in [k]$. With classical preprocessing time linear in the size of the input up to polylogarithmic factors, we give:*

A classical algorithm which gives

$$\|\mathbf{C} - \mathbf{AB}\|_F \leq \varepsilon \quad \text{in time} \quad \tilde{O} \left(3^k \cdot \max_{i \neq j} n_i n_j \cdot \frac{\|\mathbf{A}_1\|_{2,2}^2 \cdots \|\mathbf{A}_k\|_{2,2}^2}{\varepsilon^2} \log \left(\frac{1}{\delta} \right) \right), \quad (16)$$

and a quantum algorithm which gives

$$\|\mathbf{C} - \mathbf{AB}\|_{\max} \leq \varepsilon \quad \text{in time} \quad \tilde{O} \left(3^{k/2} \cdot \max_{i \neq j} n_i n_j \cdot \frac{\|\mathbf{A}_1\|_{2,2} \cdots \|\mathbf{A}_k\|_{2,2}}{\varepsilon} \log \left(\frac{1}{\delta} \right) \right), \quad (17)$$

both with success probability at least $(1 - \delta)$.

Proof. The sample complexity can be found by simply using the bound on $\text{Tr}[\Sigma]$ derived in Lemma

4.8 for the classical matrix Bernstein inequality quoted in Lemma 2.3 and the quantum mean estimation subroutine of Theorem 2.5, with a choice of $L = 3^k \|\mathbf{A}_1\|_F^2 \cdots \|\mathbf{A}_k\|_F^2 \log(1/\delta)/\varepsilon^2$ and $N = 3^{k/2} \|\mathbf{A}_1\| \cdots \|\mathbf{A}_k\|_F \log(n_0 n_k / \delta) / \varepsilon$ respectively. The time complexity per sample costs $O(k \max_{i \neq j} n_i n_j)$ for both the classical and quantum algorithm. $\square$

So far we have not exploited fast matrix multiplication. We can instead insert $\mathbf{S}_\ell^T \mathbf{S}_\ell$ in between each matrix. That is, we approximate $\mathbf{A}_1 \cdots \mathbf{A}_k$ by evaluating $\left(\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{S}_\ell^T \mathbf{S}_\ell \right) \mathbf{A}_k$. Note that unlike in the case of a product of 2 matrices, there is no longer a simple interpretation of this as a sample mean over choices of $\mathbf{s}_\ell \mathbf{s}_\ell^T$. Thus, now we need to analyze $\text{Tr}[\Sigma]$ for the matrix-valued random variable $\left(\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{S}_\ell^T \mathbf{S}_\ell \right) \mathbf{A}_k$ explicitly.

Theorem 4.11 (Sketching algorithm for multiple matrices exploiting fast matrix multiplication). *Consider a matrix product of k matrices $\mathbf{A}_1 \cdots \mathbf{A}_k$, with $\mathbf{A}_\ell \in \mathbb{R}^{n_{\ell-1} \times n_\ell}$ for all $\ell \in [k]$. With classical preprocessing time linear in the size of the input up to polylogarithmic factors, we give a classical algorithm which returns $\mathbf{C}$ satisfying*

$$\|\mathbf{C} - \mathbf{A}\mathbf{B}\|_F \leq \varepsilon \quad \text{in time} \quad \tilde{O} \left(3^{k(\omega-2)} \cdot \max_{i \neq j} n_i n_j \left(\frac{\|\mathbf{A}_1\|_{2,2}^2 \cdots \|\mathbf{A}_k\|_{2,2}^2}{\varepsilon^2} \right)^{\omega-2} \log \left(\frac{1}{\delta} \right) \right), \quad (18)$$

with success probability at least $(1 - \delta)$.

Proof. We need to modify our prior analysis for the new random variable $\left(\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{S}_\ell^T \mathbf{S}_\ell \right) \mathbf{A}_k$. We can check that the trace of covariance satisfies

$$\begin{aligned} \text{Tr}[\Sigma] &= \sum_{ij} \text{Var} \left[\left[\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{S}_\ell^T \mathbf{S}_\ell \mathbf{A}_k \right]_{ij} \right] \\ &= \sum_{ij} \mathbb{E} \left[\left(\left[\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{S}_\ell^T \mathbf{S}_\ell \mathbf{A}_k \right]_{ij} - \left[\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \right]_{ij} \right)^2 \right] \\ &= \mathbb{E} \left[\sum_{ij} \left[\prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{S}_\ell^T \mathbf{S}_\ell \mathbf{A}_k - \prod_{\ell=1}^{k-1} \mathbf{A}_\ell \right]_{ij}^2 \right] \\ &= \mathbb{E} \left[\left\| \prod_{\ell=1}^{k-1} \mathbf{A}_\ell \mathbf{S}_\ell^T \mathbf{S}_\ell \mathbf{A}_k - \prod_{\ell=1}^{k-1} \mathbf{A}_\ell \right\|_F^2 \right], \end{aligned}$$

which shows that it controls the expected (squared) error. Second, a modification of Lemma 4.8 yields

$$\begin{aligned} \text{Tr}[\Sigma] &= \sum_{\substack{r_0=u_0, \\ q_k=t_k}} \sum_{(q_1, r_1), \dots, (q_{k-1}, r_{k-1})} \prod_{\ell=1}^k [\mathbf{A}_\ell]_{r_{\ell-1} q_\ell} [\mathbf{A}_\ell]_{u_{\ell-1} t_\ell} \\ &\quad \cdot \left(\prod_{\ell=1}^{k-1} \left(\text{Cov} \left[[\mathbf{S}_\ell^T \mathbf{S}_\ell]_{q_\ell, r_\ell}, [\mathbf{S}_\ell^T \mathbf{S}_\ell]_{t_\ell, u_\ell} \right] + \delta_{q_\ell r_\ell} \delta_{t_\ell u_\ell} \right) - \prod_{\ell=1}^{k-1} \delta_{q_\ell r_\ell} \delta_{t_\ell u_\ell} \right), \end{aligned}$$

and so similar to before, we can the approximation error by determining the quantities $\text{Cov} \left[[\mathbf{S}_\ell^T \mathbf{S}_\ell]_{q_\ell, r_\ell}, [\mathbf{S}_\ell^T \mathbf{S}_\ell]_{t_\ell, u_\ell} \right]$. Recall that $\mathbb{E}[\mathbf{S}^T \mathbf{S}] = \mathbf{I}$. Similar to in Lemma 4.2 we have

$$\begin{aligned} \text{Var} \left[[\mathbf{S}^T \mathbf{S}]_{ii} \right] &= \mathbb{E} \left[[\mathbf{S}^T \mathbf{S}]_{ii}^2 \right] - 1 \\ &= \frac{1}{c^2} \mathbb{E} \left[\sum_{k\ell} h_k(i)^2 h_\ell(i)^2 \right] - 1 \end{aligned}$$

$$\begin{aligned}
&= \sum_{k\ell} \frac{1}{c^2} - 1 \\
&= 0.
\end{aligned}$$

The variance of off-diagonal terms ($i \neq j$) satisfy

$$\begin{aligned}
\text{Var} [[\mathbf{S}^T \mathbf{S}]_{ij}] &= \mathbb{E}[(\mathbf{S}^T \mathbf{S})_{ij}^2] - 0 \\
&= \frac{1}{c^2} \mathbb{E} \left[\sum_{k\ell} h_k(i) h_k(j) h_\ell(i) h_\ell(j) \right] \\
&= \frac{1}{c^2} \sum_k h_k(i) h_k(j) h_k(i) h_k(j) \\
&= \frac{1}{c}.
\end{aligned}$$

Thus, we can write $\text{Var} [[\mathbf{S}^T \mathbf{S}]_{ij}] = \frac{1}{c}(1 - \delta_{ij})$. Finally, the covariance (for $(i, j) \neq (k, \ell)$) can be evaluated as

$$\text{Cov} [[\mathbf{S}^T \mathbf{S}]_{ij}, [\mathbf{S}^T \mathbf{S}]_{k\ell}] = \mathbb{E}[(\mathbf{S}^T \mathbf{S})_{ij}(\mathbf{S}^T \mathbf{S})_{k\ell}] - \mathbb{E}[[\mathbf{S}^T \mathbf{S}]_{ij}] \mathbb{E}[[\mathbf{S}^T \mathbf{S}]_{k\ell}] \quad (19)$$

$$= \frac{1}{c^2} \mathbb{E} \left[\sum_{qr} h_q(i) h_q(j) h_r(k) h_r(\ell) \right] - \delta_{ij} \delta_{k\ell} \quad (20)$$

$$= \delta_{ij} \delta_{k\ell} + \frac{1}{c} \delta_{i\ell} \delta_{jk} - \delta_{ij} \delta_{k\ell} \quad (21)$$

$$= \frac{1}{c} \delta_{i\ell} \delta_{jk}, \quad (22)$$

So that overall, we have $\text{Cov} [[\mathbf{S}^T \mathbf{S}]_{ij}, [\mathbf{S}^T \mathbf{S}]_{k\ell}] = \frac{1}{c}(\delta_{i\ell} \delta_{jk} + \delta_{ik} \delta_{j\ell})(1 - \delta_{ij})$, which is an analogous result for the covariances to Lemma 4.9, now with a factor $\frac{1}{c}$. We can retrace the proof of Lemma 4.8 to obtain the bound

$$\text{Tr}[\Sigma] \leq \frac{3^k \|\mathbf{A}_1\|_F \cdots \|\mathbf{A}_k\|_F}{c}.$$

Now we observe that the classical matrix product costs $O(k \max_{i \neq j} n_i n_j c^{\omega-2})$ time to evaluate, where we take $c = 3^k \|\mathbf{A}_1\|_{2,2}^2 \cdots \|\mathbf{A}_k\|_{2,2}^2 / \varepsilon^2$ to attain the desired approximation guarantee of ε . The success probability is controlled by Markov's inequality, and we can attain any arbitrary constant success probability by increasing c by a constant factor. By performing $O(\log(1/\delta))$ experiments, we can attain success probability at least $(1 - \delta)$ by using a median-of-means strategy. $\square$

We can substitute this into the expressions in Lemma 4.1 to give

$$\begin{aligned}
\text{Tr}[\Sigma] &= \sum_{ij} \left(\sum_{k\ell} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{\ell j}^2 + [\mathbf{AB}]_{ij}^2 - 2 \sum_{k\ell} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{kj}^2 \right) \\
&= \|\mathbf{A}\|_F^2 \|\mathbf{B}\|_F^2 + \|\mathbf{AB}\|_F^2 - 2 \sum_{ijk} [\mathbf{A}]_{ik}^2 [\mathbf{B}]_{kj}^2.
\end{aligned}$$

The variance of off-diagonal terms ($i \neq j$) satisfy

$$\begin{aligned}
\text{Var} [s_i s_j] &= \mathbb{E}[(s_i s_j)^2] \\
&= \mathbb{E} [h(i)^2 h(j)^2] \\
&= 1.
\end{aligned}$$

Thus, we can write $\text{Var}[s_i s_j] = (1 - \delta_{ij})$. Finally, the covariance (for $\{i, j\} \neq \{k, \ell\}$) can be evaluated as

$$\begin{aligned} \text{Cov}[s_i s_j, s_k s_\ell] &= \mathbb{E}[s_i s_j s_k s_\ell] - \mathbb{E}[s_i s_j] \mathbb{E}[s_k s_\ell] \\ &= \mathbb{E}[h(i)h(j)h(k)h(\ell)] - \delta_{ij}\delta_{k\ell} \\ &= \delta_{ij}\delta_{k\ell} - \delta_{ij}\delta_{k\ell} \\ &= 0. \end{aligned}$$

This means that overall, we have

$$\text{Cov}[s_i s_j, s_k s_\ell] = (\delta_{ik}\delta_{j\ell} + \delta_{i\ell}\delta_{jk})(1 - \delta_{ij}). \quad (23)$$

We can substitute this into the expressions in Lemma 4.1 to give

$$\begin{aligned} \text{Tr}[\Sigma] &= \sum_{ij} \left(\sum_{k\ell} [A]_{ik}^2 [B]_{\ell j}^2 + [AB]_{ij}^2 - 2 \sum_{k\ell} [A]_{ik}^2 [B]_{k\ell}^2 \right) \\ &= \|A\|_F^2 \|B\|_F^2 + \|AB\|_F^2 - 2 \sum_{ijk} [A]_{ik}^2 [B]_{kj}^2, \end{aligned}$$

5 Complexity comparison

In this section we give a more detailed complexity comparison, as well as compare our time complexity results against prior work whilst relaxing assumptions about square matrices or inability to perform fast matrix multiplication. In Table 5 we present a comparison of time complexities for rectangular matrices, where the matrix multiplication exponent is set to ω . We see that exploiting $\omega < 3$ can lead to greatly improved complexities for sketching-based algorithms.

In order to compare the expressions, we recall the following sequence of inequalities that follow directly from Hölder's inequality (see Lemma 2.2):

$$\frac{1}{\sqrt{n}} \|X\|_{1,1} \leq \left\{ \begin{array}{l} \|X\|_{1,2} \\ \|X^T\|_{1,2} \\ \|X^T\|_{2,1} \\ \|X\|_{2,1} \end{array} \right\} \leq \sqrt{n} \|X\|_{2,2} = \sqrt{n} \|X^T\|_{2,2} \leq \left\{ \begin{array}{l} n \|X\|_{2,\infty} \\ n \|X^T\|_{2,\infty} \\ n \|X^T\|_{\infty,2} \\ n \|X\|_{\infty,2} \end{array} \right\}.$$

We prove one more inequality here which is stated in Figure 1.

Lemma 5.1. *Consider $A \in \mathbb{C}^{n \times k}$ and $B \in \mathbb{C}^{k \times m}$, and denote $\bar{A}$ as the matrix with element-wise absolute values. We have*

$$\|\bar{A}\bar{B}\|_{1,1} \leq \|A\|_{2,1} \|B^T\|_{2,1}. \quad (24)$$

Proof. The proof follows from Cauchy-Schwarz, as

$$\|\bar{A}\bar{B}\|_{1,1} = \sum_{j,k,\ell} |[A]_{j\ell}| \cdot |[B]_{\ell k}| \leq \sum_{j,k} \|[A]_{j,\bullet}\|_2 \cdot \|[B]_{\bullet,k}\|_2 = \|A\|_{2,1} \|B^T\|_{2,1}. \quad \square$$

At a high level, it seems that there is a trade-off between the strength of the computational model, and the run-time of the algorithms for achieving the same precision. Indeed, the random-walk-based algorithms, building on the ideas of Cohen & Lewis, achieve the best run-times, but at the same time they require the strongest computational model assumptions (RAM classically, and QRAM quantumly). On the other hand, the sketching-based algorithms have a somewhat worse run-time, especially if we do not consider fast matrix-multiplication algorithms (i.e. $\omega = 3$). However, the assumptions on the computational model are a lot weaker, since we only need QROM-access quantumly for the approach based on Drineas et al., and for the approach based on Sarlós we can even implement it in the circuit model, both classically and quantumly.

		Runtime ($\tilde{O}(\cdot)$)		Data model	Multiple matrices?
		$\ \cdot\ _{\max}$	$\ \cdot\ _F$		
Classical	[CL99]	$\frac{\ \bar{\mathbf{A}}\bar{\mathbf{B}}\ _{1,2}^2}{\varepsilon^2}$	$\sqrt{nm} \frac{\ \bar{\mathbf{A}}\bar{\mathbf{B}}\ _{1,2}^2}{\varepsilon^2}$	RAM	✓
	[Sar06]	$n_* n'_* \left(\frac{\ \mathbf{A}\ _{2,\infty}^2 \ \mathbf{B}^T\ _{2,\infty}^2}{\varepsilon^2} \right)^{\omega-2}$	$n_* n'_* \left(\frac{\ \mathbf{A}\ _{2,2}^2 \ \mathbf{B}\ _{2,2}^2}{\varepsilon^2} \right)^{\omega-2}$	circuit	✓ (Thm. 4.11)
	[DKM06]	$n_* n'_* \left(\frac{\ \mathbf{A}^T\ _{\infty,2}^2 \ \mathbf{B}\ _{\infty,2}^2}{\varepsilon^2} \right)^{\omega-2}$	$n_* n'_* \left(\frac{\ \mathbf{A}\ _{2,2}^2 \ \mathbf{B}\ _{2,2}^2}{\varepsilon^2} \right)^{\omega-2}$	circuit	
	[DKM06]	$nm \left(\frac{\ \mathbf{A}^T\ _{\infty,2}^2 \ \mathbf{B}\ _{\infty,2}^2}{\varepsilon^2} \right)^{\omega-2}$	$nm \left(\frac{\ \mathbf{A}\ _{2,2}^2 \ \mathbf{B}\ _{2,2}^2}{\varepsilon^2} \right)^{\omega-2}$	RAM	
	Thm. 3.4	$\frac{\ \bar{\mathbf{A}}\bar{\mathbf{B}}\ _{\infty,\infty} \ \bar{\mathbf{A}}\bar{\mathbf{B}}\ _{1,1}}{\varepsilon^2}$	$\frac{\ \bar{\mathbf{A}}\bar{\mathbf{B}}\ _{1,1}^2}{\varepsilon^2}$	RAM	✓
Quantum	[Sha18]	$\frac{\ \mathbf{A}\ _{2,1} \ \mathbf{B}^T\ _{2,1}}{\varepsilon}$	$\sqrt{nm} \frac{\ \mathbf{A}^T\ _{2,1} \ \mathbf{B}\ _{2,1}}{\varepsilon}$	QROM	
	Thms. 4.6 & 4.10	$n_* n'_* \frac{\ \mathbf{A}\ _{2,2} \ \mathbf{B}\ _{2,2}}{\varepsilon}$		quantum circuit	✓
	Thm. 4.7	$nm \frac{\ \mathbf{A}\ _{2,2} \ \mathbf{B}\ _{2,2}}{\varepsilon}$		QROM	
	Thm. 3.5	$\frac{\ \bar{\mathbf{A}}\bar{\mathbf{B}}\ _{1,1}}{\varepsilon}$	$\sqrt{nm} \frac{\ \bar{\mathbf{A}}\bar{\mathbf{B}}\ _{1,1}}{\varepsilon}$	QRAM	✓

Table 2: **Detailed time complexity comparison.** Here we consider multiplication of two rectangular matrices $\mathbf{A} \in \mathbb{R}^{n \times q}$, $\mathbf{B} \in \mathbb{R}^{q \times m}$. We denote $n_* n'_* = \max(nm, nq, qm)$ as the largest pair of dimensions. In all constructions that require memory usage (i.e., either RAM, QROM, or QRAM), we require a classical preprocessing routine that initializes the memory (and hence runs in the RAM-model), and runs in time linear in the size of the input.

The situation becomes slightly more complicated when we additionally take fast matrix-multiplication algorithms into account, i.e., when $\omega < 3$. It seems that there is little hope to beat any classical sketching-based algorithm with these quantum sketching techniques, if we take $\omega < 2.5$. However, the algorithms that achieve a matrix multiplication constant below 2.5 are very complicated and introduce very high constant factors, so we believe there is still utility in designing simple classical and quantum algorithms that behave slightly worse asymptotically than the classical state-of-the-art, but might beat most straightforward classical implementations.

Acknowledgements

SA was supported in part by the European QuantERA project QOPT (ERA-NET Cofund 2022-25), the French PEPR integrated projects EPiQ (ANR-22-PETQ0007) and HQI (ANR-22-PNCQ-0002), and the French ANR project QUOPS (ANR-22-CE47-0003-01). AC is supported by a Simons-CIQC postdoctoral fellowship through NSF QLCI Grant No. 2016245. SW acknowledges funding provided by the Institute for Quantum Information and Matter, an NSF Physics Frontiers Center (NSF Grant PHY-23171100).

References

- [AG23] S. Apers and S. Gribling. “Quantum speedups for linear programming via interior point methods”. arXiv:2311.03215. 2023 (Page 9).
- [All+23] J. Allcock, J. Bao, A. Belovs, T. Lee, and M. Santha. “On the quantum time complexity of divide and conquer”. arXiv:2311.16401. 2023 (Page 8).

- [Alm+25] J. Alman, R. Duan, V. V. Williams, Y. Xu, Z. Xu, and R. Zhou. “More asymmetry yields faster matrix multiplication”. In: *Proceedings of the 2025 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*. arXiv:2404.16349. SIAM. 2025, pp. 2005–2039 (Page 1).
- [AMS96] N. Alon, Y. Matias, and M. Szegedy. “The space complexity of approximating the frequency moments”. In: *Proceedings of the 28th Annual ACM Symposium on Theory of Computing (STOC)*. 1996, pp. 20–29 (Pages 5, 15).
- [Ape20] J. van Apeldoorn. “A Quantum View on Convex Optimization”. PhD thesis. Universiteit van Amsterdam, 2020 (Page 8).
- [Ber+24] A. Bernasconi, A. Berti, G. M. Del Corso, and A. Poggiali. “Quantum subroutine for efficient matrix multiplication”. In: *IEEE Access* (2024) (Page 2).
- [CHJ22] A. Cornelissen, Y. Hamoudi, and S. Jerbi. “Near-optimal quantum algorithms for multivariate mean estimation”. In: *Proceedings of the 54th ACM Symposium on the Theory of Computing (STOC)*. arXiv:2111.09787. 2022, pp. 33–43 (Pages 2, 6, 8, 9, 12).
- [CL99] E. Cohen and D. D. Lewis. “Approximating matrix multiplication for pattern recognition tasks”. In: *Journal of Algorithms* 30.2 (1999). Earlier version in *SODA’97*, pp. 211–252 (Pages 1–4, 9–12, 25).
- [Cor25] A. Cornelissen. “Quantum algorithms through graph composition”. arXiv:2504.02115. 2025 (Page 8).
- [CW90] D. Coppersmith and S. Winograd. “Matrix multiplication via arithmetic progressions”. In: *Journal of Symbolic Computation* 9.3 (1990). Computational algebraic complexity editorial, pp. 251–280. ISSN: 0747-7171 (Page 6).
- [Dal+25] A. M. Dalzell, A. Gilyén, C. T. Hann, S. McArdle, G. Salton, Q. T. Nguyen, A. Kubica, and F. G. Brandão. “A distillation-teleportation protocol for fault-tolerant QRAM”. arXiv:2505.20265. 2025 (Page 8).
- [DKM06] P. Drineas, R. Kannan, and M. W. Mahoney. “Fast Monte Carlo algorithms for matrices I: Approximating matrix multiplication”. In: *SIAM Journal on Computing* 36.1 (2006). Open access version on *M. W. Mahoney’s webpage*, pp. 132–157 (Pages 1–3, 5, 6, 13, 16, 17, 25).
- [Gro11] D. Gross. “Recovering low-rank matrices from few coefficients in any basis”. In: *IEEE Transactions on Information Theory* 57.3 (2011). arXiv:0910.1879, pp. 1548–1566 (Page 7).
- [Ili89] C. S. Iliopoulos. “Worst-case complexity bounds on algorithms for computing the canonical structure of finite abelian groups and the Hermite and Smith normal forms of an integer matrix”. In: *SIAM Journal on Computing* 18.4 (1989), pp. 658–669 (Page 6).
- [JR23] S. Jaques and A. G. Rattew. “QRAM: A survey and critique”. arXiv:2305.10310. 2023 (Page 8).
- [Kup13] G. Kuperberg. “Another Subexponential-time Quantum Algorithm for the Dihedral Hidden Subgroup Problem”. In: *8th Conference on the Theory of Quantum Computation, Communication and Cryptography*. arXiv:1112.3333. 2013, p. 20 (Page 8).
- [Mon15] A. Montanaro. “Quantum speedup of Monte Carlo methods”. In: *Proc. R. Soc. A* 471.2181 (2015). arXiv:1504.06987 (Page 13).
- [NS20] M. Naya-Plasencia and A. Schrottenloher. “Optimal merging in quantum k-xor and k-sum algorithms”. In: *Annual International Conference on the Theory and Applications of Cryptographic Techniques*. Cryptography ePrint Archive: 2019/501. Springer. 2020, pp. 311–340 (Page 8).
- [Pag13] R. Pagh. “Compressed matrix multiplication”. In: *ACM Transactions on Computation Theory (TOCT)* 5.3 (2013). arXiv:1108.1320, pp. 1–17 (Page 6).
- [Pra14] A. Prakash. “Quantum Algorithms for Linear Algebra and Machine Learning”. PhD thesis. University of California at Berkeley, 2014 (Page 2).

- [Sar06] T. Sarlos. “Improved approximation algorithms for large matrices via random projections”. In: *Proceedings of the 47th IEEE Symposium on Foundations of Computer Science (FOCS)*. Open access version on *CiteSeerX*. IEEE. 2006, pp. 143–152 (Pages 1–3, 5, 6, 13, 15, 16, 25).
- [Sha18] C. Shao. “Quantum algorithms to matrix multiplication”. arXiv:1803.01601. 2018 (Pages 2, 3, 6, 25).
- [Sip96] M. Sipser. *Introduction to the Theory of Computation*. Cengage Learning, 1996 (Page 8).
- [Str69] V. Strassen. “Gaussian elimination is not optimal”. In: *Numerische mathematik* 13.4 (1969), pp. 354–356 (Page 6).
- [Tan25] L. Tang. “More-efficient Quantum Multivariate Mean Value Estimator from Generalized Grover Gate”. arXiv:2504.06940. 2025 (Page 8).
- [Tro15] J. A. Tropp. “An introduction to matrix concentration inequalities”. In: *Foundations and Trends® in Machine Learning* 8.1-2 (2015). arXiv:1501.01571, pp. 1–230 (Page 7).
- [Vos91] M. D. Vose. “A linear algorithm for generating random numbers with a given distribution”. In: *IEEE Transactions on Software Engineering* 17.9 (1991), p. 972 (Page 8).