

Matrix decompositions using sub-Gaussian random matrices

YARIV AIZENBUD[†]

Department of Applied Mathematics, School of Mathematical Sciences, Tel Aviv University, Israel

[†]Corresponding author. Email: aizeny@post.tau.ac.il

AND

AMIR AVERBUCH

School of Computer Science, Tel Aviv University, Israel

[Received on 3 January 2018; revised on 10 July 2018; accepted on 16 July 2018]

In recent years, several algorithms which approximate matrix decomposition have been developed. These algorithms are based on metric conservation features for linear spaces of random projection types. We present a new algorithm, which achieves with high probability a rank- r singular value decomposition (SVD) approximation of an $n \times n$ matrix and derive an error bound that does not depend on the first r singular values. Although the algorithm has an asymptotic complexity similar to state-of-the-art algorithms and the proven error bound is not as tight as the state-of-the-art bound, experiments show that the proposed algorithm is faster in practice while providing the same error rates as those of the state-of-the-art algorithms. We also show that an i.i.d. sub-Gaussian matrix with large probability of having null entries is metric conserving. This result is used in the SVD approximation algorithm, as well as to improve the performance of a previously proposed approximated LU decomposition algorithm.

Keywords: SVD; LU decomposition; low-rank approximation; random matrices; sparse matrices; sub-Gaussian matrices; Johnson–Lindenstrauss Lemma; oblivious subspace embedding.

1. Introduction

Dimensionality reduction by randomized linear maps preserves metric features. Johnson–Lindenstrauss Lemma (JL) [15] shows that there is a random distribution of linear dimensionality reduction operators that preserves, with bounded error and high probability, the norm of a set of vectors. For example, Gaussian random matrices satisfy this property.

The JL Lemma was extended in the following way: while classical formulation deals with norm conservation of sets of vectors, the JL-based extension deals with a subspace of a vector space. This extension is considered for example in [32], where it shows that Fourier-based random matrices of size $n \times \mathcal{O}(r \log r)$ conserve the norm of all the vectors from a vector space of dimension r . Similar results for distribution of sparse matrices are given in [5, 7, 16, 22].

In recent years, several randomized algorithms, which approximate matrix decomposition based on norm conservation, have been developed. The idea is roughly as follows: a randomly drawn matrix Ω , which projects the original matrix to a lower dimension, is used. The decomposition is calculated in the low dimensional space. Then, this decomposition is mapped back into a space of the original matrix dimension. It is shown in [18, 29] how to use random Gaussian matrices in order to find, with



Symbol indicates reproducible data.

high probability, an approximated interpolative decomposition, singular value decomposition (SVD) and LU decomposition. FFT-based random matrices, called SRFT (Sub-sampled Randomized Fourier Transform), which approximate matrix decompositions, are described in [36]. The special structure of the FFT-based distribution provides a fast matrix multiplication that yields a faster algorithm than the algorithms in [18]. A comprehensive review of these ideas (and many more) is given in [11]. The Sparse Embedding matrices (SEMs) random distribution is a special distribution that is used in [5]. The SEM distribution assigns one non-zero entry in each column. The algorithm in [5] uses the SEM for a new low-rank approximation algorithm that is asymptotically faster than the FFT-based algorithm.

All these algorithms rely on special properties of the randomly drawn matrix Ω . The distribution from which Ω is drawn should be metric conserving. Let $M_{n \times m}$ be a set of n by m matrices. We call a rectangular random matrix distribution $\mathcal{M}$ a *metric conserving* distribution, if for any $A \in M_{n \times m}$ and a randomly chosen $\Omega \in M_{m \times k}$ from $\mathcal{M}$, the image of $A\Omega$ is similar to the image of A . Three main parameters related to this property are the dimension k of Ω (the smaller the better), the ‘distance’ between the images of $A\Omega$ and A and the probability with which the image conservation is valid. It is obvious that these parameters are connected. Distributions, which conserve the norm allowing an error $(1 + \varepsilon)$ of the theoretical bound, are called *oblivious subspace embeddings* (OSE) [22]. The theoretical bound for a rank- r approximation of a matrix A in L_2 norm is $\sigma_{r+1}(A)$ and in Frobenius norm it is Δ_{r+1} , where $\sigma_r(A)$ is the r th largest singular value of A and $\Delta_r(A) \triangleq (\sum_{l=r}^n \sigma_l^2(A))^{1/2}$. We use the notation $\mathcal{O}_\sigma(x)$ to denote values that equal to x up to a constant that does not depend on the singular values of A .

Three important results related to the above parameters, which deal with metric conserving distributions in the context of randomized decomposition algorithms, are as follows: (1) achieving an accuracy of $\mathcal{O}_\sigma(\sigma_{r+1}(A))$ for a rank r approximation with high probability is described in [11, 18], where the error is measured in L_2 norm. To achieve this accuracy with high probability the required Ω can be an i.i.d. Gaussian matrix of size $\mathcal{O}(r)$; (2) achieving an accuracy of $\mathcal{O}_\sigma(\sigma_{r+1}(A))$ for a rank r approximation with high probability is described in [11, 36], where the error is measured in L_2 norm. To achieve this accuracy with high probability, Ω can be an FFT-based matrix of size $\mathcal{O}(r \log r)$; (3) the result in [22] achieves accuracy of $(1 + \varepsilon)\Delta_{r+1}(A)$ with high probability, where the error is measured in Frobenius norm. While Ω is drawn from a sparse distribution, its size is assumed to be not less than $\mathcal{O}(r^2/\varepsilon^2)$. In fact, for sparse matrix distributions, a lower bound for the size of Ω is provided in [23]. A summary of recent results regarding the performance of randomized decomposition algorithms appears in Table 1.1.

TABLE 1.1 Summary of the performance of randomized decomposition algorithms. Interpolative Decomposition is denoted by ID

Ref.	matrix type	alg.	norm	accuracy	computational complexity
[18]	Gaussian	SVD ID	L_2	$\mathcal{O}_\sigma(\sigma_{r+1}(A))$	SVD: $\mathcal{O}(mnr + (m+n)r^2)$ ID: $\mathcal{O}(mnr + n \log nr^2)$
[36]	SRFT	SVD ID	L_2	$\mathcal{O}_\sigma(\sigma_{r+1}(A))$	SVD: $\mathcal{O}(mn \log r + (m+n)r^2 + r^4 \log r)$ ID: $\mathcal{O}(mn \log r + nr^3)$
[29]	SRFT	LU	L_2	$\mathcal{O}_\sigma(\sigma_{r+1}(A))$	$\mathcal{O}(mn \log r + (m+n)r^2 + r^4 \log r)$
[5]	SEM	SVD	Frobenius	$(1 + \varepsilon)\sigma_{k+1}$	$\mathcal{O}(\text{nnz}(A)) + \tilde{\mathcal{O}}(nr^2 + r^3)$
[1]	SEM	LU	Frobenius	$\mathcal{O}_\sigma(\sigma_{r+1}(A))$	$\mathcal{O}(\text{nnz}(A)) + \tilde{\mathcal{O}}(nr^2 + r^3)$
This paper	sub-Gaussian, SEM	SVD LU	L_2	$\mathcal{O}_\sigma(\sigma_{r+1}(A))$	$\mathcal{O}(\text{nnz}(A)) + \tilde{\mathcal{O}}(nr^2 + r^3)$ (using SEM)

Other problems in numerical linear algebra have benefited from the norm conservation properties of random matrix distributions. When solving large least-squares problems, several randomized algorithms have been proposed (e.g. [24,28]). A randomized algorithm that improves on LAPACK's performance is derived in [2]. Based on the same theory of SEMs, [5] also proposes an algorithm for least-squares regression. The work of [33] improves [5] by using sparse matrix distribution with more non-zeros per column and an improved algorithm for approximated Least-Squares Regression. Although there are no theoretical guarantees in [33], numerical results show improvements are available. We later show that assigning several non-zeros per column improves the performance of the solution of the low-rank approximation problem.

The main contribution of this paper is twofold: first, we show that the matrices with i.i.d. sub-Gaussian entries satisfy the image conservation property even when the probability for a zero entry grows with the size of the matrix. Stronger bounds are shown in [6], but only for the special case of sparse-Bernoulli random matrices, and [5,22] consider non-i.i.d. distributions. Secondly, and independently, we construct fast SVD and LU decomposition algorithms with a bounded error and asymptotic complexity that is equal to the asymptotic complexity of the state-of-the-art algorithms. Although the asymptotic complexity is the same, the actual running times of the presented algorithms are lower than existing algorithms. Since random projections are matrices with i.i.d. entries, it is unnecessary to set in advance the dimension k of the projection. It is possible, although not elaborated in this paper, to increase k iteratively until the resulting approximation reaches the required accuracy. The described algorithms work with any metric-preserving distribution, not necessarily with a distribution with i.i.d. sub-Gaussian entries.

Metric-conserving features of random matrix distributions in different formulations have been studied extensively outside the numerical linear algebra domain discussed here starting from the JL Lemma. Metric-conserving features are strongly connected to compressed sensing (e.g. [4,9]). Roughly speaking, in compressed sensing, the goal is to reconstruct a vector from random samples (this is a different way to look at Ω from above), where the vector belongs to a special set of sparse vectors. Many results in compressed sensing consider Gaussian random matrices and some are extended to the more general class of sub-Gaussian random matrices. The results in [19,20] are maybe aimed to investigate the linear approximated reconstruction for problems in compressed sensing as can be understood from the follow-up work [21], but the results are quite general. They can be applied to the setting considered here and provide similar statements as in Section 3 (Theorems 3.5 and 3.4). For our purposes, the sparsity of the projection matrix plays a crucial role, and in this sense, the results presented here provide stronger bounds.

We show in Section 3 that for the class of matrices with i.i.d. sub-Gaussian entries, the required size of Ω to achieve an accuracy $\mathcal{O}_\sigma(\sigma_{r+1}(A))$ is measured in L_2 norm. We also show its dependency on the probability to have a zero entry. By choosing a sparse matrix distribution to be sub-Gaussian, we are able to perform a fast matrix multiplication while having the projection matrix Ω to be of a low dimension relative to the size of the problem. It is shown in [8] that this class of sub-Gaussian matrices of size $\mathcal{O}(r/\varepsilon^2)$ with constant probability distribution is an OSE. In this paper, we provide a bound for the case where the distribution depends on the size of the matrix. Note that the generality of the results in this paper has a weakness of having only asymptotic results. For more specific cases of matrices, much tighter bounds are known (see [35] for Gaussian i.i.d. matrices). For our purpose of low-rank matrix approximation, we show experimentally that the constants are not too disturbing, but obviously, a tighter bound is a stronger result which is more practical.

The state-of-the-art result for rank- r approximation algorithm appears in [5]. It describes how to use an SEM to construct an algorithm that finds for any matrix $A \in M_{m \times n}$ and any rank r , with high

probability, an SVD approximation of rank r . Namely, orthogonal matrices U, V^* and a diagonal matrix Σ are formed such that $\|A - U\Sigma V^*\|_F \leq (1+\varepsilon)\Delta_{r+1}(A)$. Although the algorithm in [11] uses projection Ω into a smaller dimension than the one used in [5], the algorithm in [5] is asymptotically faster than the algorithm in [11] because of the sparsity of the projection.

We describe in Section 4.1 an algorithm that for each $A \in M_{m \times n}$ outputs with high probability a low-rank SVD approximation that is built from U, Σ and V . The algorithm works with any metric-conserving or OSE random distribution. The size k of the random embedding in the algorithm depends on the probability p of having a zero entry. The complexity of the algorithm when using i.i.d. sub-Gaussian random matrix projections is $\mathcal{O}(\text{nnz}(A)pk + (m+n)k^2)$, where $\text{nnz}(A)$ denotes the *number of non-zeros* in A . For SEM distributions as in [22], the complexity of the algorithm in Section 4.1 is the same as in [22]. The presented algorithm guarantees with high probability that $\|A - U\Sigma V^*\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1}(A))$. Although the guaranteed error bound is less tight than [5], we show in Section 5 that in practice our algorithm reaches the same error in less time.

The randomized LU decomposition algorithm in [1] is based on ideas from [5]. We show in Section 4.2 that it is also valid when random matrices from a sub-Gaussian distribution are chosen with the complexity and error bound that is equal to those from the SVD decomposition.

The paper has the following structure: In Section 2, we present the necessary mathematical preliminaries. In Section 3, we show that i.i.d. sub-Gaussian random matrices are metric conserving and in Section 4 we describe the SVD algorithm, and show that the LU algorithm in [1] is valid with i.i.d. sub-Gaussian random matrices. In Section 5, we present numerical results obtained with the approximated SVD algorithm from Section 4.

2. Preliminaries

2.1 The ε -net

The ε -net is introduced in Definition 2.1. Its size is bounded by Lemma 2.1 which is proved in [27]. Throughout the paper, S^{n-1} denotes the $(n-1)$ -sphere in $\mathbb{R}^n$.

DEFINITION 2.1 Let (T, d) be a metric space and let $K \subset T$. A set $\mathcal{N} \subset T$ is called an ε -net of K if for all $x \in K$ there exists $y \in \mathcal{N}$ such that $d(x, y) < \varepsilon$.

LEMMA 2.1 (Proposition 2.1 in [27]). For any $\varepsilon < 1$, there exists an ε -net $\mathcal{N}$ of S^{n-1} such that

$$|\mathcal{N}| \leq 2n \left(1 + \frac{2}{\varepsilon}\right)^{n-1}.$$

Remark. It follows that for sufficiently large n we have

$$|\mathcal{N}| \leq \left(\frac{3}{\varepsilon}\right)^n.$$

Moreover, a $1/2$ -net of S^{n-1} has at most

$$2n \left(1 + \frac{2}{1/2}\right)^{n-1} = 2n \cdot 5^{n-1} \leq 6^n$$

points.

2.2 Compressible and incompressible vectors

DEFINITION 2.2 A vector $v \in \mathbb{R}^n$ with $\|v\| = 1$ is called (ε, η) -incompressible if $\sum_{j:|v_j| \leq \varepsilon} |v_j|^2 \geq \eta^2$ and compressible otherwise.

LEMMA 2.2 Let $U \subset \mathbb{R}^n$ be a subspace of dimension r . There exist $\mathcal{N}$, an ε_{net} -net of the set of (ε_c, η) -compressible vectors in U , such that

$$|\mathcal{N}| \leq r^{\frac{1}{\varepsilon_c^2}} \eta^{r - \frac{1}{\varepsilon_c^2}} \left(\frac{c_{\text{net}}}{\varepsilon_{\text{net}}} \right)^r,$$

for an absolute constant c_{net} .

Proof. The (η, ε_c) -compressible vectors lie at a distance of η from a sparse vector with no more than ε_c^{-2} non-zero coordinates. For small enough η , the volume of η -balls around ε_c^{-2} -sparse vectors is $r^{\varepsilon_c^{-2}} \eta^{r - \frac{1}{\varepsilon_c^2}}$. The same arguments from the proof of Lemma 2.1 show that the number of points in an ε_{net} -net of this volume does not exceed $r^{\frac{1}{\varepsilon_c^2}} \eta^{r - \frac{1}{\varepsilon_c^2}} \left(\frac{c_{\text{net}}}{\varepsilon_{\text{net}}} \right)^r$. $\square$

2.3 Sub-Gaussian random variables

In this section, we introduce sub-Gaussian random variables and discuss some of their properties. Sub-Gaussian variables are an important class of random variables that have strong tail decay properties. This class contains, for example, all the bounded random variables and all the normal variables.

DEFINITION 2.3 A random variable X is called sub-Gaussian if there exist constants v and C such that, for any $t > 0$, $\mathbb{P}(|X| > t) \leq Ce^{-vt^2}$ and X has a non-zero variance. A random variable X is called centred if $\mathbb{E}X = 0$ and normalized if, in addition, $\mathbb{E}(X^2) = 1$.

Remark. For convenience, we use the term *sub-Gaussian matrix* for a matrix with i.i.d. sub-Gaussian entries.

Many non-asymptotic results on sub-Gaussian matrix distributions have been obtained recently. A survey of this topic appears in [25, 34].

The following facts, proved in [17, 25–27, 34], are used in the paper:

1. Linear combinations of centred sub-Gaussian variables are also sub-Gaussian. This is stated in Lemma 2.3. The inequality in this theorem is similar to the Hoeffding inequality [13].
2. A bound for the first singular value of a sub-Gaussian random matrix is given in Lemma 2.4.
3. A bound on the probability of a sum of centred sub-Gaussian variables to be small is given in Lemma 2.6.

Formally,

LEMMA 2.3 Let $X_1, \dots, X_n$ be independent centred sub-Gaussian random variables. Then, for any $a_1, \dots, a_n \in \mathbb{R}$,

$$\mathbb{P} \left(\left| \sum_{j=1}^n a_j X_j \right| > t \right) \leq 2 \exp \left(- \frac{ct^2}{\sum_{j=1}^n a_j^2} \right).$$

LEMMA 2.4 Let Ω be a $k \times n$, $n \geq k$, random matrix whose entries are i.i.d. centred sub-Gaussian random variables. Then, $\mathbb{P}(\sigma_1(\Omega) > t\sqrt{n}) \leq e^{-c_0 t^2 n}$ for $t \geq C_0$, where C_0 and c_0 are universal constants.

Since we are interested in sparse matrices, Observation 2.4 is useful.

OBSERVATION 2.4 Any normalized sub-Gaussian random variable X is representable as a combination of a centred sub-Gaussian random variable $\frac{1}{\sqrt{p}}Z$ with $\mathbb{P}(Z = 0) = 0$, $\mathbb{E}(Z^2) = 1$ with probability p and 0 otherwise. Note that $\mathbb{E}(X) = 0$, $\mathbb{E}(X^2) = 1$, $\mathbb{E}(X^3) = \frac{\mathbb{E}(Z^3)}{\sqrt{p}}$ and $\mathbb{E}(X^4) = \frac{\mathbb{E}(Z^4)}{p}$.

LEMMA 2.5 Let X and Z be centred sub-Gaussian random variables as in Observation 2.4. Let $X_1, \dots, X_n$ be i.i.d. copies of X . Then, for any $(a_1, \dots, a_n) \in S^{n-1}$ the third and fourth moment (skewness and kurtosis) of $\sum_{i=1}^n a_i X_i$ are bounded by

$$\mathbb{E}\left(\sum_{i=1}^n a_i X_i\right)^3 \leq \frac{\mathbb{E}(Z^3)}{\sqrt{p}}$$

and

$$\mathbb{E}\left(\sum_{i=1}^n a_i X_i\right)^4 \leq \frac{\mathbb{E}(Z^4) + p}{p} \triangleq \frac{z_4}{p}.$$

Proof. Rewriting the Left hand side we have,

$$\mathbb{E}\left(\sum_{i=1}^n a_i X_i\right)^3 = \sum_{i=1}^n a_i^3 \mathbb{E}(X_i^3) + \sum_{i,j=1, i \neq j}^n a_i^2 a_j \mathbb{E}(X_i^2) \mathbb{E}(X_j) \leq \mathbb{E}(X^3) = \frac{\mathbb{E}(Z^3)}{\sqrt{p}}$$

and

$$\mathbb{E}\left(\sum_{i=1}^n a_i X_i\right)^4 = \sum_{i=1}^n a_i^4 \mathbb{E}(X_i^4) + \sum_{i,j=1, i \neq j}^n a_i^2 a_j^2 \leq \mathbb{E}(X^4) + 1 = \frac{\mathbb{E}(Z^4) + p}{p}.$$

Since $p \leq 1$, the proof is complete. $\square$

LEMMA 2.6 Let X and Z be sub-Gaussian random variables as in Observation 2.4. Let $X_1, \dots, X_n$ be i.i.d. copies of X . For every coefficients vector (in particular, for a compressible vector) $a = (a_1, \dots, a_n) \in S^{n-1}$, the random sum $S = \sum_{i=1}^n a_i X_i$ satisfies $\mathbb{P}(|S| < \lambda) \leq 1 - p \frac{(1-\lambda^2)^2}{z_4}$.

Proof. Let $0 < \lambda < (\mathbb{E}S^2)^{1/2} = 1$. By the Cauchy–Schwarz inequality,

$$\mathbb{E}S^2 = \mathbb{E}S^2 \mathbf{1}_{[-\lambda, \lambda]}(S) + \mathbb{E}S^2 \mathbf{1}_{\mathbb{R} \setminus [-\lambda, \lambda]}(S) \leq \lambda^2 + \left(\mathbb{E}S^4\right)^{1/2} \mathbb{P}(|S| > \lambda)^{1/2}.$$

This leads to the Paley–Zygmund inequality

$$\mathbb{P}(|S| > \lambda) \geq \frac{(\mathbb{E}S^2 - \lambda^2)^2}{\mathbb{E}S^4} = \frac{(1 - \lambda^2)^2}{\mathbb{E}S^4}.$$

By Theorem 2.3, the random variable S is sub-Gaussian. By Lemma 2.5, $\mathbb{E}S^4 \leq \frac{z_4}{p}$, where $z_4 = \mathbb{E}Z^4 + 1$. To complete the proof, observe that

$$\mathbb{P}(|S| < \lambda) \leq 1 - \frac{(1 - \lambda^2)^2}{\mathbb{E}S^4} = 1 - p \frac{(1 - \lambda^2)^2}{z_4}.$$

□

COROLLARY 2.7 By substituting $\lambda = 1/2$ in Lemma 2.6 we have $\mathbb{P}(|S| < 1/2) \leq 1 - z'_4 p$ for $z'_4 = \frac{9}{16z_4}$.

LEMMA 2.8 For any $0 < \alpha < 1$, there exists a constant c_s such that, for any k ,

$$\binom{k}{\alpha k} < c_s^{\alpha k \ln(\frac{1}{\alpha} - 1)} < c_s^{\alpha k \ln \frac{1}{\alpha}}.$$

Proof. We use the Stirling formula to estimate $\ln \binom{k}{\alpha k}$

$$\begin{aligned} \ln \binom{k}{\alpha k} &= k \ln k - k - (\alpha k \ln(\alpha k) - \alpha k) - ((k - \alpha k) \ln(k - \alpha k) - (k - \alpha k)) + \mathcal{O}(\ln k) \\ &= k \ln k - \alpha k \ln(\alpha k) - k \ln(k - \alpha k) + \alpha k \ln(k - \alpha k) + \mathcal{O}(\ln k) \\ &= k \ln k - \alpha k \ln \alpha k - k \ln k(1 - \alpha) + \alpha k \ln \left(\alpha k \left(\frac{1}{\alpha} - 1 \right) \right) + \mathcal{O}(\ln k) \\ &= \alpha k \ln \left(\frac{1}{\alpha} - 1 \right) - k \ln(1 - \alpha) + \mathcal{O}(\ln k) \\ &\sim \alpha k \ln \left(\frac{1}{\alpha} - 1 \right) - k \left(-\alpha - \frac{\alpha^2}{2} - \dots \right) \sim \alpha k \ln \left(\frac{1}{\alpha} - 1 \right). \end{aligned}$$

□

Lemma 2.9 follows from the Berry–Essen theorem [3] in much the same way as done in [31].

LEMMA 2.9 Assume $S = \sum_{i=1}^n a_i X_i$, where X_i are i.i.d. random variables with $E(X_i) = 0, E(X_i^2) = 1$. Then, for all r

$$\sup_t \mathbb{P} \left(\left| \sum_{i=1}^n a_i X_i - t \right| < r \right) \leq \frac{r}{\|a\|} + 2C_{\text{BE}} \frac{\mathbb{E}(X_1^3) \sum_{i=1}^n |a_i|^3}{\|a\|^3} \quad (2.1)$$

holds, where $\|a\| = \sqrt{\sum_{i=1}^n a_i^2}$.

Proof. Let N be a standard normal variable. From the Berry–Essen theorem it follows that for all r

$$\left| \mathbb{P} \left(\frac{\sum_{i=1}^n a_i X_i}{\|a\|} < \frac{r}{\|a\|} \right) - \mathbb{P} \left(N < \frac{r}{\|a\|} \right) \right| \leq C_{\text{BE}} \frac{\mathbb{E}(X_1^3) \sum_{i=1}^n |a_i|^3}{\|a\|^3}.$$

Thus, for any t ,

$$\begin{aligned}
 & \left| \mathbb{P} \left(\frac{|\sum_{i=1}^n a_i X_i - t|}{\|a\|} < \frac{r}{\|a\|} \right) - \mathbb{P} \left(|N - t| < \frac{r}{\|a\|} \right) \right| \\
 &= \left| \mathbb{P} \left(\frac{\sum_{i=1}^n a_i X_i}{\|a\|} < \frac{t+r}{\|a\|} \right) - \mathbb{P} \left(N < \frac{t+r}{\|a\|} \right) \right. \\
 &\quad \left. - \mathbb{P} \left(\frac{\sum_{i=1}^n a_i X_i}{\|a\|} < \frac{t-r}{\|a\|} \right) - \mathbb{P} \left(N < \frac{t-r}{\|a\|} \right) \right| \quad (2.2) \\
 &\leq 2C_{\text{BE}} \frac{\mathbb{E}(X_1^3) \sum_{i=1}^n |a_i|^3}{\|a\|^3}.
 \end{aligned}$$

By rewriting (2.2), we have

$$\mathbb{P} \left(\left| \sum_{i=1}^n a_i X_i - t \right| < r \right) \leq \mathbb{P} \left(|N - t| < \frac{r}{\|a\|} \right) + 2C_{\text{BE}} \frac{\mathbb{E}(X_1^3) \sum_{i=1}^n |a_i|^3}{\|a\|^3}. \quad (2.3)$$

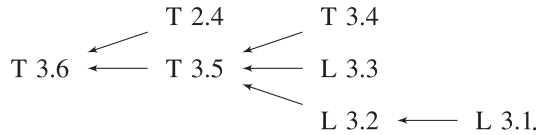
For any t , $\mathbb{P} \left(|N - t| < \frac{r}{\|a\|} \right) \leq \mathbb{P} \left(|N| < \frac{r}{\|a\|} \right) < \frac{r}{\|a\|}$. □

3. Metric conservation of sub-Gaussian random matrices

The main goal of this section is to show that for any matrix A and for a sub-Gaussian matrix Ω , the image of $A\Omega$ is ‘close’ to the image of A with high probability. In other words, Ω preserves the geometry. More precisely, if Q is an orthogonal basis for $A\Omega$, then $\|A - QQ^*A\|_2$ is small.

In order to show that the application of a random sub-Gaussian matrix preserves the geometry of A , we have to bound its behaviour in any subspace of a given dimension r . We show in Theorem 3.4 that the norm of a random sub-Gaussian matrix in a subspace of dimension r is bounded from above with high probability. In Lemmas 3.2 and 3.3, it is shown that Ω conserves incompressible vectors and arbitrary vectors, respectively, from a subspace of dimension r . In Theorem 3.5, these results are combined to show that the minimal singular value is bounded from below with high probability. The flow of the proof is based on ideas from the proof of bounds on singular values of Bernoulli random matrices in [31] and ideas from [25]. In Theorem 3.6, these results and the fact that the norm of a random matrix is also bounded (Theorem 2.4) are used to show that a sub-Gaussian matrix preserves the geometry.

The dependencies among the different theorems in this section are shown next.



Note that despite the similar look of Theorem 5.39 in [34] and Theorems 3.5 and 3.4, Theorems 3.5 and 3.4 are essentially different since they deal with the values of the operator on a small subspace and not on the whole space.

LEMMA 3.1 Let X and Z be centred sub-Gaussian random variables as in Observation 2.4. Let $X_1, \dots, X_n$ be i.i.d. copies of X . Denote $\varepsilon_c = \varepsilon_0 \eta \sqrt{p}$. For any (η, ε_c) -incompressible $a = (a_1, \dots, a_n) \in S^{n-1}$, the random sum $\sum_{i=1}^n a_i X_i$ satisfies

$$P\left(\left|\sum_{i=1}^n a_i X_i\right| \leq 2C\varepsilon_0 \eta\right) \leq 2C_1(C)\varepsilon_0 \mathbb{E}(Z^3),$$

where C is a constant that will be chosen later, and $C_1(C)$ depends only on C .

Proof. We recall that a is incompressible if $\sum_{j:|a_j| \leq \varepsilon_c} |a_j|^2 \geq \eta^2$. Using Lemma 2.9, we get

$$\sup_t \mathbb{P}\left(\left|\sum_{i=1}^n a_i X_i - t\right| \leq r\right) \leq \frac{r}{\|a\|} + 2C_{\text{BE}} \frac{\mathbb{E}(X^3) \sum_{i=1}^n |a_i|^3}{\|a\|^3}.$$

Note that we can condition out variables

$$\sup_t \mathbb{P}\left(\left|\sum_{i=1}^n a_i X_i - t\right| \leq r\right) \leq \sup_t \mathbb{P}\left(\left|\sum_{i=1}^n a_i X_i - t\right| \leq r \mid X_1 = x_1\right).$$

If we condition out all the X_i for which $a_i > \varepsilon_c$, we get

$$\sup_t \mathbb{P}\left(\left|\sum_{i=1}^n a_i X_i - t\right| \leq r\right) \leq \frac{r}{\sqrt{\sum_{j:|a_j| \leq \varepsilon_c} a_j^2}} + 2C_{\text{BE}} \frac{\mathbb{E}(X^3) \sum_{j:|a_j| \leq \varepsilon_c} |a_j|^3}{\left(\sum_{j:|a_j| \leq \varepsilon_c} a_j^2\right)^{3/2}} \leq \frac{r}{\eta} + 2C_{\text{BE}} \frac{\mathbb{E}(X^3)\varepsilon_c}{\eta}.$$

Substituting $r = 2C\varepsilon_0 \eta$, we have

$$\sup_t \mathbb{P}\left(\left|\sum_{i=1}^n a_i X_i - t\right| \leq 2C\varepsilon_0 \eta\right) \leq \frac{2C\varepsilon_0 \eta}{\eta} + 2C_{\text{BE}} \frac{\mathbb{E}(X^3)\varepsilon_c}{\eta}.$$

By using Lemma 2.5 and by substituting $\varepsilon_c = \varepsilon_0 \eta \sqrt{p}$ the proof is completed. $\square$

LEMMA 3.2 (Ω conserves incompressible vectors in a subspace). Let X and Z be centred sub-Gaussian random variables as in Observation 2.4. Let Ω be a $k \times n$ ($n \geq k$) random matrix whose entries are i.i.d. copies of X . Denote $\varepsilon_c = \varepsilon_0 \eta \sqrt{p}$. Then, for any (ε_c, η) -incompressible $x \in S^{n-1}$,

$$\mathbb{P}(\|\Omega x\|_2 < 2C\alpha\varepsilon_0 \eta \sqrt{k}) \leq (C_1(C)\varepsilon_0 \mathbb{E}(Z^3))^{k/4}$$

for a constant α .

Proof. The coordinates of the vector Ωx are independent linear combinations of i.i.d. sub-Gaussian random variables with incompressible coefficients $(x_1, \dots, x_n) \in S^{n-1}$. Hence, by Lemma 3.1, $\mathbb{P}(|(\Omega x)_j| < 2C\varepsilon_0 \eta) \leq C_1(C)\varepsilon_0 \mathbb{E}(Z^3) = \mu$ for all $j = 1, \dots, k$.

Assume that $\|\Omega x\|_2 < 2C\varepsilon_0\eta\alpha\sqrt{k}$. Then, $|(\Omega x)_j| < 2C\varepsilon_0\eta$ for at least $\lfloor(1 - \alpha^2)k\rfloor$ coordinates. Thus,

$$\begin{aligned} \mathbb{P}(\|\Omega x\|_2 < 2C\varepsilon_0\eta\alpha\sqrt{k}) &\leq \mathbb{P}\left(\text{at least } \lfloor(1 - \alpha^2)k\rfloor \text{ coordinates satisfy } |(\Omega x)_j| < 2C\varepsilon_0\eta\right) \\ &= \sum_{l=\lfloor(1-\alpha^2)k\rfloor}^k \binom{k}{l} \mathbb{P}(|(\Omega x)_1| < 2C\varepsilon_0\eta)^l (1 - \mathbb{P}(|(\Omega x)_1| < 2C\varepsilon_0\eta))^{k-l} \\ &\leq \sum_{l=\lfloor(1-\alpha^2)k\rfloor}^k \binom{k}{l} \mu^l \\ &\leq \mu^{\lfloor(1-\alpha^2)k\rfloor} \sum_{l=\lfloor(1-\alpha^2)k\rfloor}^k \binom{k}{l}. \end{aligned}$$

If α is sufficiently small, then $\lfloor(1 - \alpha^2)k\rfloor > k/2$ and

$$\mathbb{P}(\|\Omega x\|_2 < 2C\varepsilon_0\eta\alpha\sqrt{k}) \leq \mu^{k/2} \alpha^2 k \binom{k}{\lfloor(1 - \alpha^2)k\rfloor}.$$

For α sufficiently small, $\alpha^2 k \binom{k}{\lfloor(1-\alpha^2)k\rfloor} \leq \mu^{-k/4}$. Thus, $\mathbb{P}(\|\Omega x\|_2 < 2C\varepsilon_0\eta\alpha\sqrt{k}) \leq \mu^{-k/4}$. $\square$
 $\mu^{k/2} \leq \mu^{k/4}$.

LEMMA 3.3 (Ω conserves any vector in a subspace). Let Ω be a $k \times n$ ($n \geq k$) random matrix whose entries are i.i.d. copies of a normalized sub-Gaussian random variable X with variance 1, where X, Z and p are as in Observation (2.4). Let C be a constant that will be chosen later, and let η be sufficiently small such that $\eta^2 \ln \frac{1}{\eta} < c_4 p$. Then, for every $x \in S^{n-1}$, $\mathbb{P}(\|\Omega x\|_2 < 2C\eta\sqrt{k}) \leq (1 - z'_4 p)^{k/4}$ holds for a constant c_4 .

Proof. The coordinates of the vector Ωx are independent linear combinations of i.i.d. sub-Gaussian random variables with coefficients $(x_1, \dots, x_n) \in S^{n-1}$. Hence, Corollary 2.7 yields $\mathbb{P}(|(\Omega x)_j| < 1/2) \leq 1 - z'_4 p, j = 1, \dots, k$.

Assume that $\|\Omega x\|_2 < 2C\eta\sqrt{k}$. Then, $|(\Omega x)_j| < 1/2$ for at least $\lfloor(1 - 4 \cdot 2^2 C^2 \eta^2)k\rfloor$ coordinates. Thus,

$$\begin{aligned} \mathbb{P}(\|\Omega x\|_2 < 2C\eta\sqrt{k}) &\leq \mathbb{P}(\# \lfloor(1 - 16C^2\eta^2)k\rfloor \text{ coordinates satisfy } |(\Omega x)_j| < 1/2) \\ &= \sum_{l=\lfloor(1-16C^2\eta^2)k\rfloor}^k \binom{k}{l} \mathbb{P}(|(\Omega x)_1| < 1/2)^l (1 - \mathbb{P}(|(\Omega x)_1| < 1/2))^{k-l} \\ &\leq \sum_{l=\lfloor(1-16C^2\eta^2)k\rfloor}^k \binom{k}{l} (1 - z'_4 p)^l \\ &\leq (1 - z'_4 p)^{\lfloor(1-16C^2\eta^2)k\rfloor} \sum_{l=\lfloor(1-16C^2\eta^2)k\rfloor}^k \binom{k}{l}. \end{aligned}$$

If η is sufficiently small, then $\lfloor (1 - 16C^2\eta^2)k \rfloor > k/2$ and

$$\mathbb{P}(\|\Omega x\|_2 < 2C\eta\sqrt{k}) \leq (1 - z'_4 p)^{k/2} 16C^2 k \binom{k}{\lfloor (1 - 16C^2\eta^2)k \rfloor}.$$

From Lemma 2.8 it follows that for η sufficiently small,

$$16C^2 k \binom{k}{\lfloor (1 - 16C^2\eta^2)k \rfloor} \leq 16C^2 k c_s^{16C^2\eta^2 k \ln\left(\frac{1}{16C^2\eta^2} - 1\right)} < c_1^{\eta^2 k \ln\left(\frac{1}{16C^2\eta^2} - 1\right)} < c_2^{k\eta^2 \ln \frac{1}{\eta^2}}.$$

Additionally, $(1 - z'_4 p)^{-k/4} > c_3^{c_2 p k}$. Thus, for η such that $\eta^2 \ln \frac{1}{\eta} < c_4 p$, $16C^2 k \binom{k}{\lfloor (1 - 16C^2\eta^2)k \rfloor} \leq (1 - z'_4 p)^{-k/4}$ holds. Thus, $\mathbb{P}(\|\Omega x\|_2 < 2C\eta\sqrt{k}) \leq (1 - z'_4 p)^{-k/4} \cdot (1 - z'_4 p)^{k/2} \leq (1 - z'_4 p)^{k/4}$. $\square$

The proof of Theorem 3.4 is similar to the proof of Theorem 2.4.

THEOREM 3.4 (Maximum value in a subspace). Let $U \subset \mathbb{R}^n$ be a linear subspace of dimension r . Let Ω be a $k \times n$ random matrix where $n \geq k > r$ and $k = \mathcal{O}(r)$ is sufficiently large. Assume the entries of Ω are i.i.d. centred sub-Gaussian random variables. Then, for $t \geq C_0$ we have

$$\mathbb{P}\left(\max_{x \in U, \|x\|=1} \|\Omega x\| > t\sqrt{k}\right) \leq e^{-c_0 t^2 k}.$$

Proof. Let $\mathcal{N}$ be a $(1/2)$ -net of the r -dimensional unit sphere of the image of U . Let $\mathcal{M}$ be a $(1/2)$ -net of the k -dimensional unit sphere of the image of Ω . For any $u \in U$, where $\|u\| = 1$, we can choose $x \in \mathcal{N}$ such that $\|x - u\|_2 < 1/2$. Then,

$$\|\Omega u\|_2 \leq \|\Omega x\|_2 + \|x - u\|_2 \max_{u_1 \in U, \|u_1\|=1} \|\Omega u_1\|.$$

Thus,

$$\max_{u_1 \in U, \|u_1\|=1} \|\Omega u_1\| \leq \|\Omega x\|_2 + \frac{1}{2} \max_{u_1 \in U, \|u_1\|=1} \|\Omega u_1\|.$$

This shows that $\|\Omega\| \leq 2 \sup_{x \in \mathcal{N}} \|\Omega x\|_2 = 2 \sup_{x \in \mathcal{N}} \sup_{v \in S^{k-1}} \langle \Omega x, v \rangle$. In a similar way, by approximating v with an element from $\mathcal{M}$ we get

$$\sup_{x \in \mathcal{N}, v \in S^{k-1}} \langle \Omega x, v \rangle \leq \sup_{x \in \mathcal{N}, v \in \mathcal{M}} \langle \Omega x, v \rangle + \frac{1}{2} \sup_{x \in \mathcal{N}, v' \in S^{k-1}} \langle \Omega x, v' \rangle.$$

We obtain $\|\Omega\| \leq 4 \max_{x \in \mathcal{N}, y \in \mathcal{M}} |\langle \Omega x, y \rangle|$. By Lemma 2.1, we can choose these nets to be such that $|\mathcal{N}| \leq 6^r$ and $|\mathcal{M}| \leq 6^k$.

By Lemma 2.3, for every $x \in \mathcal{N}$ and $y \in \mathcal{M}$, the random variable $\langle \Omega x, y \rangle = \sum_{j=1}^k \sum_{k=1}^n a_{j,k} y_j x_k$ is sub-Gaussian, i.e. for $t > 0$

$$\mathbb{P}(|\langle \Omega x, y \rangle| > t\sqrt{k}) \leq C_2 e^{-c_1 t^2 k}.$$

By taking the union bound, we get

$$\begin{aligned} \mathbb{P}(\|\Omega\|_2 > t\sqrt{k}) &\leq |\mathcal{N}||\mathcal{M}|\mathbb{P}(|\langle \Omega x, y \rangle| > t\sqrt{k}/4, x \in \mathcal{N}, y \in \mathcal{N}) \\ &\leq 6^k \cdot 6^r \cdot C_2 e^{-c_1/16t^2k} \leq C_2 e^{-c_0 t^2 k}, \end{aligned}$$

provided that $t \geq C_0$ for an appropriately chosen constant $C_0 > 0$. This completes the proof. $\square$

By combining Theorem 3.4 with the ε -net argument and Lemmas 3.2 and 3.3, we obtain an estimate for the smallest value of $\|\Omega v\|$, where v in a subspace of dimension r .

THEOREM 3.5 (Smallest value on a subspace). Let X and Z be centred sub-Gaussian random variables as in Observation 2.4 with parameter p . There are constants M, C and D such that for any $n, r \in \mathbb{N}$, $p \in \mathbb{R}$, $0 < p < 1$, $n > r$ and for any r -dimensional linear subspace $U \subset \mathbb{R}^n$, if $k > D \log\left(\frac{1}{p}\right)\left(r + \frac{1}{p^3}\right)$, then for $\Omega \in M_{k \times n}$ random matrix whose entries are i.i.d. copies of X ,

$$\mathbb{P}\left(\min_{x \in U, \|x\|=1} \|\Omega x\|_2 \leq M\eta\sqrt{k}\right) < e^{-Cpk} \quad (3.1)$$

holds for $\eta < \mathcal{O}(\sqrt{p})$.

Proof. The proof is divided into three steps. In steps 1 and 2, Ω is bounded on incompressible and compressible vectors, respectively, and in step 3 these results are combined to complete the proof. We take M so that $M > C_0$, where C_0 comes from Theorem 3.4, and α from Lemma 3.2 such that $e^{-c_0 \frac{M^2}{\alpha} k}$ from Theorem 3.4 is sufficiently small.

Step 1: Let $\mathcal{N}$ be an $\alpha\eta$ -net of the set of (ε_c, η) -incompressible vectors in the image of U . By Lemma 2.1, the number of vectors in $\mathcal{N}$ is bounded by $\left(\frac{3}{\alpha\eta}\right)^r$. From Lemma 3.2 with $C = \frac{M}{\alpha\varepsilon_0}$ it follows that for any vector $x \in \mathcal{N}$ and for $\varepsilon_c = \varepsilon_0\eta\sqrt{p}$,

$$\mathbb{P}(\|\Omega x\|_2 < 2M\eta\sqrt{k}) \leq (C_1\varepsilon_0 E(Z^3))^{k/4}.$$

Thus, by the union bound with failure probability of not more than

$$\left(\frac{3}{\alpha\eta}\right)^r \cdot (C_1\varepsilon_0 E(Z^3))^{k/4}, \quad (3.2)$$

we have that

$$\min_{x \in \mathcal{N}} \|\Omega x\|_2 \geq 2M\alpha\eta\sqrt{k}. \quad (3.3)$$

Since $\mathcal{N}$ is an $\alpha\eta$ -net of the set of (ε_c, η) -incompressible vectors in the image of U , with the probability given in Eq. (3.2), then Eq. (3.3) holds. By Theorem 3.4, for any incompressible vector y ,

$$\|\Omega y\| \geq \min_{x \in \mathcal{N}} \|\Omega x\|_2 - \alpha\eta\|\Omega\| \geq 2M\alpha\eta\sqrt{k} - M\alpha\eta\sqrt{k} = M\alpha\eta\sqrt{k}.$$

Step 2: Let $\mathcal{M}$ be an η -net of the set of (ε_c, η) -compressible vectors in the image of U . The number of vectors in $\mathcal{M}$ is bounded by Lemma 2.2 with $r^{\frac{1}{\varepsilon_c}} \eta^{-\frac{1}{\varepsilon_c}} \left(\frac{1}{\eta}\right)^r = r^{\frac{1}{\varepsilon_c}} \eta^{-\frac{1}{\varepsilon_c}}$. From Lemma 3.3 with $C = M$ it follows that for any vector $x \in \mathcal{M}$,

$$\mathbb{P}(\|\Omega x\|_2 < 2M\eta\sqrt{k}) \leq (1 - z'_4 p)^{k/4}.$$

Thus, by the union bound with failure probability not exceeding

$$r^{\frac{1}{\varepsilon_c}} \eta^{-\frac{1}{\varepsilon_c}} \cdot (1 - z'_4 p)^{k/4}, \quad (3.4)$$

we have

$$\min_{x \in \mathcal{M}} \|\Omega x\|_2 \geq 2M\eta\sqrt{k}. \quad (3.5)$$

For any compressible vector y we have

$$\|\Omega y\| \geq \min_{x \in \mathcal{N}} \|\Omega x\|_2 - \eta \|\Omega\| \geq 2M\eta\sqrt{k} - M\eta\sqrt{k} = M\eta\sqrt{k}.$$

Thus, if the numbers in Eqs. (3.4) and (3.2) are small enough, then,

$$\min_{y \in U, \|y\|=1} \|\Omega y\| \geq M\eta\sqrt{k}.$$

Step 3: The probabilities in Eqs. (3.2) and (3.4) are analysed next. We have

$$\left(\frac{3}{\alpha\eta}\right)^r \cdot \left(C_1 \varepsilon_0 E(Z^3)\right)^{k/4} = \exp\left(r \log\left(\frac{3}{\alpha\eta}\right) - k/4 \log\left(\frac{1}{C_1 \varepsilon_0 z_3}\right)\right)$$

$$r^{\frac{1}{\varepsilon_c}} \eta^{-\frac{1}{\varepsilon_c}} \cdot (1 - z'_4 p)^{k/4} \leq \exp\left(\frac{1}{\varepsilon_0^2 \eta^2 p} \log(r) + \frac{1}{\varepsilon_0^2 \eta^2 p} \log\left(\frac{1}{\eta}\right) - c_1 p k\right)$$

for some c_1 that depends only on z_4 . For any $\epsilon > 0$, Lemma 3.3 holds for $\eta = p^{1/2-\epsilon}$ and the probabilities in Eqs. (3.2) and (3.4) are less than

$$\exp\left(r \log\left(\frac{3}{\alpha\eta}\right) - k/4 \log\left(\frac{1}{M \varepsilon_0 z_3}\right)\right) < \exp\left(\frac{1}{2} r \log\left(\frac{c_5}{p}\right) - c_6 k\right)$$

and

$$\exp\left(\frac{1}{\varepsilon_0^2 \eta^2 p} \log(r) + \frac{1}{\varepsilon_0^2 \eta^2 p} \log\left(\frac{1}{\eta}\right) - c_1 p k\right) < \exp\left(c_7 \frac{1}{p^2} \log(r) + c_8 \frac{1}{p^2} \log\left(\frac{1}{\eta}\right) - c_1 p k\right)$$

for constants c_i . Thus, for Eq. (3.1) to hold, k has to satisfy

$$c_9 r \log \left(\frac{c_5}{p} \right) \ll k \quad (3.6)$$

and

$$c_{10} \frac{1}{p^3} \log(r) + c_{11} \frac{1}{p^3} \log \left(\frac{1}{p} \right) \ll k. \quad (3.7)$$

Note that $\frac{1}{p^3} \log(r)$ is bounded by $\mathcal{O} \left(r \log \left(\frac{c_5}{p} \right) \right)$ or $\mathcal{O} \left(\frac{1}{p^3} \log \left(\frac{1}{p} \right) \right)$. Thus, there exists a constant D such that Eqs. (3.6) and (3.7) are equivalent to

$$D \log \left(\frac{1}{p} \right) \left(r + \frac{1}{p^3} \right) \ll k.$$

□

We showed that the following known result (e.g. see [11]) can be used with sub-Gaussian matrices.

THEOREM 3.6 (Theorem 11.2 in [11]). Let A be an $m \times n$ matrix with singular values $\sigma_1, \dots, \sigma_n$ in descending order. For any integer $0 < r < m$, let Ω be an $n \times k$ random matrix. Denote $Y = A\Omega$ and $Y = QR$, where Q is a matrix with orthonormal columns and R is a full rank triangular matrix. If for any subspace $U \subset \mathbb{R}^n$ of dimension k , $\min_{x \in U} \|\Omega x\|_2$ and $\|\Omega\|_2$ are bounded from below and from above, respectively, with high probability, then, with high probability,

$$\|A - QQ^*A\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1}) \quad (3.8)$$

and

$$\|A - QQ^*A\|_F \leq \mathcal{O}_\sigma(\Delta_{r+1}). \quad (3.9)$$

Remark. Note that the notation $\mathcal{O}_\sigma(\sigma_{r+1})$ means that the error does not depend on the singular values except of having a linear dependency on σ_{r+1} . Dependency exists on n and k .

REMARK 3.7 Note that if $\Omega_1 \in M_{k \times n}$ satisfies the conditions of Theorem 3.6 with failure probability δ_1 and $\Omega_2 \in M_{l \times k}$ satisfies the conditions of Theorem 3.6 with failure probability δ_2 , then $\Omega = \Omega_2 \Omega_1$ also satisfies the conditions of Theorem 3.6 with failure probability of at most $\delta_1 + \delta_2$. This fact is important, since it enables us to combine embedding matrices in order to achieve an additional dimensionality reduction. A similar statement is introduced in [5] as Fact 45.

4. Approximated matrix decompositions

4.1 Randomized SVD using sparse projections

We present an algorithm that approximates the SVD decomposition of any matrix A . From Theorem 3.6 it follows that the randomized SVD Algorithm 5.1 in [11] is valid for sub-Gaussian matrices. This algorithm does not take advantage of the fact that Ω can be a sparse matrix. Thus, Algorithm 5.1 can be adapted similarly to the algorithm in Theorem 47 [5] and to the LU decomposition algorithm [1]. For the SVD approximation to be of rank r , we use the following version of Weyl's inequality (for proof see [14, Corollary 7.3.5]):

THEOREM 4.1 (Weyl's inequality for singular values). Let $A, B \in M_{m \times n}$. If $\|A - B\|_2 \leq \varepsilon$, then for $1 \leq k \leq \min(m, n)$, $|\sigma_k(A) - \sigma_k(B)| \leq \varepsilon$.

COROLLARY 4.2 If $\|A - U\Sigma V^*\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1}(A))$, then $\|A - U[\Sigma]_r V^*\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1}(A))$, where $[\Sigma]_r$ is the best rank- r approximation of Σ .

Proof.

$$\begin{aligned}
 \|A - U[\Sigma]_r V^*\|_2 &= \|A - U\Sigma V^* + U\Sigma V^* - U[\Sigma]_r V^*\|_2 \\
 &\leq \|A - U\Sigma V^*\|_2 + \|U\Sigma V^* - U[\Sigma]_r V^*\|_2 \\
 &\leq \mathcal{O}_\sigma(\sigma_{r+1}(A)) + \sigma_{r+1}(B) \\
 &\leq \mathcal{O}_\sigma(\sigma_{r+1}(A)) + \sigma_{r+1}(A) + \mathcal{O}_\sigma(\sigma_{r+1}(A)) \\
 &= \mathcal{O}_\sigma(\sigma_{r+1}(A)).
 \end{aligned}$$

□

Algorithm 4.1 describes a randomized SVD decomposition which provides a rank- r approximation. This approximation has an error $\mathcal{O}_\sigma(\sigma_{r+1}(A))$. Theorem 4.3 proves that the algorithm is correct for any matrix distribution that satisfies the conditions of Theorem 3.6. Its complexity is evaluated in Section 4.1.1. Numerical results are given in Section 5.

Algorithm 4.1 Sub-Gaussian-based Randomized SVD

Input: Matrix A of size $m \times n$ to decompose, r desired rank, k_1, k_2, l number of columns to use.

Output: Matrices U, Σ, V such that $\|A - U\Sigma V^*\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1}(A))$, where U and V are matrices with orthonormal columns and Σ is a diagonal matrix.

- 1: Create a random sub-Gaussian matrix Ω_1 of size $k_1 \times n$.
 - 2: Create a random Gaussian matrix Ω'_1 of size $l \times k_1$.
 - 3: Compute $B = A\Omega_1^* \Omega'^{*}_1$ ($B \in M_{m \times l}$).
 - 4: Compute the QR decomposition: $B = QR$, $Q \in M_{m \times l}$ with orthonormal columns, $R \in M_{l \times l}$ is a full rank upper triangular matrix.
 - 5: Create a random sub-Gaussian matrix Ω_2 of size $k_2 \times m$.
 - 6: Compute $\Omega_2 Q$, $\Omega_2 A$ and $(\Omega_2 Q)^\dagger$.
 - 7: Compute the SVD of $(\Omega_2 Q)^\dagger \Omega_2 A = \tilde{U}_1 \Sigma_1 V_1^*$.
 - 8: $U_1 \leftarrow Q \tilde{U}_1$.
 - 9: $U \leftarrow U_1(:, 1:r)$.
 - 10: $\Sigma \leftarrow \Sigma_1(1:r, 1:r)$.
 - 11: $V \leftarrow V_1(:, 1:r)$.
-

THEOREM 4.3 Assume that A is a matrix of size $m \times n$ where $m < n$, and $r < m$. Then for $k_1, k_2 = \mathcal{O}\left(\log\left(\frac{1}{p}\right)\left(r + \frac{1}{p^3}\right)\right)$ and $l = \mathcal{O}(r)$, Algorithm 4.1 outputs, with high probability, matrices U, Σ and V such that $\|A - U\Sigma V^*\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1}(A))$.

Proof. For a matrix $A \in M_{m \times n}$, let $\Omega_1 \in M_{k_1 \times n}$ be a sub-Gaussian matrix and let $\Omega'_1 \in M_{l \times k_1}$ be a random Gaussian matrix. Denote the QR-decomposition of $A\Omega_1^* \Omega_1'^* \in M_{m \times l}$ by $QR = A\Omega_1^* \Omega_1'^*$. From Theorem 3.6, Remark 3.7 and Theorem 10.8 of [11], it follows that for $l = \mathcal{O}(r)$

$$\|QQ^*A - A\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1}). \quad (4.1)$$

From Theorem 3.5 it follows that for a sub-Gaussian matrix $\Omega_2 \in M_{k_2 \times m}$, where $k_2 = \mathcal{O}(k_1)$, the matrix $\Omega_2 Q$ is left invertible with as high probability as needed, namely $(\Omega_2 Q)^\dagger \Omega_2 Q = I_{k_1 \times k_1}$. Thus, $\|QQ^*A - A\|_2 = \|Q(\Omega_2 Q)^\dagger (\Omega_2 Q)Q^*A - A\|_2$. By construction, $\|U_1 \Sigma_1 V_1^* - A\|_2 = \|Q\tilde{U}_1 \Sigma_1 V_1^* - A\|_2 = \|Q(\Omega_2 Q)^\dagger \Omega_2 A - A\|_2$. $\|Q(\Omega_2 Q)^\dagger \Omega_2 A - A\|_2$ is bounded as follows:

$$\begin{aligned} \|Q(\Omega_2 Q)^\dagger \Omega_2 A - A\|_2 &= \|Q(\Omega_2 Q)^\dagger \Omega_2 A - Q(\Omega_2 Q)^\dagger (\Omega_2 Q)Q^*A \\ &\quad + Q(\Omega_2 Q)^\dagger (\Omega_2 Q)Q^*A - A\|_2 \\ &= \|Q(\Omega_2 Q)^\dagger \Omega_2 (A - QQ^*A) + QQ^*A - A\|_2 \\ \text{by the triangle inequality} &\leq \|Q(\Omega_2 Q)^\dagger \Omega_2 (A - QQ^*A)\|_2 + \|QQ^*A - A\|_2 \\ \text{since } \|AB\|_2 &\leq \|A\|_2 \|B\|_2 \leq \|Q(\Omega_2 Q)^\dagger \Omega_2\|_2 \|A - QQ^*A\|_2 + \|QQ^*A - A\|_2 \\ &= ((\Omega_2 Q)^\dagger \Omega_2)_2 + 1 \|A - QQ^*A\|_2 \\ \text{since } \|AB\|_2 &\leq \|A\|_2 \|B\|_2 \leq ((\Omega_2 Q)^\dagger \Omega_2)_2 \|A - QQ^*A\|_2. \end{aligned} \quad (4.2)$$

Theorem 3.5 shows that $\|(\Omega_2 Q)^\dagger\|_2 \leq 1/(c_1 \sqrt{k_2})$ with high probability, which depends on the constant in the asymptotics of k_1 . From Theorem 2.4 we deduce that $\|\Omega_2\|_2 \leq C_0 \sqrt{n}$ with high probability. The probability can be made as close to 1 as required by altering the parameter C_0 . Thus, from Eq. (4.2) it follows that

$$\|Q(\Omega_2 Q)^\dagger \Omega_2 A - A\|_2 \leq \left(\frac{C_0}{c_1} \sqrt{\frac{n}{k_2}} + 1\right) \|A - QQ^*A\|_2.$$

In conjunction with Eq. (4.1), this yields that $\|U_1 \Sigma_1 V_1^* - A\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1})$.

From Corollary 4.2 it follows that $\|U\Sigma V^* - A\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1})$. $\square$

Remark. A bound for the Frobenius norm $\|A - U\Sigma V^*\|_F \leq \mathcal{O}_\sigma(\Delta_{r+1}(A))$ is obtained similarly by using Eq. (3.9).

4.1.1 Computational complexity of Algorithm 4.1 For computational complexity estimation and implementation, the internal random matrix distribution of the algorithm is selected as a subclass of sparse sub-Gaussian matrices. We chose sparse-Gaussian matrices. Sparse-Gaussian matrices are sparse matrices, each entry of which is i.i.d. with probability $1 - p$ to be zero and standard Gaussian otherwise. The complexity of each step in Algorithm 4.1 is shown in Table 4.1.

TABLE 4.1 *Complexity of Algorithm 4.1*

Step in Algorithm 4.1	A sparse	A dense
Creation of sparse matrix Ω_1 of size $k_1 \times n$	$\mathcal{O}(n)$	$\mathcal{O}(n)$
Creation of Gaussian matrix Ω'_1 of size $l \times k_1$	$\mathcal{O}(k_1 l)$	$\mathcal{O}(k_1 l)$
Computation of $B = A\Omega_1^* \Omega'^{*}_1$	$\mathcal{O}(\text{nnz}(A)pk_1 + mk_1 l)$	$\mathcal{O}(mnpk_1 + mk_1 l)$
Computation of its QR-decomposition, $B = QR$	$\mathcal{O}(mk_1^2)$	$\mathcal{O}(mk_1^2)$
Creation of sparse matrix Ω_2 of size $k_2 \times m$	$\mathcal{O}(m)$	$\mathcal{O}(m)$
Computation of $\Omega_2 Q$, $\Omega_2 A$	$\mathcal{O}(mk_2 + \text{nnz}(A)pk_2)$	$\mathcal{O}(mk_2 + mnpk_2)$
Computation of $(\Omega_2 Q)^\dagger \Omega_2 A$	$\mathcal{O}(k_1 k_2^2 + nk_2)$	$\mathcal{O}(k_1 k_2^2 + nk_2)$
Computation of the SVD of $(\Omega_2 Q)^\dagger \Omega_2 A$	$\mathcal{O}(nk_2^2)$	$\mathcal{O}(nk_2^2)$

The total complexity is $\mathcal{O}(\text{nnz}(A)pk + (m + n)k^2)$ for $k = \max(k_1, k_2)$. Note that k_1 and k_2 are of the same asymptotic order, but since Ω_2 approximates a subspace of dimension $l = \mathcal{O}(r)$ and Ω_1 approximates a subspace of dimension exactly r , we choose $k_1 < k_2$. For example, for sub-Gaussian random matrices with $p = \mathcal{O}(r^{-1/3})$ and $k = \mathcal{O}(r \log r)$, the complexity is $\mathcal{O}(\text{nnz}(A)r^{2/3} \log r + (m + n)(r \log r)^2)$.

For the OSE defined in [22], the asymptotic complexity is the same as in [22]. We show in Section 5 that although the asymptotic complexity is the same, Algorithm 4.1 is faster in practice.

REMARK 4.4 Classical SVD algorithms are iterative. In some sense, they are also approximate. Nevertheless, they converge extremely fast to machine precision, and thus it is valid to say that the complexity of SVD of an $n \times n$ matrix is $\mathcal{O}(n^3)$. See [30] for further details on classical SVD algorithms.

4.1.2 Numerical stability We argue that Algorithm 4.1 is stable since all its individual steps in the algorithm are stable. Indeed, the matrix multiplication steps are obviously stable. The SVD and QR decompositions, which appear in steps 4 and 7, are also stable (if appropriate algorithms are used, see e.g. [12]). The last part that should be examined is the computation of the pseudo-inverse in step 6 of Algorithm 4.1. Since we show in Theorem 3.5 that with high probability the norm of $(\Omega_2 Q)^\dagger$ is bounded, the calculation of the pseudo-inverse is also numerically stable.

4.2 Sub-Gaussian-based randomized LU decomposition

Theorem 3.6 is equivalent to Theorem 3.1 in [1] where the L_2 norm is used instead of using the Frobenius norm. A sub-Gaussian distribution can be used instead of the SEM distribution. Since the proof of the correctness of the algorithm in [1] is based on Theorem 3.1, it is also applicable for sub-Gaussian matrices.

THEOREM 4.5 Assume that sub-Gaussian random matrices are used instead of SEMs in the approximated LU decomposition of rank- r in [1]. Then, for any $r \in \mathcal{N}$, and for any matrix $A \in M_{m \times n}$, the approximated LU decomposition of rank- r results in matrices L and U and permutations P and Q such that $\|PAQ - LU\|_2 \leq \mathcal{O}_\sigma(\sigma_{r+1}(A))$.

The complexity of the algorithm in [1] is $\mathcal{O}(\text{nnz}(A)pk + (m + n)k^2)$.

TABLE 5.1 Algorithm 4.1 applied with different parameters to a sparse matrix A of size 5000×5000 with $\text{nnz}(A) = 36683$. Each row denotes a different k_1 and k_2 . Each column represents a different number of non-zero entries in a line of the random projection that appears in Algorithm 4.1. The numbers in the table are the approximation error where the best possible error is $\sigma_{101} \approx 0.0067$. The colour represents the running time, as the running time increases, the colour becomes darker, where $\text{greyscale} = 45 \times \text{running time in seconds}$.

k_1	k_2	projection sparsity (nnz in each row)				
		2	3	4	10	20
200	300	0.7921	0.0227	0.0184	0.0174	0.0176
200	500	0.0328	0.0156	0.0134	0.0122	0.0115
200	700	0.5086	0.0129	0.0115	0.0110	0.0108
300	400	0.0228	0.0101	0.0085	0.0080	0.0080
300	600	0.0134	0.0070	0.0068	0.0068	0.0067
300	900	0.2859	0.0076	0.0067	0.0067	0.0067
500	600	0.0080	0.0067	0.0067	0.0067	0.0067
500	900	0.0067	0.0067	0.0067	0.0067	0.0067

5. Numerical results

The results in this paper are valid for all types of i.i.d. sub-Gaussian matrices and OSE distributions. In the current implementation, whose software is available as supplementary materials, we used sparse-Gaussian matrices where each entry in the matrix is i.i.d. with probability p to be standard Gaussian and zero otherwise. Note that this distribution is like the distribution in Observation 2.4 up to a multiplicative constant that does not affect Algorithm 4.1. Other sub-Gaussian random variables (e.g. Rademacher) were tested and provided similar results.

5.1 Parameter selection

Since the bounds in this paper are asymptotic, we show in this section the influence of the parameters on the results. The matrix A in this section is real of size 5000×5000 in double precision. The first 100 singular values of A are 1, and the other decay from e^{-5} to e^{-50} , which means that $\sigma_{101}(A) \approx 0.0067$. We show the obtained results when we use projections of different sizes and different sparsity. The colour of a cell in Tables 5.1 and 5.2 represents the running time of Algorithm 4.1 such that $\text{greyscale} = 50 \times \text{running time in seconds}$. Table 5.1 shows the results of an experiment done on a sparse matrix A with about 10^5 non-zeros, and Table 5.2 shows the results of an experiment done on a dense matrix A .

One can notice that for sparse matrices the number of non-zeros in a row affects the accuracy significantly. One non-zero in a row, as suggested in [5], requires large k s and has stability issues for too small k s. Even two non-zero entries degrade the accuracy significantly, or require large k s that results in a longer running time. Thus, it is better to choose at least a few non-zeros in each row in the projection, as was also suggested in [22].

TABLE 5.2 Algorithm 4.1 applied with different parameters to a dense matrix A of size 5000×5000 . The numbers in the table are the errors of the approximation where the best possible error is ≈ 0.0067 . The colour represents the running time, as the running time increases, the colour becomes darker, where greyscale = $45 \times$ running time in seconds

k_1	k_2	projection sparsity (nnz in each row)				
		2	3	4	10	20
200	300	0.0165	0.0167	0.0163	0.0166	0.0169
200	500	0.0115	0.0116	0.0116	0.0114	0.0116
200	700	0.0103	0.0106	0.0105	0.0106	0.0107
300	400	0.0077	0.0078	0.0077	0.0079	0.0080
300	600	0.0068	0.0067	0.0067	0.0067	0.0068
300	900	0.0067	0.0067	0.0067	0.0067	0.0067
500	600	0.0067	0.0067	0.0067	0.0067	0.0067
500	900	0.0067	0.0067	0.0067	0.0067	0.0067

5.2 Performance comparison on random matrices with rapidly decaying singular values

We describe the results from three different experiments. All the experiments were implemented in Matlab on a single core of Intel Xeon CPU X5560 2.8GHz. Although the current Matlab implementations are not optimal, they demonstrate the performance when they are compared to a Matlab code. All the following experiments compare between the running time and the generated error of the following three algorithms in different scenarios: 1. The FFT-based algorithm given in [36]; 2. The Algorithm from [5] and 3. Algorithm 4.1 with three non-zero entries in each row. Although the proven error bounds for Algorithm 4.1 are less tight than the bounds in the other algorithms, we see that in practice Algorithm 4.1 reaches the same error. In all the experiments, the parameters for the different algorithms are chosen such that the reconstruction error rates are similar and aligned to the error from [5] and [36]. The slowest algorithm has an error that is not smaller than that of the fastest algorithm.

The experiments that were carried out are as follows:

1. Rank- r approximation is computed for a randomly generated full matrix $A \in M_{3000 \times 3000}$ with singular values that decay exponentially fast from 1 to e^{-50} . Figure 5.1 displays the comparison between the running time and the error from the rank- r approximation of the three algorithms mentioned above. The x-axis denotes the rank and the y-axis denotes the running time. The results show that for a small rank range [5] is faster than the FFT-based algorithm [36]. For a larger rank range, the FFT-based algorithm is faster. For all ranks, Algorithm 4.1 is the fastest.
2. Rank- r approximation is computed for a randomly generated full matrix $A \in M_{3000 \times 3000}$, where the first r singular values are 1 and the other singular values decay exponentially fast from e^{-5} to e^{-50} . Figure 5.2a displays the comparison between the running time for rank- r approximation of the three algorithms mentioned above. The x-axis denotes the rank and y-axis denotes the running time. As in experiment 1, for a small rank range, [5] is faster than the FFT-based algorithm [36]. For a higher rank range, the FFT-based algorithm is faster than [5]. For all ranks, Algorithm 4.1 is the fastest.

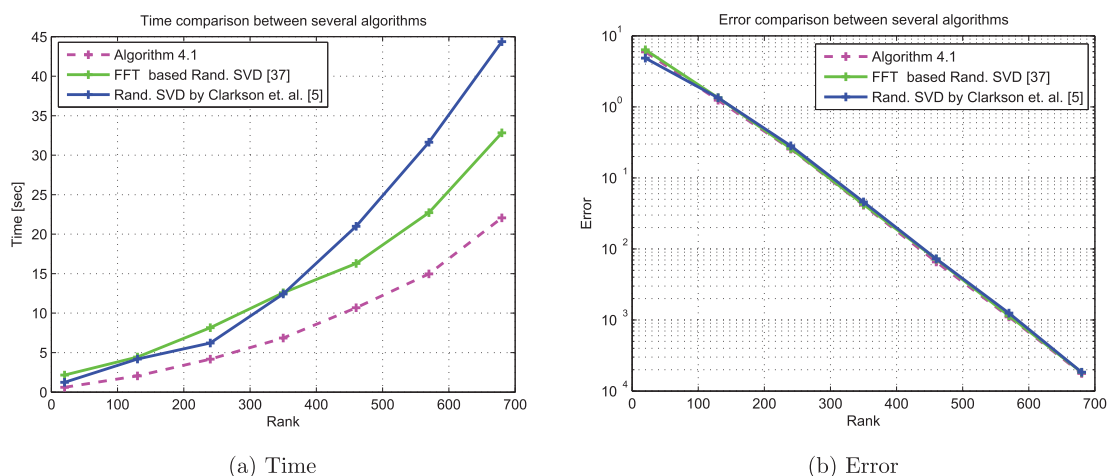


FIG. 5.1. Results from the approximation of a matrix of size 3000×3000 with exponentially decaying singular values. The x-axes in (a) and (b) denote the rank of the approximation. The y-axis in (a) denotes the running time. The y-axis in (b) denotes the error from the rank approximation measured in the operator norm.

3. Rank-300 approximation of a randomly generated full matrix $A \in M_{n \times n}$ is computed when the first 300 singular values are 1 and the other singular values decay exponentially fast from e^{-5} to e^{-50} . Figure 5.3 displays the comparison between the run time for rank-300 approximation of the three algorithms mentioned above. The x-axis denotes the rank and y-axis denotes the running time. This experiment shows that the computation of the sparse SVD in [5] is faster

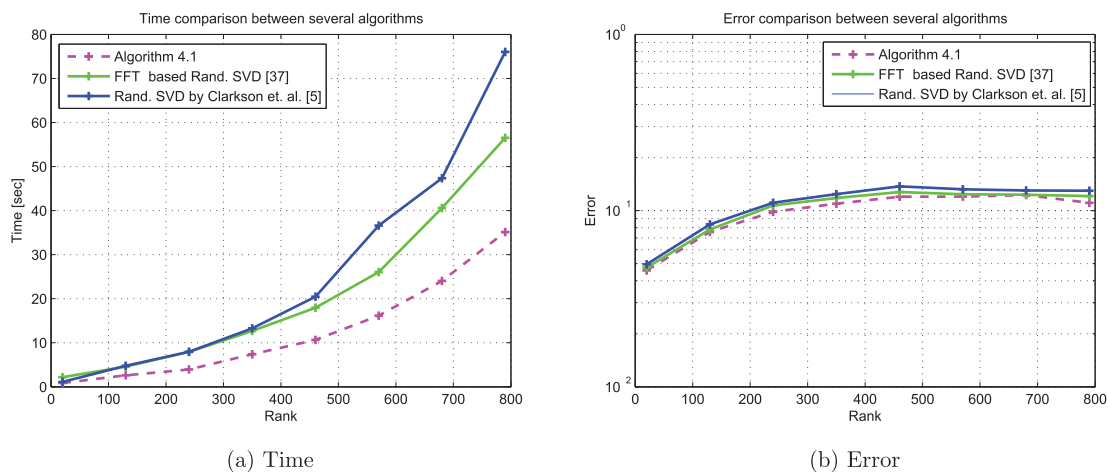


FIG. 5.2. Results from the approximation of a matrix of size 3000×3000 with different numerical ranks. The x-axes in (a) and (b) denote the numerical rank. The y-axis in (a) denotes the running time. The y-axis in (b) denotes the rank approximation error measured in the operator norm.

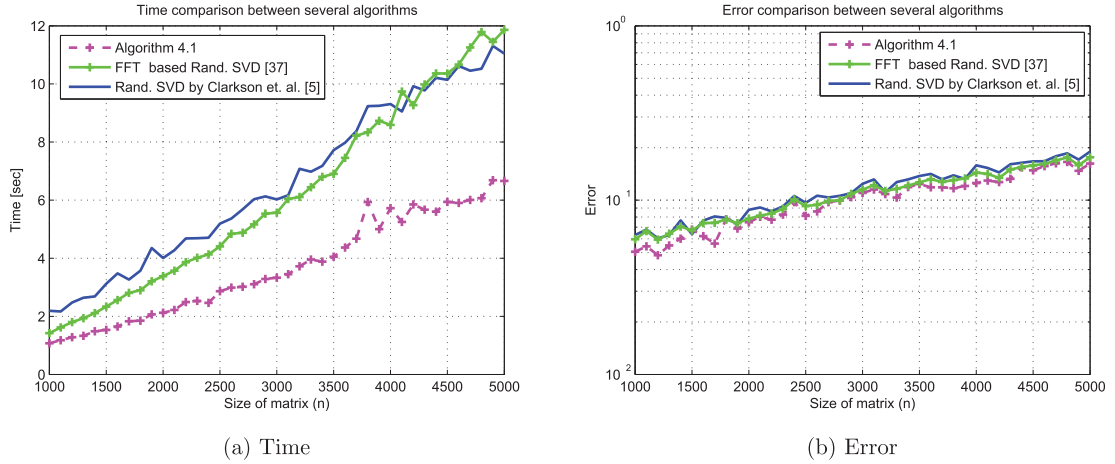


FIG. 5.3. Results from the approximation of a matrix of size $n \times n$, $n = 1000, \dots, 5000$ with numerical rank 300. The x-axis in (a) and (b) denotes n . The y-axis in (a) denotes the running time. The y-axis in (b) denotes the approximation error measured in the operator norm.

than the computation of the FFT-based algorithm [36] when n increases. For rank 300 and for $n \approx 4500$, the algorithm from [5] is faster than the FFT-based algorithm. For ranks higher than 300, a large n is required for the algorithm from [5] to be faster than the FFT-based algorithm. The Sparse SVD Algorithm 4.1 presented in this paper is the fastest for all n .

In Algorithm 4.1, it is only necessary to apply the matrix A once from the left and once from the right; therefore, A does not have to be stored in memory. Table 5.3 shows the running time for large matrices that cannot be stored in a computer memory. The matrices we chose have a similar form to the choice in [10]. We chose $A = F \Sigma F$, where F is the DFT matrix and Σ is a diagonal matrix with singular values σ_i that decay linearly until $i = 200$ and exponentially from thereon. In this experiment, we set $\sum_{i=201}^n \sigma_i$ to be constant. Algorithm 4.1 is applied to rank 200 with $k_1 = 500$ and $k_2 = 700$.

5.3 A large, sparse and noisy matrix arising in image processing

Next, we follow the lines of the experiment done in Chapter 7.2 of [11], and provide a low-rank approximation to the graph Laplacian associated with an image. We begin with a grayscale image of size 120×120 . We form for each pixel i a vector $x_i \in \mathbb{R}^{25}$ by gathering the intensities of the pixels in a 5×5 neighbourhood centred at pixel i (near the edges, x_i is padded with zeros). Next, we form the 14400×14400 weight matrix W that reflects the similarities between patches such that

$$w_{ij} = \exp \left\{ -\|x_i - x_j\|^2 / \sigma^2 \right\}.$$

The objective is to construct the low frequency eigenvectors of the graph Laplacian matrix

$$L = I - D^{-1/2} W D^{-1/2},$$

TABLE 5.3 Comparing running time of Algorithm 4.1 with the standard SVD that are applied to large matrices of size $n \times n$. The Relative Error is the ratio between the error from the rank- r decomposition and from the $(r + 1)$'s singular value. We were unable to apply the classical SVD algorithm to $n \geq 32,768$

Size (n)	Relative error from Algorithm 4.1	Time for Algorithm 4.1 (sec)	Time for full SVD
1,024	1.5465	1.0011	1.5232
2,048	1.5645	1.6236	11.3702
4,096	1.5422	2.7653	94.6345
8,192	1.5571	5.2999	578.2982
16,384	1.4846	12.1065	4324.683
32,768	1.5686	26.4022	
65,536	1.5074	50.6191	
131,072	1.4838	109.8185	
262,144	1.5357	205.0068	
524,288	1.4854	418.4137	
1,048,576	1.5240	847.8211	

where I is the identity matrix and D is the diagonal matrix with entries $d_{ii} = \sum_j w_{ij}$. These are the eigenvectors associated with the dominant eigenvalues of the auxiliary matrix $A = D^{-1/2}WD^{-1/2}$. Like in [11], we approximate the eigenvalues of the matrix A . Since the eigenvalues of A decay slowly, in [11] a power scheme was used. In other words, the constructed column space is based on $(AA^*)A$ (instead of just A), this improves the accuracy significantly in the case of slowly decaying spectrum. For the sake of comparison of Algorithm 4.1 to the FFT-based randomized SVD of [36], and the algorithm of [5], although it does not make ‘computational sense’, we applied all the algorithms to the matrix $(AA^*)A$.

The approximation of A by the results of Algorithm 4.1, the FFT-based algorithm [36] and the sparse SVD in [5] are shown in Fig. 5.4. This experiment was averaged five times for ranks in $[10, 2000]$. In Fig. 5.4a, the running times are compared and in Fig. 5.4b the accuracies are compared. Here, as before, Algorithm 4.1 outperforms the alternatives. In Fig. 5.4b, the gap between the approximation error of Algorithm 4.1 and the optimal approximation (the singular values) is very small in low ranks, and gets bigger for ranks that approach 2000. This is due to the rate of decay of the singular values. In order to get better approximation in reigns where the singular values decay slowly, higher orders in the power scheme are required.

Conclusion

We showed that matrices with i.i.d. sub-Gaussian entries conserve subspaces and the connection between the distribution of the entries and the required size of the matrix established. A new algorithm is presented, which yields with high probability, a rank- r SVD approximation for an $m \times n$ matrix that achieves an asymptotic complexity of $\mathcal{O}(\text{nnz}(A)pk + (m + n)k^2)$. Additionally, we showed that the approximated LU algorithm in [1], which uses sub-Gaussian random matrices, has a computational complexity of $\mathcal{O}(\text{nnz}(A)pk + (m + n)k^2)$. We demonstrated in the experiments that, although the derived error bounds are not as tight as the bounds from the algorithms in [5, 11], in practice, the algorithm in this paper reaches the same error in less time.

Future work includes non-asymptotic estimation of the algorithm parameters, including error estimation improvement to get tighter bounds.

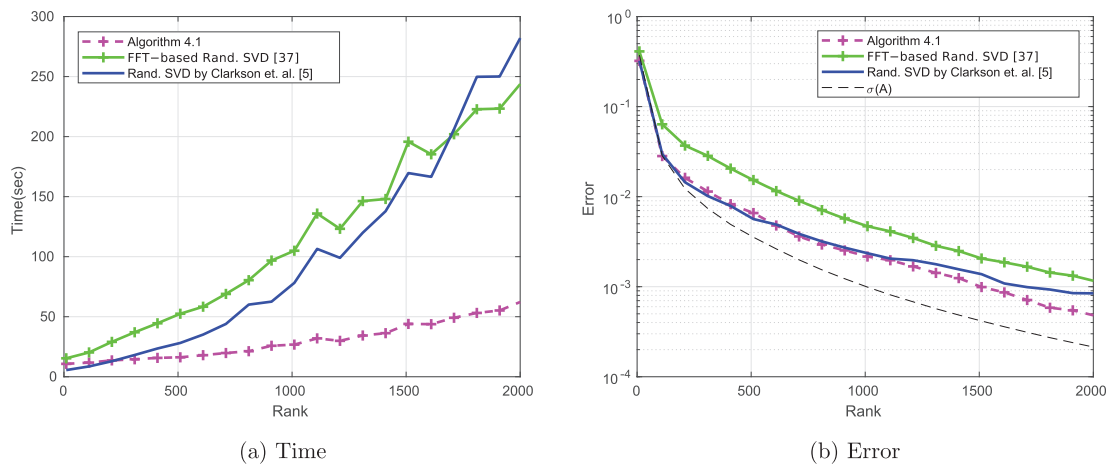


FIG. 5.4. Low-rank approximation of a graph Laplacian of an image for different ranks (x-axis). The approximated matrix is of size 14400×14400 . The accuracies of the FFT-based algorithm [36], the algorithm of [5] and Algorithm 4.1 are shown in Fig. 5.4b together with the optimal error (dotted line). In Fig. 5.4a, the y-axis denotes the running time of different algorithms.

Acknowledgments

We thank Prof. Jelani Nelson and Dr. Haim Avron for their continuous support and constructive remarks. We also thank Prof. Joel Tropp for his remarks regarding numerical stability in a previous version of this paper.

Supplementary material

Matlab implementation of Algorithm 4.1 is located in:
<http://www.math.tau.ac.il/~aizeny/publications.html>

Funding

Israeli Ministry of Science & Technology Indian R&D (Grant No. 1); Israel Science Foundation (ISF 1556/17); Blavatnik Computer Science Research Fund; Blavatnik ICRC Funds; and Fellowships from Jyväskylä University and the Clore Foundation.

REFERENCES

1. AIZENBUD, Y., SHABAT, G. & AVERBUCH, A. (2016) Randomized LU decomposition using sparse projections *Comput. Math. Appl.*, **72**, 2525–2534.
2. AVRON, H., MAYMOUNKOV, P. & TOLEDO, S. (2010) Blendenpik: supercharging lapack's least-squares solver. *SIAM J. Sci. Comput.*, **32**, 1217–1236.
3. BERRY, A. C. (1941) The accuracy of the gaussian approximation to the sum of independent variates. *Trans. Amer. Math. Soc.*, **49**, 122–136.
4. CANDÈS, E. J. & TAO, T. (2006) Near-optimal signal recovery from random projections: universal encoding strategies? *IEEE Trans. Inf. Theory*, **52**, 5406–5425.

5. CLARKSON, K. L. & WOODRUFF, D. P. (2013) Low rank approximation and regression in input sparsity time. *Proceedings of the 45th Annual ACM Symposium on Theory of Computing*. New York, NY, USA:ACM, pp. 81–90.
6. COHEN, M. B. (2016) Nearly tight oblivious subspace embeddings by trace inequalities. *Proceedings of the Twenty-Seventh Annual ACM-SIAM Symposium on Discrete Algorithms*, pp. 278–287.
7. DASGUPTA, A., KUMAR, R. & SARLÓS, T. (2010) A sparse johnson: Lindenstrauss transform. *Proceedings of the Forty-Second ACM Symposium on Theory of Computing*. New York, NY, USA:ACM, pp. 341–350.
8. DIRKSEN, S. (2014) Dimensionality reduction with subgaussian matrices: a unified theory. *Foundations of Computational Mathematics*. arXiv preprint arXiv:1402.3973, **16**, 1367–1396.
9. DONOHO, D. L. (2006) Compressed sensing. *IEEE Trans. Inf. Theory*, **52**, 1289–1306.
10. HALKO, N., MARTINSSON, P.-G., SHKOLNISKY, Y. & TYGERT, M. (2011) An algorithm for the principal component analysis of large data sets. *SIAM J. Sci. Comput.*, **33**, 2580–2594.
11. HALKO, N., MARTINSSON, P.-G. & TROPP, J. A. (2011) Finding structure with randomness: probabilistic algorithms for constructing approximate matrix decompositions. *SIAM Review*, **53**, 217–288.
12. HIGHAM, N. J. (2002) *Accuracy and Stability of Numerical Algorithms*. SIAM.
13. Hoeffding, W. (1963) Probability inequalities for sums of bounded random variables *J. Amer. Stat. Assoc.*, **58**, 13–30.
14. HORN, R. A. & JOHNSON, C. R. (1990) *Matrix Analysis*. New York, NY, USA:Cambridge University Press.
15. JOHNSON, W. B. & LINDENSTRAUSS, J. (1984) Extensions of Lipschitz mappings into a Hilbert space. *Contemp. Math.*, **26**, 1.
16. KANE, D. M. & NELSON, J. (2014) Sparser Johnson–Lindenstrauss transforms. *Journal of the ACM*, **61**, 4.
17. LITVAK, A. E., PAJOR, A., RUDELSON, M. & TOMCZAK-JAEGERMANN, N. (2005) Smallest singular value of random matrices and geometry of random polytopes. *Adv. Math.*, **195**, 491–523.
18. MARTINSSON, P.-G., ROKHLIN, V. & TYGERT, M. (2011) A randomized algorithm for the decomposition of matrices. *Appl. Comput. Harmon. Anal.*, **30**, 47–68.
19. MENDELSON, S., PAJOR, A. & TOMCZAK-JAEGERMANN, N. (2005) Reconstruction and subgaussian processes. *C. R. Math. Acad. Sci. Soc. R. Can.*, **340**, 885–888.
20. MENDELSON, S., PAJOR, A. & TOMCZAK-JAEGERMANN, N. (2007) Reconstruction and subgaussian operators in asymptotic geometric analysis. *Geom. Funct. Anal.*, **17**, 1248–1282.
21. MENDELSON, S., PAJOR, A. & TOMCZAK-JAEGERMANN, N. (2008) Uniform uncertainty principle for bernoulli and subgaussian ensembles. *Constr. Approx.*, **28**, 277–289.
22. NELSON, J. & NGUYEN, H. L. (2013) Osnap: faster numerical linear algebra algorithms via sparser subspace embeddings. *IEEE 54th Annual Symposium on Foundations of Computer Science (FOCS)*. IEEE, pp. 117–126.
23. NELSON, J. & NGUYEN, H. L. (2013) Sparsity lower bounds for dimensionality reducing maps. *Proceedings of the Forty-Fifth Annual ACM Symposium on Theory of Computing*. New York, NY, USA:ACM, pp. 101–110.
24. ROKHLIN, V. & TYGERT, M. (2008) A fast randomized algorithm for overdetermined linear least-squares regression. *Proc. Natl. Acad. Sci.*, **105**, 13212–13217.
25. RUDELSON, M. (2014) Recent developments in non-asymptotic theory of random matrices. *Mod. Aspects Random Matrix Theory*, vol. **72**. p. 83.
26. RUDELSON, M. & VERSHYNIN, R. (2008) The Littlewood–Offord problem and invertibility of random matrices. *Adv. Math.*, **218**, 600–633.
27. RUDELSON, M. & VERSHYNIN, R. (2009) Smallest singular value of a random rectangular matrix. *Commun. Pure Appl. Math.*, **62**, 1707–1739.
28. SARLOS, T. (2006) Improved approximation algorithms for large matrices via random projections. *2006 47th Annual IEEE Symposium on Foundations of Computer Science, FOCS'06*. IEEE, pp. 143–152.
29. SHABAT, G., SHMUELI, Y., AIZENBUD, Y. & AVERBUCH, A. (2018) Randomized LU decomposition. *Applied and Computational Harmonic Analysis*, **44**, 246–272, arXiv preprint arXiv:1310.7202.
30. STEWART, G. W. (2001) *Matrix Algorithms Volume 2: Eigensystems*, vol. **2**. SIAM.

31. TAO, T. (2012) *Topics in Random Matrix Theory*, vol. 132 of Graduate Studies in Mathematics. Providence, RI: American Mathematical Soc.
32. TROPP, J. A. (2011) Improved analysis of the subsampled randomized Hadamard transform. *Adv. Adapt. Data Anal.*, **3**, 115–126.
33. URANO, Y. (2013) A fast randomized algorithm for linear least-squares regression via sparse transforms, *Master's Thesis*, New York, USA: New York University, .
34. VERSHYNIN, R (2010) *Introduction to the non-asymptotic analysis of random matrices*, arXiv preprint arXiv:1011.3027.
35. WITTEN, R. & CANDÈS, E. (2015) Randomized algorithms for low-rank matrix factorizations: sharp performance bounds. *Algorithmica*, **72**, 264–281.
36. WOOLFE, F., LIBERTY, E., ROKHLIN, V. & TYGERT, M. (2008) A fast randomized algorithm for the approximation of matrices. *Appl. Comput. Harmon. Anal.*, **25**, 335–366.