

Robust CUR Decomposition: Theory and Imaging Applications*

HanQin Cai[†], Keaton Hamm[‡], Longxiu Huang[†], and Deanna Needell[†]

Abstract. This paper considers the use of robust principal component analysis (RPCA) in a CUR decomposition framework and applications thereof. Our main algorithms produce a robust version of column-row factorizations of matrices $D = L + S$, where L is low-rank and S contains sparse outliers. These methods yield interpretable factorizations at low computational cost and provide new CUR decompositions that are robust to sparse outliers, in contrast to previous methods. We consider two key imaging applications of RPCA: video foreground-background separation and face modeling. This paper examines the qualitative behavior of our robust CUR decompositions on the benchmark videos and face datasets and finds that our method works as well as standard RPCA while being significantly faster. Additionally, we consider hybrid randomized and deterministic sampling methods which produce a compact CUR decomposition of a given matrix and apply this to video sequences to produce canonical frames thereof.

Key words. CUR decomposition, RPCA, robust CUR, low-rank matrix approximation, interpolative decompositions, robust algorithms

AMS subject classifications. 15A23, 65F30, 68P20, 68W20, 68W25, 68Q25

DOI. 10.1137/20M1388322

1. Introduction. The scale of data being collected is increasingly massive, which requires better and more computationally efficient tools to handle associated data matrices. One method of handling large data matrices is to extract a meaningful column submatrix and to either represent the full data matrix via a factorization with respect to the columns or to carry out any desired analysis on the column submatrix and then reconstruct the full data matrix from the columns. This technique is motivated by some foundational results that tell one when and how well column submatrices can represent a data matrix. In particular, for a low-rank matrix $L \in \mathbb{R}^{m \times n}$ with $\text{rank}(L) = r$, one can write $L = CX$, where C is a (much smaller) column submatrix of L as long as $\text{rank}(C) = \text{rank}(L)$. Hence, one obtains a representation of L in terms of a few representative data vectors.

*Received by the editors December 28, 2020; accepted for publication (in revised form) July 20, 2021; published electronically October 18, 2021.

<https://doi.org/10.1137/20M1388322>

Funding: The work of the second author was partially supported by the Army Research Office under grant number W911NF-20-1-0076. The views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies, either expressed or implied, of the Army Research Office or the U.S. Government. The U.S. Government is authorized to reproduce and distribute reprints for Government purposes notwithstanding any copyright notation herein. The work of the third and fourth authors was supported by NSF BIGDATA DMS-1740325 and NSF DMS-2011140.

[†]Department of Mathematics, University of California, Los Angeles, Los Angeles, CA 90095 USA (hqcai@math.ucla.edu, huangl3@math.ucla.edu, deanna@math.ucla.edu).

[‡]Department of Mathematics, University of Texas at Arlington, Arlington, TX 76019 USA (keaton.hamm@uta.edu).

In general, such factorizations are called *interpolative decompositions* [50], while for particular choices of $\mathbf{X}$ including $\mathbf{X} = \mathbf{C}^\dagger \mathbf{L}$ or $\mathbf{X} = \mathbf{U}^\dagger \mathbf{R}$, where $\mathbf{R}$ is a row submatrix of $\mathbf{L}$ and $\mathbf{U}$ is the overlap of the column and row indices of $\mathbf{C}$ and $\mathbf{R}$, they are termed *CUR decompositions* [7, 20, 22, 33, 37, 44]. By some, the latter examples are called *cross-approximations* [40] or *(pseudo-)skeleton approximations* [16, 26, 27, 28]. The problem of finding the best set of k columns with which to approximate $\mathbf{L}$ is the column subset selection problem [6].

Since their advent, column factorizations such as these have been shown to have great flexibility and approximation power as a general low-rank approximation tool. Specifically, they provide advantages over traditional factorization methods. For example, (1) column factorizations can yield compact representations of large data matrices; (2) for incoherent matrices, the complexity of generating the uniform sampling pattern is $\mathcal{O}(1)$ ¹ while yielding a factorization which well approximates the SVD (which requires $\mathcal{O}(mnr)$ time for a rank r matrix); and (3) the representation of a matrix in a basis consisting of actual data vectors provides a more interpretable representation of the data than abstract factorizations such as QR or SVD.

In many image and video processing applications, the observed data can be interpreted as follows:

1. Low-rank information of interest with some sparse outliers. For example, in the task of face modeling [51], the face data of the same person in different photos form a low-rank matrix that we want to recover, while the facial occlusions (accessories, shadows, and facial expressions) are considered sparse outliers.
2. A combination of two pieces of information of interest that have low rank and sparse properties, respectively. For example, in the task of video background subtraction [34], the static background has the low-rank property, and the foreground (i.e., the moving objects) is sparse.

Robust principal component analysis (RPCA) is an apt tool for solving the aforementioned problems and various other tasks in this arena [9, 14, 52].

In the RPCA framework, we assume we observe $\mathbf{D} = \mathbf{L} + \mathbf{S} \in \mathbb{R}^{m \times n}$, where $\mathbf{L}$ is a low-rank matrix and $\mathbf{S}$ is a sparse matrix, and the goal is to recover $\mathbf{L}$. For large-scale data problems, $\mathbf{D}$ may not be able to be stored in memory. The purpose of this paper is to elucidate how one may use column factorizations to assist in the RPCA task of finding the low-rank piece. We elaborate on our main contributions in section 1.2 below, but first we present the main ideas utilized throughout the paper.

1.1. Main ideas. There are several points to consider when trying to marry the techniques of CUR decompositions and robust low-rank matrix recovery from sparse outliers. The main goal of this paper is to do so in such a way as to guarantee speed, interpretability, and robustness. The first two goals are commonly achieved by CUR decompositions (as opposed to SVD based methods), while the third is achieved by RPCA but not typically by CUR. We elaborate more on how the method proposed here achieves these aims in comparison with existing techniques in the following subsection.

Taking these factors into consideration, we propose the following broadly applicable method: subsample columns and rows of $\mathbf{D}$, apply an existing RPCA algorithm to these

¹Here and throughout, we take $\mathcal{O}(1)$ to mean asymptotically with respect to m and n .

matrices, and use the denoised column and row matrices to form a CUR decomposition of the underlying low-rank part, $\mathbf{L}$. This procedure has several advantages: (1) it is faster than standard RPCA as long as few columns and rows may be selected, which is guaranteed by some theoretical results on submatrices of incoherent matrices; (2) it can work on matrices that do not fit in memory because it first subsamples the data matrix; (3) it returns a column-row factorization that is known to be more interpretable than the SVD as it corresponds to representing data via other actual data points; (4) there are good theoretical guarantees for its performance; and (5) it is a flexible framework that allows the user to specify the RPCA algorithm utilized. The last point allows one to determine the tradeoff between robustness and computational cost existing in the various RPCA methods in use.

1.2. Main contributions. We consider two variations of algorithms for producing robust CUR decompositions. In particular, given $\mathbf{D} = \mathbf{L} + \mathbf{S}$, our first algorithm (Algorithm 3.1) uniformly randomly selects a column submatrix $\mathbf{C} = \mathbf{D}(:, J)$ and row submatrix $\mathbf{R} = \mathbf{D}(I, :)$ of $\mathbf{D}$, then performs RPCA on both of these matrices to produce $\hat{\mathbf{C}}$ and $\hat{\mathbf{R}}$ which approximate $\mathbf{L}(:, J)$ and $\mathbf{L}(I, :)$, respectively. Then we obtain an approximate CUR decomposition of the underlying low-rank matrix $\mathbf{L}$; namely, $\mathbf{L} \approx \hat{\mathbf{C}}(\hat{\mathbf{C}}(I, :))^{\dagger} \hat{\mathbf{R}} \approx \mathbf{L}(:, J)(\mathbf{L}(I, J))^{\dagger} \mathbf{L}(I, :)$. We call such decompositions robust CUR (RCUR) decompositions.

The theoretical contributions of this procedure are as follows:

1. We prove bounds for all of the relevant properties of RPCA algorithms (incoherence, sparsity, and condition number) of randomly selected column and row submatrices of $\mathbf{L}$ in terms of the corresponding properties of $\mathbf{L}$.
2. Our algorithms provide a novel robust version of CUR decompositions.
3. Our error analysis shows that RCUR decompositions are capable of well approximating low-rank matrices with sparse, but possibly large magnitude noise, a task for which standard CUR decompositions are not designed.
4. Our RCUR decomposition yields relative error spectral norm guarantees, which have not been shown previously for standard CUR decompositions.
5. The column/row subsampling procedure allows our algorithm to work on large data matrices which cannot be stored in memory and typically has less computational cost than RPCA on the whole data matrix.
6. Our algorithm is flexible, as it allows the user to specify the RPCA algorithm used and thus to specify the tradeoff between robustness and computational cost.

For independent interest, we also consider how properties such as incoherence pass to submatrices under the greedy column selection procedure of [3]. Under greedy column selection, one may have adversarial deterministic sparse noise which prevents the column submatrix from being sparse.

Our second main algorithm is a hybrid random + deterministic method for choosing exactly $\text{rank}(\mathbf{L})$ columns and rows of $\mathbf{D}$ to produce the most compact CUR decomposition of $\mathbf{L}$. We show that with high probability, this algorithm produces a good CUR approximation of $\mathbf{L}$ in the spectral norm.

1.3. Prior art. Both CUR decompositions and RPCA algorithms have been well developed in the literature. However, our work is one of the first to combine them. Randomized sampling algorithms for CUR decompositions and the column subset selection problem have

been studied in [16, 20, 22, 32, 37, 46, 47], for example, while deterministic sampling methods have been explored in [3, 5, 35], with hybrid methods explored in [6, 7]. For a general overview, of CUR decompositions and approximations, see [33]. In contrast to the current paper, CUR decompositions typically exhibit error guarantees in the Frobenius norm, while ours hold for the spectral norm. Additionally, typical upper bounds for CUR decompositions of matrices hold in the worst case when the matrix is full rank, and the CUR decomposition approximates the truncated SVD. Thus CUR decompositions are primarily a low-rank approximation tool. This is in contrast to RPCA algorithms, which are primarily concerned with estimating a low-rank matrix which is corrupted by sparse outliers, i.e., RPCA is a low-rank recovery tool.

The algorithms developed herein apply RPCA in a CUR decomposition framework [12, 14, 39]. Our Algorithm 3.1 is faster than merely applying a RPCA algorithm directly to $\mathbf{D}$ and is usable for matrices which are too large to store in memory; however, this comes at a small cost in robustness to extreme outliers. We note that the authors in [13] used CUR decompositions as an approximation to the SVD within an iterative RPCA framework; however, while that algorithm achieved empirical success, there is currently no theoretical result for the algorithm convergence or outlier toleration.

2. Preliminaries and layout. We will always assume that $\mathbf{L} = \mathbf{W}_L \mathbf{\Sigma}_L \mathbf{V}_L^T$ is the compact SVD of $\mathbf{L} \in \mathbb{R}^{m \times n}$. The symbol $[n]$ denotes the set of integers $\{1, \dots, n\}$ for any $n \in \mathbb{N}$. We use $\kappa(\mathbf{L}) := \|\mathbf{L}\|_2 \|\mathbf{L}^\dagger\|_2 = \frac{\sigma_{\max}(\mathbf{L})}{\sigma_{\min}(\mathbf{L})}$ to denote the spectral condition number of $\mathbf{L}$, where $\sigma_{\max}(\mathbf{L})$ and $\sigma_{\min}(\mathbf{L})$ are the maximum and minimum nonzero singular values of $\mathbf{L}$, respectively.

Below are two fundamental assumptions which are typically made in the RPCA framework, which we also utilize.

Definition 2.1. Let $\mathbf{L} \in \mathbb{R}^{m \times n}$ with $\text{rank}(\mathbf{L}) = r$, and let $\mathbf{L} = \mathbf{W}_L \mathbf{\Sigma}_L \mathbf{V}_L^T$ be its compact SVD. Then $\mathbf{L}$ has $\{\mu_1(\mathbf{L}), \mu_2(\mathbf{L})\}$ -incoherence (i.e., $\mu_1(\mathbf{L})$ -column incoherence and $\mu_2(\mathbf{L})$ -row incoherence) for some constants $1 \leq \mu_1(\mathbf{L}) \leq \frac{m}{r}$, $1 \leq \mu_2(\mathbf{L}) \leq \frac{n}{r}$ provided

$$(A1) \quad \max_i \|\mathbf{W}_L^T \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu_1(\mathbf{L})r}{m}} \quad \text{and} \quad \max_i \|\mathbf{V}_L^T \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu_2(\mathbf{L})r}{n}}.$$

Definition 2.2. Let $\mathbf{S} \in \mathbb{R}^{m \times n}$ be an α -sparse matrix, and let $\mathbf{e}_i$ be the i th canonical unit basis vector of the appropriate dimension. Then $\mathbf{S}$ has at most αn nonzero entries in each row and at most αm nonzero entries in each column. In the other words,

$$(A2) \quad \max_i \|\mathbf{S}^T \mathbf{e}_i\|_0 \leq \alpha n \quad \text{and} \quad \max_j \|\mathbf{S} \mathbf{e}_j\|_0 \leq \alpha m,$$

where $\|\cdot\|_0$ denotes the ℓ_0 -norm.

Remark 2.3. For a successful reconstruction, we expect $\mu_i(\mathbf{L}) = \mathcal{O}(1)$, $i = 1, 2$, which is generally true in real-world RPCA applications [9, 11, 12]. For the purpose of theoretical analysis, if $\mu_i(\mathbf{L})$ is too large, then the RPCA problem may become too difficult to solve since the toleration of α usually depends on $\text{poly}(\frac{1}{\mu_i(\mathbf{L})})$ [12, 14, 39, 53]. Additionally, the rank-sparsity uncertainty principle [15] states that a matrix cannot be simultaneously incoherent and sparse; that is, small $\mu_i(\mathbf{L})$ ensures $\mathbf{L}$ is not sparse.

2.1. Organization. The rest of the paper is laid out as follows: section 3 contains statements of our main results as well as our main algorithm; section 4 contains the proof of our results concerning how incoherence, sparsity, and condition number change via column selection (Theorem 3.1). Section 5 illustrates our methods via experiments on video background-foreground separation tasks. Finally, the technical details of proofs of the main results are contained in the appendices; specifically, the proof of Theorem 3.5 is contained in Appendix A, the combined error analysis of CUR decompositions and RPCA leading to the proof of Theorem 3.10 is in Appendix B, the proofs for the combined random and deterministic sampling result of Theorem 3.11 are in Appendix C, and the purely deterministic sampling results are contained in Appendices D and E.

3. Main results. To develop a framework of combining RPCA and CUR decompositions, one must first understand how properties of incoherence and sparsity are inherited by submatrices. This, in turn, requires estimating pseudoinverses of submatrices of singular vectors of $\mathbf{L}$, which is generally a difficult task.

Recall the first step in the RCUR decomposition described above: we select column and row submatrices of the observed data $\mathbf{D} = \mathbf{L} + \mathbf{S}$, $\tilde{\mathbf{C}} = \mathbf{D}(:, J)$, and $\tilde{\mathbf{R}} = \mathbf{D}(I, :)$. We then apply RPCA algorithms to $\tilde{\mathbf{C}}$, and $\tilde{\mathbf{R}}$ to find “good” approximations for the submatrices $\mathbf{C} = \mathbf{L}(:, J)$ and $\mathbf{R} = \mathbf{L}(I, :)$ of the underlying low-rank matrix $\mathbf{L}$. To apply the RPCA algorithms on $\tilde{\mathbf{C}}$ and $\tilde{\mathbf{R}}$, we must ensure that $\mathbf{C}$ and $\mathbf{R}$ are incoherent and that $\mathbf{S}(I, :)$, $\mathbf{S}(:, J)$ are sparse. In the next subsection, we provide quite general bounds for the incoherence of submatrices for various column selection methods.

Second, to carry out an error analysis for RCUR decompositions, we must consider a concrete sampling method for which the submatrices $\mathbf{C}$ and $\mathbf{R}$ of $\mathbf{L}$ have low incoherence and the column and row submatrices of $\mathbf{S}$ are sparse, and we must consider the error of the CUR decomposition $\mathbf{L} - \hat{\mathbf{C}}(\hat{\mathbf{C}}(I, :))^{\dagger} \hat{\mathbf{R}}$ where $\hat{\mathbf{C}}$ is the output of a given RPCA algorithm whose input is $\tilde{\mathbf{C}}$. The main results pertaining to this step are in section 3.2.

3.1. How incoherence passes to submatrices. We begin with our first main result that gives generic bounds on incoherence and condition numbers of column submatrices of low-rank matrices. Since this matter is of general interest, in this subsection, we will suppose that we have access to a low-rank matrix $\mathbf{L}$ itself, which is not true in the RPCA problem. We stress that this assumption is not at all necessary, but we will show that one can obtain concrete bounds for incoherence of submatrices compared to incoherence of the full matrix for a variety of column sampling schemes.

Theorem 3.1. Suppose $\mathbf{L} \in \mathbb{R}^{m \times n}$ has rank r and satisfies (A1). Let $J \subseteq [n]$ such that $\mathbf{C} = \mathbf{L}(:, J)$ has rank r , and denote $\beta := \sqrt{\frac{|J|}{n}} \|\mathbf{V}_{\mathbf{L}}(J, :)^{\dagger}\|_2$. Then

1. $\max_i \|\mathbf{W}_{\mathbf{C}}^T \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu_1(\mathbf{L})r}{m}},$
2. $\max_i \|\mathbf{V}_{\mathbf{C}}^T \mathbf{e}_i\|_2 \leq \beta \kappa(\mathbf{L}) \sqrt{\frac{\mu_2(\mathbf{L})r}{|J|}},$
3. $\kappa(\mathbf{C}) \leq \beta \sqrt{\mu_2(\mathbf{L})r} \kappa(\mathbf{L}).$

In particular,

$$\mu_1(\mathbf{C}) \leq \mu_1(\mathbf{L}) \quad \text{and} \quad \mu_2(\mathbf{C}) \leq \beta^2 \kappa(\mathbf{L})^2 \mu_2(\mathbf{L}).$$

Remark 3.2. We assume the condition number $\kappa(\mathbf{L}) = \mathcal{O}(1)$, which is commonly assumed in many RPCA papers [11, 12, 53]. According to Theorem 3.1, a well-conditioned $\mathbf{L}$ with good incoherence ensures the corresponding properties of its submatrices, and this is a key to the success of RCUR.

Naturally, the value of β will depend heavily on the sampling method utilized to select the column indices J . Thus, we desire algorithms for selecting column submatrices of $\mathbf{L}$ for which $\beta \ll \mathcal{O}(n)$. However, it is typically the case that $|J|$ and β are inversely proportional and also that judicious sampling of smaller column indices can be significantly more costly than, say, uniform random sampling. Consequently, we will seek a trade-off for which both $|J|$ and β are small compared to n , while simultaneously maintaining low complexity in the sampling scheme that produces the column indices J .

Let us now better illustrate the character of the bounds given in Theorem 3.1 by some specific examples. We first consider deterministic sampling methods and begin with an existential bound based on maximum volume sampling. The *volume* of a matrix $M \in \mathbb{R}^{s \times t}$ is $\prod_{i=1}^{\min\{s,t\}} \sigma_i(M)$.

Proposition 3.3 ([42, Lemma 1]). *Invoke the notations and assumptions of Theorem 3.1, and assume that $|J| = \ell$. If J is chosen so that $\mathbf{V}_{\mathbf{L}}(J, :)$ is the maximal volume submatrix of $\mathbf{V}_{\mathbf{L}}$ of size $\ell \times r$, then*

$$\beta \leq \sqrt{\frac{\ell}{n} + \frac{r\ell(n-\ell)}{n(\ell-r+1)}}.$$

Note that if $\ell = r$, then $\beta = \mathcal{O}(r)$, whereas if $\ell = n$, then $\beta = 1$. Moreover, the function on the right-hand side of the estimate in Proposition 3.3 is decreasing as ℓ increases for fixed r and n . Thus maximal volume sampling achieves a good bound for β . However, these bounds come at substantial cost: finding maximal volume submatrices is NP-hard [4]. There are approximation algorithms available [25, 38, 41], but analysis of how close their output is to the optimal solution is unknown.

Regarding deterministic column sampling methods, Avron and Boutsidis prove upper bounds on β for a greedy column selection algorithm [3]. We will discuss this algorithm in more detail later, but for now, we mention the following.

Theorem 3.4 ([3, Theorem 3.1]). *Let $\mathbf{L} \in \mathbb{R}^{m \times n}$ with $\text{rank}(\mathbf{L}) = r$. Suppose $\mathbf{L}$ satisfies (A1). Then Algorithm C.1 ([3, Algorithm 1]) applied to $\mathbf{V}_{\mathbf{L}}$ yields column indices $J \subseteq [n]$ with $|J| = r$ for which $\mathbf{C} = \mathbf{L}(:, J)$ satisfies $\text{rank}(\mathbf{C}) = \text{rank}(\mathbf{L})$, and $\beta \leq r$.*

The greedy algorithm of [3] applied to $\mathbf{V}_{\mathbf{L}}$ requires $\mathcal{O}(nr^2 + nr(n-r))$ operations to select exactly r columns and yield the bound $\beta \leq r$. Thus, this method is more feasible than maximal volume sampling. Both of the above procedures deterministically select exactly r columns of $\mathbf{L}$ and yield a good bound for β , but at nontrivial computational cost. We now turn our gaze to random sampling methods, which we will see can yield good bounds on β with high probability at trivial computational cost.

By far the easiest and cheapest sampling method is to choose J from $[n]$ via uniform sampling. While uniform sampling may not be successful for arbitrary matrices, it is known to be typically successful for incoherent matrices. The following theorem, derived from estimates

of pseudoinverses of orthogonal matrices due to Tropp [48], provides a bound for β under this paradigm.

Theorem 3.5. *Suppose that $\mathbf{L} \in \mathbb{R}^{m \times n}$ has rank r and satisfies condition (A1), and $\mathbf{V}_{\mathbf{L}} \in \mathbb{R}^{n \times r}$ are its first r right singular vectors. Suppose that $J \subseteq [n]$ is chosen by sampling uniformly without replacement to yield $\mathbf{C} = \mathbf{L}(:, J)$ and that $|J| \geq \gamma \mu_2(\mathbf{L})r$ for some $\gamma > 0$. Then the quantity $\beta = \sqrt{|J|/n} \|(\mathbf{V}_{\mathbf{L}}(J, :))^{\dagger}\|_2$ satisfies*

1. $\beta \leq \frac{1}{\sqrt{1-\delta}}$,
2. $\max_i \|\mathbf{V}_{\mathbf{C}}^T \mathbf{e}_i\|_2 \leq \beta \kappa(\mathbf{L}) \sqrt{\frac{\mu_2(\mathbf{L})r}{|J|}}$, and
3. $\kappa(\mathbf{C}) \leq \beta \sqrt{\mu_2(\mathbf{L})r} \kappa(\mathbf{L})$

with probability of at least $1 - \frac{r}{e^{\gamma(\delta+(1-\delta)\log(1-\delta))}}$. In particular, $\mu_1(\mathbf{C}) \leq \mu_1(\mathbf{L})$ and $\mu_2(\mathbf{C}) \leq \beta^2 \kappa(\mathbf{L})^2 \mu_2(\mathbf{L})$ with the given success probability.

Let us provide a concrete choice of parameters for illustration, though we stress that the estimate in Theorem 3.5 is quite flexible.

Corollary 3.6. *With the notations and assumptions of Theorem 3.5, put $\delta = 0.99$. Then sampling $|J| \geq 1.06 \mu_2(\mathbf{L})r \log(rn)$ columns from $[n]$ uniformly without replacement yields*

1. $\beta \leq 10$,
2. $\max_i \|\mathbf{V}_{\mathbf{C}}^T \mathbf{e}_i\|_2 \leq 10 \kappa(\mathbf{L}) \sqrt{\frac{\mu_2(\mathbf{L})r}{|J|}} \leq 10 \kappa(\mathbf{L}) (\log(rn))^{-\frac{1}{2}}$, and
3. $\kappa(\mathbf{C}) \leq 10 \sqrt{\mu_2(\mathbf{L})r} \kappa(\mathbf{L})$

with probability at least $1 - \frac{1}{n}$. In particular, $\mu_1(\mathbf{C}) \leq \mu_1(\mathbf{L})$ and $\mu_2(\mathbf{C}) \leq 100 \kappa(\mathbf{L})^2 \mu_2(\mathbf{L})$ with the given success probability.

A brief remark here is needed on sampling order in Corollary 3.6—the $r \log(rn)$ term is mild but necessary to achieve success probability which decays with the problem size, n ; one can derive a similar result in which the success probability decays polynomially with the rank, r , while requiring only $|J| \gtrsim r \log(r)$; such a bound could be useful in problems in which the rank is not minuscule. Later, this sampling order may be compared with those of CUR approximations and will be seen to be favorable.

For ease of exposition above, we have focused solely on the case of column selection; however, we note that analogous results are easily obtained for row selection by taking transpose in the formulae above. For clarity, since we will subsequently be interested in selecting both column and row submatrices to form CUR decompositions, we state the combined bound on incoherence here.

Corollary 3.7. *With the notations and assumptions of Corollary 3.6, suppose also that $I \subseteq [m]$ with $|I| \geq 1.06 \mu_1(\mathbf{L})r \log(rm)$ is sampled uniformly without replacement, and that $\mathbf{R} = \mathbf{L}(I, :)$. If $\beta' := \sqrt{|I|/m} \|(\mathbf{W}_{\mathbf{L}}(I, :))^{\dagger}\|_2$, then with probability at least $(1 - \frac{1}{n})(1 - \frac{1}{m})$, the conclusion of Corollary 3.6 holds, and so do the following:*

1. $\beta' \leq 10$,
2. $\max_i \|\mathbf{V}_{\mathbf{R}}^T \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu_2(\mathbf{L})r}{n}}$,
3. $\max_i \|\mathbf{W}_{\mathbf{R}}^T \mathbf{e}_i\|_2 \leq 10 \kappa(\mathbf{L}) \sqrt{\frac{\mu_1(\mathbf{L})r}{|I|}}$,
4. $\kappa(\mathbf{R}) \leq 10 \sqrt{\mu_1(\mathbf{L})r} \kappa(\mathbf{L})$.

In particular, $\mu_1(\mathbf{R}) \leq 100 \kappa(\mathbf{L})^2 \mu_1(\mathbf{L})$, and $\mu_2(\mathbf{R}) \leq \mu_2(\mathbf{L})$.

Remark 3.8. Note that if one used more sophisticated random sampling techniques, such as leverage score sampling [37] or volume sampling [18], one could potentially obtain good bounds on β with larger success probability than in the results stated here. After all, these distributions take into account more information about the structure of the column and row space of the matrix. However, using these comes at a nontrivial computational cost for computation of the probability distributions ($\mathcal{O}(mnk)$ for leverage scores and approximately $\mathcal{O}(\max\{m^2, n^2\})$ for volume sampling). Given that a constant bound for β can be obtained with relatively large probability via uniform sampling, which requires no computation, this method should be preferred for incoherent matrices.

There are faster methods than those mentioned above that approximate leverage score distributions, e.g., [21], but to be effective in the RPCA model, one needs to approximate leverage scores of $\mathbf{L}$ from observations of $\mathbf{L} + \mathbf{S}$. Leverage scores are unstable under addition of sparse, arbitrary magnitude outliers, so without imposing additional restrictions on $\mathbf{S}$, such approximations are likely unsuitable for the RPCA problem.

3.2. RCUR decomposition. Now we are ready to state our main algorithm. The idea of Algorithm 3.1 is simple yet effective due to the computational ease of uniform sampling and the known success of CUR decompositions. The algorithm may be considered a proto-type in that it has an RPCA algorithm as a subroutine in lines 5 and 6. We emphasize that any RPCA algorithm can be used in these lines, and we will quantify the accuracy and computational complexity of our algorithm based on those of the RPCA subroutine.

Remark 3.9. We note that Algorithm 3.1 and all of the theoretical results related to it here hold for both sampling column and row indices uniformly *with* or *without* replacement. For simplicity, we state all results here for sampling *with* replacement. The only minor change to any results here for sampling *without* replacement is that some of the constants in the big- $\mathcal{O}$ notations and success probabilities will change, but typically the change is small.

Now let us state a sample result illustrating the type of guarantee Algorithm 3.1 gives. To do so, we assume that the AltProj RPCA algorithm of [39] is used in lines 5 and 6.

Theorem 3.10. *Let $\mathbf{L} \in \mathbb{R}^{m \times n}$ and $\mathbf{S} \in \mathbb{R}^{m \times n}$ satisfy assumptions (A1) and (A2), respectively, with $\text{rank}(\mathbf{L}) = r$. Let $\varepsilon, \delta \in (0, 1)$ and $\alpha = \mathcal{O}(\frac{1-\delta}{\kappa(\mathbf{L})^2(\mu_1(\mathbf{L}) \vee \mu_2(\mathbf{L}))})$. Suppose that*

$$|I| \geq c_1 \max\{\mu_1(\mathbf{L})r \log(m), \log(m)/\alpha\} \quad \text{and} \quad |J| \geq c_2 \max\{\mu_2(\mathbf{L})r \log(n), \log(n)/\alpha\}$$

Algorithm 3.1. Uniform sampling RCUR

- 1: **Input:** $\mathbf{D} = \mathbf{L} + \mathbf{S} \in \mathbb{R}^{m \times n}$ with $\text{rank}(\mathbf{L}) = r$; RPCA: the chosen RPCA solver.
 - 2: **Output:** $\hat{\mathbf{L}}$: approximation of $\mathbf{L}$.
 - 3: **Initialization:** Draw sampling indices $I \subseteq [m]$, $J \subseteq [n]$ uniformly.
 - 4: $\tilde{\mathbf{C}} = \mathbf{D}(:, J)$, $\tilde{\mathbf{R}} = \mathbf{D}(I, :)$
 - 5: $[\hat{\mathbf{C}}, \hat{\mathbf{S}}_C] = \text{RPCA}(\tilde{\mathbf{C}}, r)$
 - 6: $[\hat{\mathbf{R}}, \hat{\mathbf{S}}_R] = \text{RPCA}(\tilde{\mathbf{R}}, r)$
 - 7: $\hat{\mathbf{L}} = \hat{\mathbf{C}}(\hat{\mathbf{C}}(I, :))^{\dagger} \hat{\mathbf{R}}$
-

rows and columns, respectively, are chosen uniformly without replacement and that AltProj [39] is used as the RPCA subroutine in lines 5 and 6 of Algorithm 3.1. Then by running $\mathcal{O}(r \log(\sqrt{\frac{mn}{|I||J|}} \frac{\kappa(\mathbf{L})}{(1-\delta)\varepsilon}))$ iterations of AltProj, the output of Algorithm 3.1 satisfies

$$\frac{\|\mathbf{L} - \hat{\mathbf{L}}\|_2}{\|\mathbf{L}\|_2} \leq \varepsilon \kappa(\mathbf{L})^{-1}$$

with probability at least $1 - \frac{r}{n^{c_2(\delta+(1-\delta)\log(1-\delta))}} - \frac{r}{m^{c_1(\delta-(1-\delta)\log(1-\delta))}} - \frac{1}{n} - \frac{1}{m}$.

Moreover, the total complexity of Algorithm 3.1 with AltProj as a subroutine is

$$\mathcal{O}\left((mr^3 \log(n) + nr^3 \log(m)) \log\left(\sqrt{\frac{mn}{|I||J|}} \frac{\kappa(\mathbf{L})}{(1-\delta)\varepsilon}\right)\right).$$

3.3. Hybrid algorithms for RCUR approximation. The output of Algorithm 3.1 is an approximation of the form $\mathbf{L} \approx \hat{\mathbf{C}}\hat{\mathbf{U}}^\dagger\hat{\mathbf{R}}$, where $\hat{\mathbf{C}}$ and $\hat{\mathbf{R}}$ are cleaned versions of columns and rows, respectively, of $\mathbf{D}$. Put a different way, $\hat{\mathbf{C}}$ approximates the column space of $\mathbf{L}$, and $\hat{\mathbf{R}}$ approximates the row space of $\mathbf{L}$. However, these consist of $\mathcal{O}(r \log(rn))$ and $\mathcal{O}(r \log(rm))$ vectors, respectively, while the dimensions of the column and row space of $\mathbf{L}$ are both r ; that is, we have redundant approximate bases for these spaces. We are interested in determining from these the most relevant vectors that represent the largest subspaces of both the column and row space of $\mathbf{L}$. In particular, from the randomly sampled indices I, J , we would prefer to choose *exactly* r elements of each. This problem in general is called the *column subset selection problem*: given $\hat{\mathbf{C}} \in \mathbb{R}^{m \times \ell}$, choose $I_1 \subseteq [\ell]$ (and set $\hat{\mathbf{C}}_1$) with $|I_1| = r$ which minimizes $\|\hat{\mathbf{C}} - \hat{\mathbf{C}}_1\hat{\mathbf{C}}_1^\dagger\hat{\mathbf{C}}\|_2$ over all possible column submatrices of $\hat{\mathbf{C}}_1$ consisting of exactly r columns. This problem is NP-complete [43] (see also [17] for other complexity results), but there are both randomized and deterministic approximation algorithms; a small subsampling of prior art on this problem is [3, 5, 6, 19, 47].

We consider now what happens when we apply a deterministic column subset selection algorithm to the matrices $\hat{\mathbf{C}}$ and $\hat{\mathbf{R}}$ given by the output of lines 5 and 6 of Algorithm 3.1 to choose exactly r columns and rows thereof. We state this algorithm generally but specialize our analysis to the greedy algorithm of [3] (described fully here in the appendix as Algorithm C.1) to state a concrete theorem. For clarity, we state the general algorithm here as Algorithm 3.2. Again, RPCA is any such algorithm the user specifies, and DeterministicCS($\mathbf{A}, r$) is any algorithm that deterministically selects exactly r columns of $\mathbf{A}$ (I_1 and J_1 in Algorithm 3.2 are the indices chosen by the deterministic column selection algorithm).

The following theorem describes the accuracy of the output of Algorithm 3.2.

Theorem 3.11. *Take the notations and assumptions of Theorem 3.10. Let $\hat{\mathbf{C}}$ and $\hat{\mathbf{R}}$ be the intermediate outputs of Algorithm 3.2 in lines 5 and 6 with AltProj being used as the RPCA subroutine, where the latter runs $\mathcal{O}(r \log(\frac{229\kappa(\mathbf{L})^5 r \sqrt{mn\mu_1(\mathbf{L})\mu_2(\mathbf{L})}}{\varepsilon(1-\delta)^2 \sigma_{\max}(\mathbf{L})}))$ iterations. Let $\hat{\mathbf{C}}_1 = \hat{\mathbf{C}}(:, J_1)$, $\hat{\mathbf{U}}_1 = \hat{\mathbf{C}}(I_1, :)$, and $\hat{\mathbf{R}}_1 = \hat{\mathbf{R}}(I_1, :)$, where I_1 and J_1 are obtained by applying Algorithm C.1 on $\hat{\mathbf{R}}$ and $\hat{\mathbf{C}}$ with $|I_1| = |J_1| = r$ in lines 8 and 9. Then with probability at least $1 - \frac{r}{n^{c_2(\delta+(1-\delta)\log(1-\delta))}} - \frac{r}{m^{c_1(\delta-(1-\delta)\log(1-\delta))}} - \frac{1}{n} - \frac{1}{m}$,*

$$\|\mathbf{L} - \hat{\mathbf{C}}_1\hat{\mathbf{U}}_1^\dagger\hat{\mathbf{R}}_1\|_2 \leq \frac{1}{3}\varepsilon\sigma_{\min}(\mathbf{L}).$$

Algorithm 3.2. Uniform + deterministic RCUR

-
- 1: **Input:** $\mathbf{D} = \mathbf{L} + \mathbf{S} \in \mathbb{R}^{m \times n}$ with $\text{rank}(\mathbf{L}) = r$; RPCA: the chosen RPCA solver; DeterministicCS: the chosen deterministic row/column selection method.
 - 2: **Output:** $\hat{\mathbf{L}}$: approximation of $\mathbf{L}$.
 - 3: **Initialization:** Draw sampling indices $I \subseteq [m]$, $J \subseteq [n]$ uniformly.
 - 4: $\tilde{\mathbf{C}} = \mathbf{D}(:, J)$, $\tilde{\mathbf{R}} = \mathbf{D}(I, :)$
 - 5: $[\hat{\mathbf{C}}, \hat{\mathbf{S}}_C] = \text{RPCA}(\tilde{\mathbf{C}}, r)$
 - 6: $[\hat{\mathbf{R}}, \hat{\mathbf{S}}_R] = \text{RPCA}(\tilde{\mathbf{R}}, r)$
 - 7: $[\hat{\mathbf{C}}_1, \hat{\mathbf{S}}_{C_1}, J_1] = \text{DeterministicCS}(\hat{\mathbf{C}}, r)$
 - 8: $[\hat{\mathbf{R}}_1, \hat{\mathbf{S}}_{R_1}, I_1] = \text{DeterministicCS}(\hat{\mathbf{R}}^T, r)$
 - 9: $\hat{\mathbf{L}} = \hat{\mathbf{C}}_1(\hat{\mathbf{C}}_1(I_1, :))^{\dagger} \hat{\mathbf{R}}_1$
-

Remark 3.12. We note for the reader two things here. First, the proofs of the approximation theorems here are valid for the Frobenius norm as well as the smaller spectral norm; for presentation purposes we choose the latter. Second, the bounds in Theorems 3.10 and 3.11 are the same by dividing by $\|\mathbf{L}\|_2$; e.g., the bound in Theorem 3.11 is the same as the relative error bound

$$\frac{\|\mathbf{L} - \hat{\mathbf{C}}_1 \hat{\mathbf{U}}_1^{\dagger} \hat{\mathbf{R}}\|_2}{\|\mathbf{L}\|_2} \leq \frac{1}{3} \varepsilon \kappa(\mathbf{L})^{-1}.$$

3.4. Deterministic column selection. In Appendices D and E, we study deterministic column sampling methods, in particular via the greedy algorithm of Avron and Boutsidis [3]. We are able to prove that, with proper preconditioning, sampling columns and rows via the greedy algorithm on the singular vectors of $\mathbf{D}$ produce submatrices of $\mathbf{L}$ with similar incoherence (see Remark D.2). In particular, we find that the method described in Appendix D produces deterministically a submatrix $\mathbf{C}$ for which $\beta \leq 4\kappa(\mathbf{L})\sqrt{r}$. However, this procedure does not guarantee $\mathcal{O}(\alpha)$ -sparsity of the corresponding submatrices of $\mathbf{S}$. This we leave as an open problem for future work.

4. Passing properties to submatrices. In this section, we study how three properties of a matrix pass to column submatrices thereof: incoherence, condition number, and sparsity. These results are of independent interest but are also necessary to study the use of RPCA algorithms on submatrices.

4.1. Incoherence. In this subsection, we study the relation of the incoherence of the matrix $\mathbf{L} \in \mathbb{R}^{m \times n}$ to that of a column submatrix of it, say $\mathbf{C} = \mathbf{L}(:, J) \in \mathbb{R}^{m \times \ell}$. We find that column incoherence is unchanged, while row incoherence inflates by no more than a factor of $\beta\kappa(\mathbf{L})$.

Let us begin by the observation that if $\mathbf{L}$ has column incoherence as in the first inequality of (A1), then any column submatrix $\mathbf{C}$ has the same left incoherence.

Lemma 4.1. Suppose $\mathbf{L} \in \mathbb{R}^{m \times n}$ satisfies (A1). For any $J \subseteq [n]$ such that $\text{rank}(\mathbf{C}) = \text{rank}(\mathbf{L})$ with $\mathbf{C} = \mathbf{L}(:, J)$, then $\mathbf{C}$ satisfies

$$\max_i \|\mathbf{W}_{\mathbf{C}}^T \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu_1(\mathbf{L})r}{m}}.$$

Proof. Notice that $\mathbf{C} = \mathbf{W}_L \boldsymbol{\Sigma}_L (\mathbf{V}_L(J, :))^T = \mathbf{W}_L \widetilde{\mathbf{W}}_C \boldsymbol{\Sigma}_C \mathbf{V}_C^T$, where $\boldsymbol{\Sigma}_L (\mathbf{V}_L(J, :))^T$ has the compact SVD decomposition $\widetilde{\mathbf{W}}_C \boldsymbol{\Sigma}_C \mathbf{V}_C^T$ with $\widetilde{\mathbf{W}}_C \in \mathbb{R}^{r \times r}$ being an orthonormal matrix and $\mathbf{V}_C^T \in \mathbb{R}^{r \times |J|}$. Therefore, $\mathbf{W}_C = \mathbf{W}_L \widetilde{\mathbf{W}}_C$. Thus, we have

$$\max_i \|\mathbf{W}_C^T \mathbf{e}_i\|_2 = \max_i \|\widetilde{\mathbf{W}}_C^T \mathbf{W}_L^T \mathbf{e}_i\|_2 = \max_i \|\mathbf{W}_L^T \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu_1(\mathbf{L})r}{m}}. \quad \blacksquare$$

Lemma 4.1 implies that we need not be concerned with the column incoherence of any column submatrix $\mathbf{C}$ of $\mathbf{L}$ given that it always matches that of the matrix it is obtained from. Now we need to understand the incoherence for $\mathbf{V}_C$, which requires some more work. First, let us make two preliminary observations.

Lemma 4.2. *Let $\mathbf{L} \in \mathbb{R}^{m \times n}$ with rank r satisfy (A1). Then for any $J \subseteq [n]$, $\mathbf{C} = \mathbf{L}(:, J)$ satisfies*

$$\max_i \|\mathbf{V}_C^T \mathbf{e}_i\|_2 \leq \kappa(\mathbf{L}) \frac{\|\mathbf{C}^\dagger\|_2}{\|\mathbf{L}^\dagger\|_2} \sqrt{\frac{\mu_2(\mathbf{L})r}{n}}.$$

Proof. Notice that $\mathbf{V}_C^T = \boldsymbol{\Sigma}_C^{-1} \mathbf{W}_C^T \mathbf{C}$. Thus for any i ,

$$\begin{aligned} \|\mathbf{V}_C^T \mathbf{e}_i\|_2 &= \|\boldsymbol{\Sigma}_C^{-1} \mathbf{W}_C^T \mathbf{C} \mathbf{e}_i\|_2 \\ &\leq \|\boldsymbol{\Sigma}_C^{-1}\|_2 \|\mathbf{W}_C^T\|_2 \|\mathbf{C} \mathbf{e}_i\|_2 \\ &\leq \frac{1}{\sigma_{\min}(\mathbf{C})} \|\mathbf{C} \mathbf{e}_i\|_2. \end{aligned}$$

Since $\mathbf{C} = \mathbf{L}(:, J) = \mathbf{W}_L \boldsymbol{\Sigma}_L (\mathbf{V}_L(J, :))^T$,

$$\begin{aligned} \frac{1}{\sigma_{\min}(\mathbf{C})} \|\mathbf{C} \mathbf{e}_i\|_2 &= \frac{1}{\sigma_{\min}(\mathbf{C})} \|\mathbf{W}_L \boldsymbol{\Sigma}_L (\mathbf{V}_L(J, :))^T \mathbf{e}_i\|_2 \\ &\leq \frac{\sigma_1(\mathbf{L})}{\sigma_{\min}(\mathbf{C})} \|(\mathbf{V}_L(J, :))^T \mathbf{e}_i\|_2 \\ &\leq \frac{\sigma_1(\mathbf{L})}{\sigma_{\min}(\mathbf{C})} \|\mathbf{V}_L^T \mathbf{e}_i\|_2 \\ &\leq \kappa(\mathbf{L}) \frac{\sigma_{\min}(\mathbf{L})}{\sigma_{\min}(\mathbf{C})} \sqrt{\frac{\mu_2(\mathbf{L})r}{n}} \\ &= \kappa(\mathbf{L}) \frac{\|\mathbf{C}^\dagger\|_2}{\|\mathbf{L}^\dagger\|_2} \sqrt{\frac{\mu_2(\mathbf{L})r}{n}}. \quad \blacksquare \end{aligned}$$

Estimation of the quantity $\frac{\|\mathbf{C}^\dagger\|_2}{\|\mathbf{L}^\dagger\|_2}$ will generally depend heavily on how the column submatrix $\mathbf{C}$ is selected. One way to estimate this is to consider the norm of the pseudoinverse of the corresponding column submatrix of $\mathbf{V}_L$ as follows.

Lemma 4.3. *Let $\mathbf{L} \in \mathbb{R}^{m \times n}$ and $\mathbf{C} = \mathbf{L}(:, J)$ for some $J \subseteq [n]$. Then*

$$\frac{\|\mathbf{C}^\dagger\|_2}{\|\mathbf{L}^\dagger\|_2} \leq \|(\mathbf{V}_L(J, :))^\dagger\|_2.$$

Proof. Note that $\mathbf{C} = \mathbf{W}_L \boldsymbol{\Sigma}_L (\mathbf{V}_L(J, :))^T$. Then one has

$$\mathbf{C}^\dagger = (\boldsymbol{\Sigma}_L \mathbf{V}_L(J, :)^T)^\dagger \mathbf{W}_L^T = (\mathbf{V}_L(J, :)^T)^\dagger \boldsymbol{\Sigma}_L^{-1} \mathbf{W}_L^T,$$

where the first equality comes from the fact that $\mathbf{W}_L$ has orthonormal columns and $\mathbf{W}_L^\dagger = \mathbf{W}_L^T$, while the second equality follows from the fact that $\boldsymbol{\Sigma}_L$ has full column rank and $\mathbf{V}_L(J, :)^T$ has full row rank. With this in hand, we have

$$\|\mathbf{C}^\dagger\|_2 = \|((\mathbf{V}_L(J, :)^T)^\dagger \boldsymbol{\Sigma}_L^{-1} \mathbf{W}_L^T)\|_2 \leq \|(\mathbf{V}_L(J, :)^T)^\dagger\|_2 \|\boldsymbol{\Sigma}_L^{-1}\|_2 = \|\mathbf{L}^\dagger\|_2 \|(\mathbf{V}_L(J, :))^\dagger\|_2,$$

whence the result. ■

Combining the conclusions of Lemmata 4.1–4.3, we see that

$$(4.1) \quad \max_i \|\mathbf{V}_C^T \mathbf{e}_i\|_2 \leq \kappa(\mathbf{L}) \sqrt{\frac{|J|}{n}} \sqrt{\frac{\mu_2(\mathbf{L})r}{|J|}} \|(\mathbf{V}_L(J, :))^\dagger\|_2 = \beta \kappa(\mathbf{L}) \sqrt{\frac{\mu_2(\mathbf{L})r}{|J|}}.$$

Our primary concern now is to attempt to bound the quantity $\beta = \sqrt{\frac{|J|}{n}} \|(\mathbf{V}_L(J, :))^\dagger\|_2$. Unfortunately, there are relatively few general results for estimating pseudoinverses of submatrices of orthogonal matrices except in special instances. We make this into the parameter β of Theorem 3.1.

4.2. Condition number. Many RPCA algorithms require an estimate on the condition number of the matrix that is input into the algorithm to provide good theoretical guarantees. Here, we estimate $\kappa(\mathbf{C})$ in terms of $\kappa(\mathbf{L})$ in the case that $\mathbf{L}$ has a certain incoherence level.

Lemma 4.4. *Let $\mathbf{L} \in \mathbb{R}^{m \times n}$ have rank r and satisfy assumption (A1), and suppose $J \subseteq [n]$. Then*

1. $\|\mathbf{C}\|_2 \leq \sqrt{\frac{\mu_2(\mathbf{L})r|J|}{n}} \|\mathbf{L}\|_2$;
2. $\kappa(\mathbf{C}) \leq \sqrt{\mu_2(\mathbf{L})r} \kappa(\mathbf{L}) \sqrt{\frac{|J|}{n}} \|\mathbf{V}_L(J, :)\|_2$.

Proof. Since $\mathbf{L} = \mathbf{W}_L \boldsymbol{\Sigma}_L \mathbf{V}_L^T$ and $\mathbf{C} = \mathbf{L}(:, J)$, we have $\mathbf{C} = \mathbf{W}_L \boldsymbol{\Sigma}_L (\mathbf{V}_L(J, :))^T$. For all $\mathbf{x} \in \mathbb{R}^{|J|}$ with $\|\mathbf{x}\|_2 = 1$,

$$\begin{aligned} \|\mathbf{C}\mathbf{x}\|_2 &= \|\mathbf{W}_L \boldsymbol{\Sigma}_L (\mathbf{V}_L(J, :))^T \mathbf{x}\|_2 \\ &\leq \|\boldsymbol{\Sigma}_L\|_2 \|(\mathbf{V}_L(J, :))^T \mathbf{x}\|_2 \\ &= \|\mathbf{L}\|_2 \left\| (\mathbf{V}_L(J, :))^T \sum_{i=1}^{|J|} \mathbf{x}_i \mathbf{e}_i \right\|_2 \\ &\leq \|\mathbf{L}\|_2 \sum_{i \in J} |\mathbf{x}_i| \|(\mathbf{V}_L(J, :))^T \mathbf{e}_i\|_2 \\ &\leq \|\mathbf{L}\|_2 \sum_{i \in J} |\mathbf{x}_i| \sqrt{\frac{\mu_2(\mathbf{L})r}{n}} \\ &\leq \|\mathbf{L}\|_2 \sqrt{|J|} \|\mathbf{x}\|_2 \sqrt{\frac{\mu_2(\mathbf{L})r}{n}} \\ &= \sqrt{\frac{\mu_2(\mathbf{L})r|J|}{n}} \|\mathbf{L}\|_2. \end{aligned}$$

To prove the second inequality, note that

$$\kappa(\mathbf{C}) = \|\mathbf{C}\|_2 \|\mathbf{C}^\dagger\|_2 \leq \sqrt{\frac{\mu_2(\mathbf{L})r|J|}{n}} \|\mathbf{L}\|_2 \|\mathbf{V}_L(J, :)^{\dagger}\|_2 \|\mathbf{L}^{\dagger}\|_2,$$

which gives the required result upon rearranging terms. $\blacksquare$

The results above combine to prove the main theorem (Theorem 3.1) regarding the quantity β .

Proof of Theorem 3.1. The proof follows by combining Lemmata 4.1–4.4 and (4.1). $\blacksquare$

4.3. Sparsity. Our ability to apply a RPCA algorithm to a column submatrix of $\mathbf{D}$ requires that the corresponding submatrix of $\mathbf{S}$ remains a sparse matrix. The content of the following proposition is that uniform sampling of sufficiently many ($\mathcal{O}(r \log n)$) columns yields a 2α -sparse submatrix if $\mathbf{S}$ was α -sparse.

Proposition 4.5. *Let $\mathbf{S}$ be an α -sparse matrix. We consider two cases of uniform sampling:*

1. *(With replacement) Consider the submatrix of $\mathbf{S} \in \mathbb{R}^{m \times n}$ formed by $|I|$ rows that were uniformly sampled with replacement, namely $\mathbf{R}$. Assume $|I| = cr \log(n)$ with $c \geq \frac{16}{3\alpha r}$; then $\mathbf{R} \in \mathbb{R}^{|I| \times n}$ is 2α -sparse with probability at least $1 - n^{-1}$. Similarly, $\mathbf{C} \in \mathbb{R}^{m \times |J|}$, formed by $|J| = cr \log(m)$ columns of $\mathbf{S}$ that were uniformly sampled with replacement, is also 2α -sparse with probability at least $1 - m^{-1}$.*
2. *(Without replacement) Consider the submatrix of $\mathbf{S} \in \mathbb{R}^{m \times n}$ formed by $|I|$ rows that were uniformly sampled without replacement, namely $\mathbf{R}$. Assume $|I| = cr \log(n)$ with $c \geq \frac{8}{\alpha r}$; then $\mathbf{R} \in \mathbb{R}^{|I| \times n}$ is 2α -sparse with probability at least $1 - 2n^{-1}$. Similarly, $\mathbf{C} \in \mathbb{R}^{m \times |J|}$, formed by $|J| = cr \log(m)$ columns of $\mathbf{S}$ that were uniformly sampled without replacement, is also 2α -sparse with probability at least $1 - 2m^{-1}$.*

Proof. Part 1. We first prove the version of uniform sampling *with* replacement. By the definition of α -sparsity, each row of $\mathbf{R}$ has no more than αn nonzero entries. Moreover, assume the j th column of $\mathbf{S}$ has exactly αn nonzero entries. Let X_i indicate whether $\mathbf{R}_{i,j}$ is nonzero. One can see that $X_i \sim \text{Ber}(\alpha)$. Consequently, by Bernstein's inequality, we have

$$\begin{aligned} \mathbb{P}\left(\sum_{i \in I} X_i > 2\alpha|I|\right) &= \mathbb{P}\left(\sum_{i \in I} X_i - \mathbb{E}(X_i) > \alpha|I|\right) \\ &\leq \exp\left(-\frac{\alpha^2|I|^2}{2|I|\alpha(1-\alpha) + \frac{2}{3}\alpha|I|}\right) \\ &\leq \exp\left(-\frac{\alpha|I|}{2(1-\alpha) + \frac{2}{3}}\right) \\ &\leq \exp(-2\log(n)) \\ &= n^{-2} \end{aligned}$$

for the j th column of $\mathbf{R}$, where the last inequality uses the assumption $\alpha|I| \geq \frac{16}{3} \log(n)$. Note that the probability inequality still holds if $\mathbf{S}$ has less than αn nonzero entries. Furthermore, to have all n columns being 2α -sparse, we get

$$\mathbb{P}(\mathbf{R} \text{ is } 2\alpha\text{-sparse}) \geq (1 - n^{-2})^n \geq 1 - n^{-1}.$$

Hence, $\mathbf{R}$ is 2α -sparse with high probability; the proof that $\mathbf{C}$ is 2α -sparse with high probability is the same and so is omitted.

Part 2. Now, we prove the version of uniform sampling *without* replacement. By [29, Theorem 1], the Bernstein inequality for uniform sampling *without* replacement is

$$\mathbb{P}\left(\sum_{i=0}^m Y_i > t\right) \leq 2 \exp\left(-\frac{t^2}{4m\sigma^2}\right),$$

where Y_i is zero-mean random variables with uniform probability and σ is the variance. Taking the same argument as in Part 1 except sampling uniformly *without* replacement, we have

$$\begin{aligned} \mathbb{P}\left(\sum_{i \in I} X_i > 2\alpha|I|\right) &= \mathbb{P}\left(\sum_{i \in I} X_i - \mathbb{E}(X_i) > \alpha|I|\right) \\ &\leq 2 \exp\left(-\frac{\alpha^2|I|^2}{4|I|\alpha(1-\alpha)}\right) \\ &\leq 2 \exp\left(-\frac{\alpha|I|}{4(1-\alpha)}\right) \\ &\leq 2 \exp(-2\log(n)) \\ &= 2n^{-2} \end{aligned}$$

for the j th column of $\mathbf{R}$, where the last inequality uses the assumption $\alpha|I| \geq 8\log(n)$. Therefore,

$$\mathbb{P}(\mathbf{R} \text{ is } 2\alpha\text{-sparse}) \geq (1 - 2n^{-2})^n \geq 1 - 2n^{-1}$$

for $\mathbf{R}$ generated by uniform sampling *without* replacement. A similar result holds for $\mathbf{C}$. ■

5. Simulations. In this section, the effectiveness and computational efficiency of RCUR (i.e., Algorithm 3.1) is illustrated on two key imaging applications, video background-foreground separation and face modeling, and can be summarized in the following three ways:

1. We show that running time of RCUR on a video background subtraction task is typically an order of magnitude faster than the state-of-the-art RPCA algorithm. Qualitatively, we also show that the output produced by RCUR is similar to that of RPCA benchmarks.
2. We illustrate the interpretable nature of the dictionary vectors produced by RCUR.
3. We show that the running time of RCUR on face modeling is typically faster than that of RPCA and the output produced by RCUR is comparable to that of RPCA.

Moreover, for fair comparison, we use AccAltProj² as the RPCA algorithm in all the numerical experiments, including the RPCA subroutine that is used in lines 5–6 of Algorithm 3.1 and 3.2.

²AccAltProj [12] is a fast state-of-the-art RPCA algorithm that accelerates the popular AltProj [39].

5.1. Application to video background-foreground separation. Background-foreground separation is an important tool to extract moving objects from a static background in a video [8]. In this section, we compare the performance of the RCUR decomposition with that of RPCA on three public benchmark videos: *Shoppingmall*, *Restaurant*, and *OSU* datasets. The size of each frame of *Shoppingmall* is 256×320 , that of *Restaurant* is 120×160 , and that of *OSU* is 240×320 . The total number of frames are 1000 for *Shoppingmall*, 3055 for *Restaurant*, and 1506 for *OSU*. Each video can be represented as a matrix $\mathbf{D} \in \mathbb{R}^{m \times n}$ by vectorizing and stacking the frames as columns of the matrix; thus the *Shoppingmall* video produces a matrix in $\mathbb{R}^{81,920 \times 1,000}$ while *Restaurant*'s is in $\mathbb{R}^{19,200 \times 3,055}$ and *OSU*'s is in $\mathbb{R}^{76,800 \times 1506}$. Then $\mathbf{D} = \mathbf{L} + \mathbf{S}$, where $\mathbf{L}$ is the background, and $\mathbf{S}$ contains the foreground elements treated as sparse outliers.

5.1.1. Runtime and quality. We apply Algorithm 3.1 to $\mathbf{D}$, the data matrix from *Shoppingmall*, *Restaurant*, and *OSU*. We take $r = 2$ to be the underlying rank of the background for each video sequence. Then according to Theorem 3.10, we randomly choose column $\tilde{\mathbf{C}} = \mathbf{D}(:, J)$ and row $\tilde{\mathbf{R}} = \mathbf{D}(I, :)$ submatrices of size $m \times 15r \log(n)$ and $25r \log(m) \times n$, respectively. Note that $\mu_1(\mathbf{L}) \neq \mu_2(\mathbf{L})$ in many real-world applications, and $\mu_1(\mathbf{L}) > \mu_2(\mathbf{L})$ under this problem setup. Then we apply the selected RPCA algorithm (i.e., AccAltProj [12]) on $\tilde{\mathbf{C}}$ and $\tilde{\mathbf{R}}$ (lines 5 and 6) and denote the corresponding output low-rank matrices by $\hat{\mathbf{C}}$ and $\hat{\mathbf{R}}$. We then approximate the static background of the video, which is a low-rank matrix $\mathbf{L}$, as $\mathbf{L} \approx \hat{\mathbf{L}} = \hat{\mathbf{C}}(\hat{\mathbf{C}}(I, :))^{\dagger} \hat{\mathbf{R}}$. The foreground $\mathbf{S}$, which represents the moving objects, is then approximated by $\mathbf{D} - \hat{\mathbf{L}}$.

Table 1 shows the runtime for Algorithm 3.1 as well as the runtime for applying RPCA directly on $\mathbf{D}$. In all cases, Algorithm 3.1 runs substantially faster; in particular, on most problems, the runtime is a full order of magnitude faster than that of RPCA.

Running time is not the whole story; as the quality of the separation of foreground and background is of prime importance as well. Note that in the background-foreground separation tasks within the RPCA framework, there is no ground truth, and thus comparisons must be done qualitatively. Figures 1, 2, and 3 illustrate the outputs of RCUR and RPCA on *Restaurant*, *Shoppingmall*, and *OSU* sequences, respectively. Each figure presents three randomly selected frames from the full video sequence as well as the background and foreground images from each algorithm. In both cases, we note that the reconstruction algorithms yield comparable results; thus, the significant speed advantage afforded by RCUR is desirable in this particular RPCA task.

Table 1
Video size and runtime.

	Frame size	Frame number	Runtime (sec)	
			RCUR	RPCA
Shoppingmall	256×320	1000	7.69	44.30
Restaurant	120×160	3055	3.48	31.63
OSU	240×320	1506	10.39	68.62

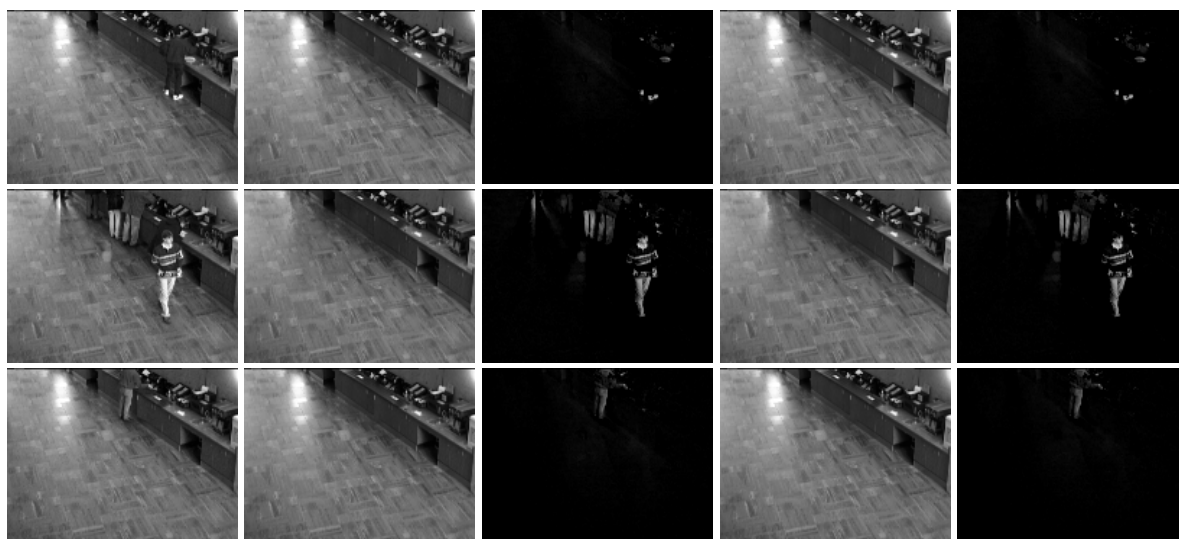


Figure 1. Restaurant: The first column contains three randomly selected frames from the original video. The middle two columns are the separated background and foreground outputs of RCUR, respectively. The right two columns are the separated background and foreground outputs of RPCA, respectively.



Figure 2. Shoppingmall: The first column contains three randomly selected frames from the original video. The middle two columns are the separated background and foreground outputs of RCUR, respectively. The right two columns are the separated background and foreground outputs of RPCA, respectively.

5.1.2. Interpretability: Extracting canonical frames. CUR decompositions have long been touted as providing better interpretability of data representations than abstract bases such as those obtained by QR or SVD (see, e.g., [33, 37, 44]), the idea being that a subset of the data itself is a good dictionary for the whole. This is also known as the self-expressivity

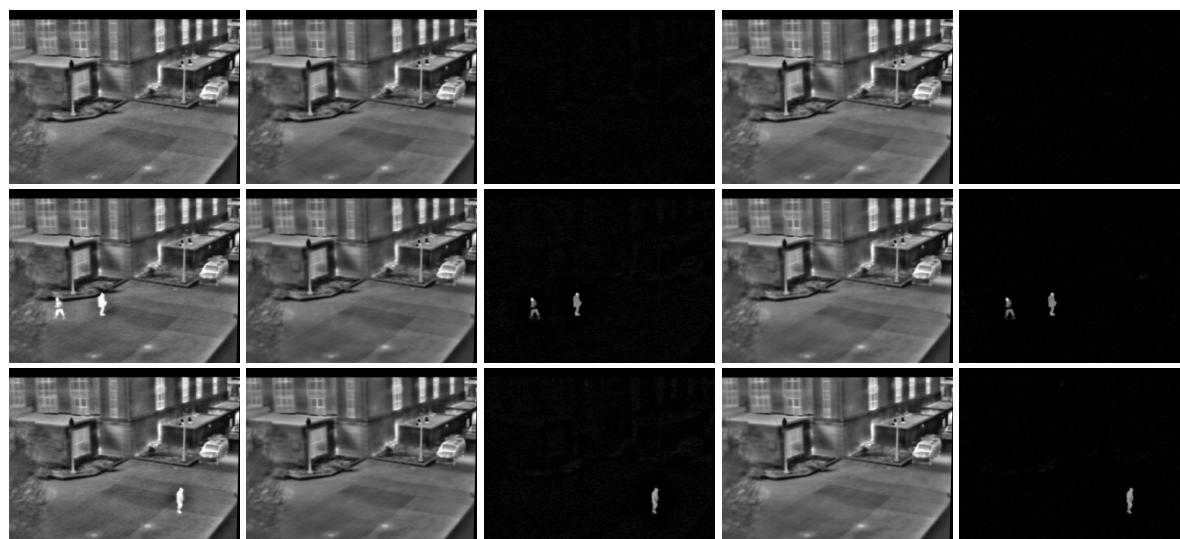


Figure 3. *OSU:* The first column contains three randomly selected frames from the original video. The middle two columns are the separated background and foreground outputs of RCUR, respectively. The right two columns are the separated background and foreground outputs of RPCA, respectively.

property of data and has been used to good effect in solving the subspace clustering problem [1, 2, 23, 31, 36, 49]. Additionally, as pointed out in [44], column selection can produce a dictionary which explains multiple variances which are not orthogonal to each other, in contrast to standard principal component analysis.

In the context of background-foreground separation, we explore this aspect of RCUR decompositions. To do so, we use Algorithm 3.2 to first uniformly sample columns (i.e., frames) of the video sequences according to the sampling requirement of Theorem 3.10; then following background-foreground separation via RPCA, we apply the deterministic sampling method of Algorithm C.1 to select exactly $r = 2$ frames from within those chosen in the first stage. Resulting from this procedure, we obtain two frames of the video background which are highly representative of the rest of the frames; in particular, they approximately minimize $\|L - CC^\dagger L\|_F$, the error when projecting the rest of the video frames onto the chosen two. Figure 4 shows the canonical background frames selected from this procedure compared with the first two basis vectors obtained by taking the SVD of the video matrix D .

Note that in the first row of Figure 4, the left two panes are the two background frames obtained by Algorithm 3.2 and represent the two canonical states of the background: one with light shining through the window onto the floor and one with the light occluded. In contrast, the second SVD basis vector captures only the light portion of the background. In the other video sequences, the background is relatively static, and so our RCUR method produces two very similar canonical background frames. In contrast to this, the second SVD basis vectors are uninformative for these sequences as one might expect. Consequently, the CUR basis vectors capture the essential features of the different states of the background and are easily interpreted by the user. In contrast to this, the SVD basis vectors may appear meaningless in terms of the actual background. This difference stems from the fact that the

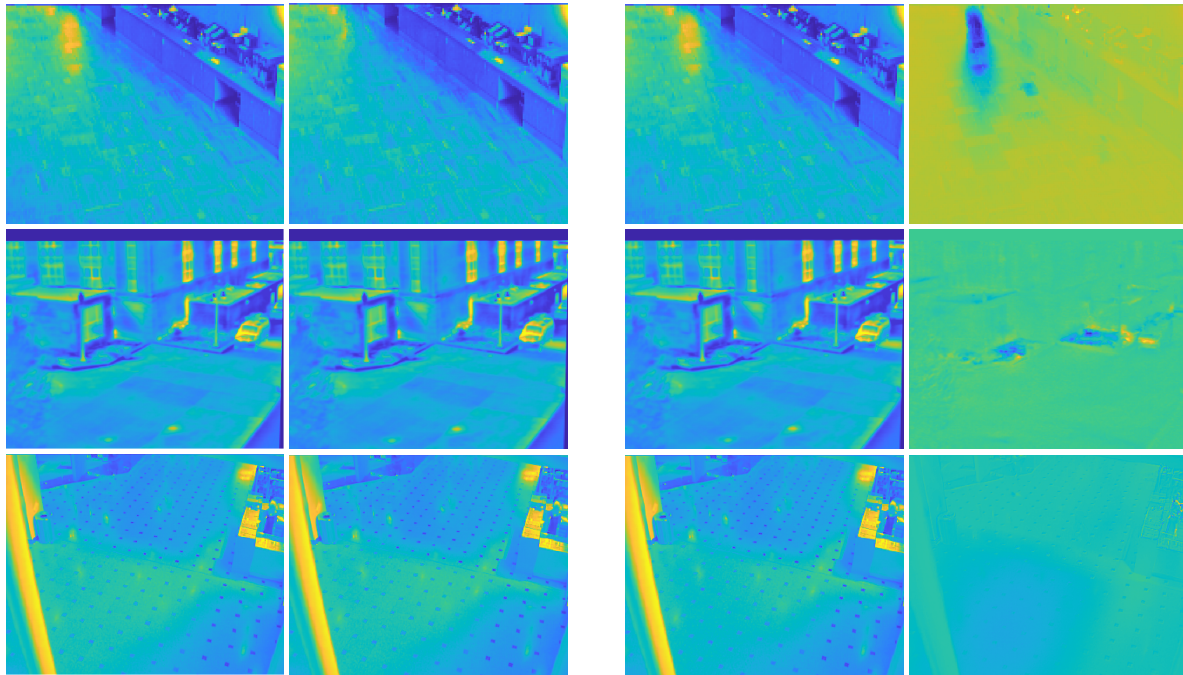


Figure 4. Basis vectors for background of video sequences (Row 1: Restaurant, Row 2: OSU, Row 3: Shoppingmall). In each row, the first two columns are the selected frames by using the deterministic CUR algorithm described above, and the last two columns are the first two orthogonal basis vectors of the SVD of the recovered low-rank matrix, respectively.

column vectors are approximating actual frames, whereas the SVD is capturing orthogonal images corresponding to decreasing variance.

5.2. Application to face modeling. In this section, we consider the problem of robust face modeling. We use the Extended Yale Face Database B (abbreviated *ExtYaleB*) [24] for a benchmark, which includes the face data of 28 human objects under 64 illumination conditions; each face image has size of 168×192 . We vectorize the face images and stack the faces of the person together to form a data matrix, so we have 28 data matrices in $\mathbb{R}^{32,256 \times 64}$. Then in terms of $\mathbf{D} = \mathbf{L} + \mathbf{S}$, the underlying face models form the low-rank $\mathbf{L}$, and the facial occlusions are sparse outlier $\mathbf{S}$. We apply Algorithm 3.1 to each of the face data matrix $\mathbf{D}$ where $r = 1$ is used to extract the face model. Again, since $\mu_2(\mathbf{L}) > \mu_1(\mathbf{L})$ in this problem, we randomly choose column $\tilde{\mathbf{C}} = \mathbf{D}(:, J)$ and row $\tilde{\mathbf{R}} = \mathbf{D}(I, :)$ submatrices of size $m \times 10r \log(n)$ and $25r \log(m) \times n$, respectively. The rest of the experimental setup is same as section 5.1.

In Table 2, we report the total runtime for Algorithm 3.1 and RPCA on the face modeling task. Clearly, Algorithm 3.1 runs faster than RPCA, but the advantage is not as large as in the video background-foreground separation task. The reason is that the face data matrix is too rectangular (i.e., $n = 64$ is too small), which gives less room for acceleration. Moreover, we present the visual face modeling results of a selected human object in Figure 5, wherein we find both algorithms achieve the desired modeling quality.

Table 2
Face data size and runtime.

	Image size	Image number per person	Person number	Total runtime (sec)	
				RCUR	RPCA
ExtYaleB	168×192	64	28	20.16	31.38



Figure 5. Face modeling on ExtYaleB: Visual comparison of the outputs by RCUR and RPCA for face modeling task. The first row contains the original face images. The second and third rows are the face models and the facial occlusions outputted by RCUR, respectively. The last two rows are the face models and the facial occlusions outputted by RPCA, respectively.

Appendix A. Uniform sampling and the proof of Theorem 3.5. Here, we use a special case of the analysis of Tropp [48] to estimate the quantity β in Theorem 3.1.

Theorem A.1 ([48, Lemma 3.4]). Suppose that $\mathbf{L} \in \mathbb{R}^{m \times n}$ has rank r and satisfies condition (A1) and that $\mathbf{V}_{\mathbf{L}} \in \mathbb{R}^{n \times r}$ are its first r right singular vectors. Suppose that $J \subseteq [n]$ is chosen by sampling uniformly without replacement and that $|J| \geq \gamma \mu_2(\mathbf{L})r$ for some $\gamma > 0$. Then the quantity $\beta = \sqrt{|J|/n} \|(\mathbf{V}_{\mathbf{L}}(J, :))^{\dagger}\|_2$ satisfies

$$\beta \leq \frac{1}{\sqrt{1-\delta}} \quad \text{with probability at least} \quad 1 - r \left(\frac{e^{-\delta}}{(1-\delta)^{1-\delta}} \right)^{\gamma} \quad \text{for all } \delta \in [0, 1).$$

Proof of Theorem 3.5. Combine Theorems A.1 and 3.1. ■

Appendix B. Error analysis and proof of Theorem 3.10. In this section we provide a combined error analysis of RPCA (in particular, AltProj [39]) and CUR approximations to obtain the main result concerning Algorithm 3.1, namely Theorem 3.10.

B.1. RPCA. Convergence has been widely studied for many RPCA algorithms [39]. Here, we list one for AltProj [39], which is one of the classic nonconvex RPCA algorithms. We emphasize that our results will stand with other RPCA convergence theories with slight variation.

Theorem B.1 ([39, Theorem 1]). *Let $\mathbf{L}$ and $\mathbf{S}$ satisfy assumptions (A1) and (A2) with $\alpha \leq \mathcal{O}(\frac{1}{(\mu_1(\mathbf{L})\sqrt{\mu_2(\mathbf{L})})^r})$, respectively. With properly chosen parameters, the output of AltProj, $\mathbf{L}_k$, satisfies*

$$\|\mathbf{L} - \mathbf{L}_k\|_2 \leq \varepsilon, \quad \|\mathbf{S} - \mathbf{S}_k\|_\infty \leq \frac{\varepsilon}{\sqrt{mn}} \quad \text{and} \quad \text{supp}(\mathbf{S}_k) \subseteq \text{supp}(\mathbf{S})$$

in $k = \mathcal{O}(r \log(\frac{\|\mathbf{L}\|_2}{\varepsilon}))$ iterations.

Note the statement of Theorem B.1 in the original paper gives a Frobenius norm bound for $\mathbf{L} - \mathbf{L}_k$, but the version stated here follows easily from the fact that $\|\cdot\|_2 \leq \|\cdot\|_F$.

B.2. CUR. Error analysis for CUR decompositions has been considered in various places, typically under assumptions on the column selection scheme. Below is a sample perturbation bound which is independent of the column selection procedure.

Theorem B.2 ([32, Remark 3.14]). *Let $\mathbf{L}$ have rank r and compact SVD $\mathbf{L} = \mathbf{W}\Sigma\mathbf{V}^T$; let $\mathbf{C} = \mathbf{L}(:, J)$, $\mathbf{R} = \mathbf{L}(I, :)$, and $\mathbf{U} = \mathbf{L}(I, J)$ for some $I \subseteq [m]$, $J \subseteq [n]$. Let $\hat{\mathbf{C}} \in \mathbb{R}^{m \times |J|}$ and $\hat{\mathbf{R}} \in \mathbb{R}^{|I| \times n}$ be arbitrary with $\hat{\mathbf{U}} = \hat{\mathbf{C}}(I, :)$. For any Schatten p -norm, if $\sigma_r(\mathbf{U}) = \sigma_{\min}(\mathbf{U}) \geq 12 \max\{\|\hat{\mathbf{R}} - \mathbf{R}\|, \|\hat{\mathbf{C}} - \mathbf{C}\|\}$, then the following holds:*

$$(B.1) \quad \|\mathbf{L} - \hat{\mathbf{C}}\hat{\mathbf{U}}_r^\dagger\hat{\mathbf{R}}\| \leq \left(\frac{7}{6} \left(\|\mathbf{W}_L(I, :)^{\dagger}\| + \|\mathbf{V}_L(J, :)^{\dagger}\| \right) + \frac{25}{6} \|\mathbf{W}_L(I, :)^{\dagger}\| \|\mathbf{V}_L(J, :)^{\dagger}\| + \frac{1}{6} \right) \max\{\|\hat{\mathbf{R}} - \mathbf{R}\|, \|\hat{\mathbf{C}} - \mathbf{C}\|\}.$$

Lemma B.3. *Suppose that $\mathbf{L}$ satisfies (A1). Let $J \subseteq [n]$ with $|J| \geq \gamma_1 \mu_2(\mathbf{L})r$ and $I \subseteq [m]$ with $|I| \geq \gamma_2 \mu_1(\mathbf{L})r$ for some $\gamma_1, \gamma_2 > 0$ be chosen by sampling $[n]$ and $[m]$ uniformly without replacement respectively. Let $\mathbf{U} = \mathbf{L}(I, J)$. Then with probability at least $1 - r(\frac{e^{-\delta}}{(1-\delta)^{1-\delta}})^{\gamma_1} - r(\frac{e^{-\delta}}{(1-\delta)^{1-\delta}})^{\gamma_2}$,*

$$(B.2) \quad \|\mathbf{U}^\dagger\|_2 \leq \frac{1}{(1-\delta)\sigma_{\min}(\mathbf{L})} \sqrt{\frac{mn}{|I||J|}}.$$

Proof. The proof follows by similar argument to the proof of [32, Lemma 3.12]. ■

Theorem B.4. *Given the notations and assumptions of Theorem 3.10, suppose that $\mathbf{C}$ and $\mathbf{R}$ satisfy conditions (A1) and (A2) with sparsity levels $\alpha \leq \mathcal{O}(\frac{1}{(\mu_1(\mathbf{C})\sqrt{\mu_2(\mathbf{C})})^r})$ and $\mathcal{O}(\frac{1}{(\mu_1(\mathbf{R})\sqrt{\mu_2(\mathbf{R})})^r})$, respectively. Suppose that AltProj is run for sufficiently many iterations*

in lines 5 and 6 of Algorithm 3.1 such that $\max\{\|\mathbf{C} - \hat{\mathbf{C}}\|_2, \|\mathbf{R} - \hat{\mathbf{R}}\|_2\} \leq \varepsilon \frac{(1-\delta)\sigma_{\min}(\mathbf{L})}{12} \sqrt{\frac{|I||J|}{mn}}$. Set $\hat{\mathbf{U}} = \hat{\mathbf{C}}(I, :)$. Then

$$(B.3) \quad \left\| \mathbf{L} - \hat{\mathbf{C}} \hat{\mathbf{U}}^\dagger \hat{\mathbf{R}} \right\|_2 \leq \frac{7}{8} \varepsilon \sigma_{\min}(\mathbf{L})$$

with probability at least $1 - \frac{r}{n^{c_2(\delta+(1-\delta)\log(1-\delta))}} - \frac{r}{m^{c_1(\delta-(1-\delta)\log(1-\delta))}}$.

Proof. Since I and J are chosen uniformly, by Lemma B.3 we have that with probability at least $1 - \frac{r}{n^{c_2(\delta+(1-\delta)\log(1-\delta))}} - \frac{r}{m^{c_1(\delta-(1-\delta)\log(1-\delta))}}$,

$$\left\| \mathbf{U}^\dagger \right\|_2 \leq \frac{1}{(1-\delta)\sigma_{\min}(\mathbf{L})} \sqrt{\frac{mn}{|I||J|}}.$$

In addition, $\max\{\|\mathbf{C} - \hat{\mathbf{C}}\|_2, \|\mathbf{R} - \hat{\mathbf{R}}\|_2\} \leq \varepsilon \frac{\sigma_{\min}(\mathbf{L})(1-\delta)}{12} \sqrt{\frac{|I||J|}{mn}}$ and $\hat{\mathbf{U}} = \hat{\mathbf{C}}(I, :)$. Combining these estimates and noting that $\|\mathbf{U}^\dagger\|_2 = \frac{1}{\sigma_{\min}(\mathbf{U})}$, we have

$$\sigma_r(\mathbf{U}) \geq (1-\delta)\sigma_{\min}(\mathbf{L}) \sqrt{\frac{|I||J|}{mn}} > 12\varepsilon \frac{(1-\delta)\sigma_{\min}(\mathbf{L})}{12} \sqrt{\frac{|I||J|}{mn}} \geq 12 \max\left\{\left\|\hat{\mathbf{C}} - \mathbf{C}\right\|_2, \left\|\hat{\mathbf{R}} - \mathbf{R}\right\|_2\right\}$$

with probability at least $1 - \frac{r}{n^{c_2(\delta+(1-\delta)\log(1-\delta))}} - \frac{r}{m^{c_1(\delta-(1-\delta)\log(1-\delta))}}$. The conclusion of the theorem follows from applying Theorem B.2 together with estimates of $\|\mathbf{W}_L(I, :)^{\dagger}\|_2$ and $\|\mathbf{V}_L(J, :)^{\dagger}\|_2$ by Theorem A.1, which implies that with the given success probability,

$$\left\| (\mathbf{V}_L(J, :))^{\dagger} \right\|_2 \leq \sqrt{\frac{n}{(1-\delta)|J|}},$$

and similarly for $\mathbf{W}_L(I, :)$ by replacing n and $|J|$ by m and $|I|$, respectively. ■

Proof of Theorem 3.10. Note that the first terms in the lower bound for $|I|$ and $|J|$ make the conclusions of Theorem 3.5 and Lemma B.3 valid, while the second terms ensure 2α -sparsity of $\mathbf{S}(I, :)$ and $\mathbf{S}(:, J)$ (Proposition 4.5). Consequently, $\mathbf{C} = \mathbf{L}(:, J)$ and $\mathbf{R} = \mathbf{L}(I, :)$ satisfy condition (A1) with the incoherence parameters in Theorem 3.1. Additionally, to apply Theorem B.1 to $\tilde{\mathbf{C}} = \mathbf{D}(:, J)$ and $\tilde{\mathbf{R}} = \mathbf{D}(I, :)$, we require that $2\alpha \leq \mathcal{O}(\frac{1}{(\mu_1(\mathbf{C})\sqrt{\mu_2(\mathbf{C})})^r})$ and $2\alpha \leq \mathcal{O}(\frac{1}{(\mu_1(\mathbf{R})\sqrt{\mu_2(\mathbf{R})})^r})$. Using the bounds of Theorem 3.5, which hold with high probability given how I and J were selected, we see that the assumption that $\alpha = \mathcal{O}(\frac{1-\delta}{\kappa(\mathbf{L})^2(\mu_1(\mathbf{L})\sqrt{\mu_2(\mathbf{L})})})$ guarantees that α is of the required order and that, in turn, Theorem B.1 holds with the given success probability.

Finally, Theorem B.4 holds on account of the above observations and the fact that the iteration count being $\mathcal{O}(r \log(\sqrt{\frac{mn}{|I||J|}} \frac{\kappa(\mathbf{L})}{(1-\delta)\varepsilon}))$ implies the bounds on $\|\mathbf{C} - \hat{\mathbf{C}}\|_2$ and $\|\mathbf{R} - \hat{\mathbf{R}}\|_2$. Hence, we conclude that the outputs $\hat{\mathbf{C}}, \hat{\mathbf{R}}$ of Algorithm 3.1 satisfy

$$\left\| \mathbf{L} - \hat{\mathbf{L}} \right\|_2 = \left\| \mathbf{L} - \hat{\mathbf{C}}(\hat{\mathbf{C}}(I, :))^{\dagger} \hat{\mathbf{R}} \right\|_F \leq \frac{7}{8} \varepsilon \sigma_{\min}(\mathbf{L})$$

with the given success probability, which completes the proof. ■

Appendix C. Uniform + deterministic: Proof of Theorem 3.11. The deterministic greedy column selection algorithm of [3] is reproduced as Algorithm C.1. Here, we provide the proof of Theorem 3.11, beginning with a bound on the pseudoinverse of $\mathbf{V}_L(:, J)$, assuming that the column submatrix of $\mathbf{L}$ and $\mathbf{D}$ are not too far apart. We will find use for this result later as well when we analyze the purely deterministic column sampling method.

Recall that Theorem 3.11 analyzes the error of approximation for the hybrid procedure of uniformly sampling column and row submatrices of $\mathbf{D} = \mathbf{L} + \mathbf{S}$, running RPCA on these to output $\hat{\mathbf{C}}$ and $\hat{\mathbf{R}}$, and finally using Algorithm C.1 to select exactly $r = \text{rank}(\mathbf{L})$ columns and rows of $\hat{\mathbf{C}}$ and $\hat{\mathbf{R}}$, respectively, to give a compact approximation of $\mathbf{L}$.

Theorem C.1. Let $\mathbf{L} \in \mathbb{R}^{m \times n}$ with $\text{rank}(\mathbf{L}) = r$. Let $\tilde{\mathbf{L}} = \mathbf{L} + \mathbf{E}$ with compact SVD $\tilde{\mathbf{L}} = \tilde{\mathbf{W}}\tilde{\Sigma}\tilde{\mathbf{V}}^T$, and suppose $\|\mathbf{E}\|_2 \leq \frac{1}{2}\sigma_{\min}(\mathbf{L})$. Suppose that $I \subseteq [m]$ and $J \subseteq [n]$ with $|I| = |J| = r$ are obtained by applying Algorithm C.1 on $\tilde{\mathbf{W}}$ and $\tilde{\mathbf{V}}$. Suppose that

$$\|\mathbf{L}(:, J) - \tilde{\mathbf{L}}(:, J)\|_2 = \|\mathbf{W}_L \Sigma_L (\mathbf{V}_L(J, :))^T - \tilde{\mathbf{W}} \tilde{\Sigma} (\tilde{\mathbf{V}}(J, :))^T\|_2 \leq \frac{1}{8} \sqrt{\frac{|J|}{rn}} \sigma_{\min}(\mathbf{L}).$$

Then $\|(\mathbf{V}_L(J, :))^\dagger\|_2 \leq 4\kappa(\mathbf{L}) \sqrt{\frac{rn}{|J|}}$.

Proof. Set $\tilde{\mathbf{C}} = \tilde{\mathbf{W}}\tilde{\Sigma}(\tilde{\mathbf{V}}(J, :))^T$ and $\mathbf{C} = \mathbf{L}(:, J) = \mathbf{W}_L \Sigma_L (\mathbf{V}_L(J, :))^T$. Then we have

$$(C.1) \quad \|(\mathbf{V}_L(J, :))^\dagger\|_2 = \|\mathbf{C}^\dagger \mathbf{W}_L \Sigma_L\|_2 \leq \sigma_{\max}(\mathbf{L}) \|\mathbf{C}^\dagger\|_2.$$

Notice that $\text{rank}(\mathbf{C}) = \text{rank}(\tilde{\mathbf{C}}) = r$. By [45, Theorems 3.1–3.4], we have

$$\|\mathbf{C}^\dagger - \tilde{\mathbf{C}}^\dagger\|_2 \leq 2 \|\mathbf{C}^\dagger\|_2 \|\tilde{\mathbf{C}}^\dagger\|_2 \|\mathbf{C} - \tilde{\mathbf{C}}\|_2.$$

Algorithm C.1. A deterministic greedy removal algorithm for subset selection [3]

- 1: **Input:** $\mathbf{X} \in \mathbb{R}^{r \times m}$ with $\text{rank}(\mathbf{X}) = r$; k : sampling size.
 - 2: **Output:** $\mathcal{S} \subseteq [m]$: set of cardinality k .
 - 3: **Initialization:** $\mathcal{S}_0 := [m]$
 - 4: Compute the SVD of $\mathbf{X}_{\mathcal{S}_0}$: $\mathbf{X}_{\mathcal{S}_0} = \mathbf{U}^{(0)} \Sigma^{(0)} \mathbf{Y}^{(0)}$
 - 5: **for** $i = 1, 2, \dots, m - k$ **do**
 - 6: Let the singular values of $\mathbf{X}_{\mathcal{S}_{i-1}}$ be $\sigma_1^{(i-1)}, \dots, \sigma_r^{(i-1)}$.
 - 7: Let the columns of $\mathbf{Y}^{(i-1)}$ be $\{y_k^{(i-1)}\}_{k \in \mathcal{S}_{i-1}}$. Denote by y_{kj} the j th element of $y_k^{(i-1)}$.
 - 8: $k_i := \underset{k \in \mathcal{S}_{i-1}, \|y_k^{(i-1)}\|_2 < 1}{\operatorname{argmin}} \left(\frac{\sum_{j=1}^r (y_{kj}^{(i-1)} / \sigma_k^{(i-1)})^2}{1 - \|y_k^{(i-1)}\|_2^2} \right)$.
 - 9: Set $\mathcal{S}_i := \mathcal{S}_{i-1} \setminus \{k_i\}$
 - 10: Downdate the SVD of $\mathbf{X}_{\mathcal{S}_{i-1}}$ to obtain an SVD of $\mathbf{X}_{\mathcal{S}_i} = \mathbf{U}^{(i)} \Sigma^{(i)} \mathbf{Y}^{(i)}$ (see [30])
 - 11: **end for**
-

In addition, by Proposition 3.3,

$$\begin{aligned}
 \|\tilde{\mathbf{C}}^\dagger\|_2 &= \left\| \left((\tilde{\mathbf{V}}(J, :))^T \right)^\dagger \tilde{\mathbf{\Sigma}}^\dagger \tilde{\mathbf{W}}^T \right\|_2 \\
 &\leq \left\| \left(\tilde{\mathbf{V}}(J, :)^T \right)^\dagger \right\|_2 \|\tilde{\mathbf{\Sigma}}^\dagger\|_2 \\
 &\leq \sqrt{\frac{rn}{|J|}} \frac{1}{\sigma_{\min}(\tilde{\mathbf{L}})} \\
 &\leq \sqrt{\frac{rn}{|J|}} \frac{1}{\sigma_{\min}(\mathbf{L}) - \|\mathbf{E}\|_2} \\
 &\leq 2\sqrt{\frac{rn}{|J|}} \frac{1}{\sigma_{\min}(\mathbf{L})}.
 \end{aligned}
 \tag{C.2}$$

Since $\|\mathbf{W}_L \mathbf{\Sigma}_L (\mathbf{V}_L(J, :))^T - \tilde{\mathbf{W}} \tilde{\mathbf{\Sigma}} (\tilde{\mathbf{V}}(J, :))^T\|_2 \leq \frac{1}{8} \sqrt{\frac{|J|}{rn}} \sigma_{\min}(\mathbf{L})$, $\|\mathbf{C}^\dagger\|_2 \leq \frac{\|\tilde{\mathbf{C}}^\dagger\|_2}{1 - 2\|\tilde{\mathbf{C}}^\dagger\|_2 \|\mathbf{C} - \tilde{\mathbf{C}}\|_2} \leq 2\|\tilde{\mathbf{C}}^\dagger\|_2$. Combining, (C.1) and (C.2), we have

$$\|\mathbf{V}_L(J, :)^T\|_2 \leq \sqrt{\frac{rn}{|J|}} \frac{4\sigma_{\max}(\mathbf{L})}{\sigma_{\min}(\mathbf{L})} = 4\kappa(\mathbf{L}) \sqrt{\frac{rn}{|J|}}. \quad \blacksquare
 \tag{C.3}$$

Proof of Theorem 3.11. The assumptions of Theorem 3.10 imply that $\tilde{\mathbf{C}} = \mathbf{D}(:, J) = \mathbf{L}(:, J) + \mathbf{S}(:, J)$ and $\tilde{\mathbf{R}} = \mathbf{D}(I, :) = \mathbf{L}(I, :) + \mathbf{S}(I, :)$ satisfy the necessary incoherence and sparsity conditions for RPCA to be successful. Thus we pick up most of the way through the proof of Theorem 3.10 above and note that the new assumption on the iteration count for AltProj in the statement of this theorem implies by Theorem B.1 that we have

$$\max \left\{ \|\mathbf{C} - \hat{\mathbf{C}}\|_2, \|\mathbf{R} - \hat{\mathbf{R}}\|_2 \right\} \leq \frac{\varepsilon(1 - \delta)^2 \sigma_{\min}(\mathbf{L})}{229\kappa(\mathbf{L})^4 r \sqrt{mn\mu_1(\mathbf{L})\mu_2(\mathbf{L})}}.
 \tag{C.4}$$

According to Theorem A.1, the following statements hold:

$$\begin{aligned}
 \|\mathbf{W}_L(I, :)^T\|_2 &\leq \sqrt{\frac{m}{(1 - \delta)|I|}}, \\
 \|\mathbf{V}_L(J, :)^T\|_2 &\leq \sqrt{\frac{n}{(1 - \delta)|J|}}
 \end{aligned}
 \tag{C.5}$$

with probability at least $1 - \frac{r}{n^{c_2(\delta + (1 - \delta) \log(1 - \delta))}} - \frac{r}{m^{c_1(\delta - (1 - \delta) \log(1 - \delta))}}$. From (C.4), we have

$$\begin{aligned}
 \|\hat{\mathbf{C}}_1 - \mathbf{C}(:, J_1)\|_2 &\leq \frac{\varepsilon(1 - \delta)^2 \sigma_{\min}(\mathbf{L})}{229\kappa(\mathbf{L})^4 r \sqrt{mn\mu_1(\mathbf{L})\mu_2(\mathbf{L})}}, \\
 \|\hat{\mathbf{R}}_1 - \mathbf{R}(I_1, :)\|_2 &\leq \frac{\varepsilon(1 - \delta)^2 \sigma_{\min}(\mathbf{L})}{229\kappa(\mathbf{L})^4 r \sqrt{mn\mu_1(\mathbf{L})\mu_2(\mathbf{L})}}.
 \end{aligned}$$

In addition,

$$\frac{1}{8} \sqrt{\frac{|J_1|}{r|J|}} \sigma_{\min}(\mathbf{C}) = \frac{1}{8\sqrt{|J|}} \frac{1}{\|\mathbf{C}^\dagger\|_2} \geq \frac{1}{8\sqrt{|J|}} \frac{1}{\|\mathbf{L}^\dagger\|_2 \left\| (\mathbf{V}_L(J, :))^\dagger \right\|_2} \geq \frac{\sqrt{1-\delta} \sigma_{\min}(\mathbf{L})}{8\sqrt{n}}.$$

Thus, $\|\hat{\mathbf{C}}_1 - \mathbf{C}(:, J_1)\|_2 \leq \frac{1}{8} \sqrt{\frac{|J_1|}{r|J|}} \sigma_{\min}(\mathbf{C})$. Similarly, we have $\|\hat{\mathbf{R}}_1 - \mathbf{R}(I_1, :)\|_2 \leq \frac{1}{8} \sqrt{\frac{|I_1|}{r|I|}} \sigma_{\min}(\mathbf{R})$. By Theorem C.1,

$$\begin{aligned} \left\| \mathbf{V}_C(J_1, :)^{\dagger} \right\| &\leq 4\kappa(\mathbf{C}) \sqrt{|J|}, \\ \left\| \mathbf{W}_R(I_1, :)^{\dagger} \right\| &\leq 4\kappa(\mathbf{R}) \sqrt{|I|}. \end{aligned}$$

Since $\mathbf{C}(:, J_1) = \mathbf{W}_C \boldsymbol{\Sigma}_C (\mathbf{V}_C(J_1, :))^T = \mathbf{W}_L \boldsymbol{\Sigma}_L (\mathbf{V}_L(\tilde{J}_1, :))^T$, we have

$$\begin{aligned} \left\| \mathbf{V}_L(\tilde{J}_1, :)^{\dagger} \right\|_2 &\leq \left\| \mathbf{V}_C(J_1, :)^{\dagger} \right\|_2 \|\boldsymbol{\Sigma}_L\|_2 \|\boldsymbol{\Sigma}_C^{-1}\|_2 \\ &\stackrel{\text{Lemma 4.3}}{=} 4\kappa(\mathbf{C}) \sqrt{|J|} \|\mathbf{L}\|_2 \left\| \mathbf{L}^\dagger \right\|_2 \left\| \mathbf{V}_L(J, :)^{\dagger} \right\|_2 \\ &\stackrel{\text{Lemma 4.4}}{\leq} 4\kappa(\mathbf{L})^2 \frac{|J| \sqrt{\mu_2 r}}{\sqrt{n}} \left\| \mathbf{V}_L(J, :)^{\dagger} \right\|_2^2 \\ &\stackrel{(\text{C.5})}{\leq} 4\kappa(\mathbf{L})^2 \frac{\sqrt{\mu_2(\mathbf{L}) r n}}{1 - \delta}. \end{aligned}$$

Similarly, we have

$$\left\| \mathbf{W}_L(I_1, :)^{\dagger} \right\| \leq 4\kappa(\mathbf{L})^2 \frac{\sqrt{\mu_1(\mathbf{L}) r m}}{1 - \delta}.$$

Thus,

$$\left\| \mathbf{U}^\dagger \right\|_2 \leq \frac{16\kappa^4(\mathbf{L}) r \sqrt{\mu_1(\mathbf{L}) \mu_2(\mathbf{L}) m n}}{(1 - \delta)^2 \sigma_{\min}(\mathbf{L})},$$

i.e., $\sigma_r(\mathbf{U}) \geq \frac{(1-\delta)^2 \sigma_{\min}(\mathbf{L})}{16\kappa^4(\mathbf{L}) r \sqrt{\mu_1(\mathbf{L}) \mu_2(\mathbf{L}) m n}} \geq 12 \max\{\|\hat{\mathbf{C}}_1 - \mathbf{L}(:, \tilde{J}_1)\|_2, \|\hat{\mathbf{R}}_1 - \mathbf{L}(\tilde{I}_1, :)\|_2\}$. Using the bounds of $\|(\mathbf{W}(I, :))^{\dagger}\|_2$, $\|(\mathbf{V}(J, :))^{\dagger}\|_2$ and Theorem B.2, we obtain

$$\left\| \mathbf{L} - \hat{\mathbf{C}}_1 \hat{\mathbf{U}}_1^\dagger \hat{\mathbf{R}}_1 \right\| \leq \frac{1}{3} \varepsilon \sigma_r(\mathbf{L})$$

after some simple calculations. ■

Appendix D. Deterministic column and row selection. In this section, we study what happens when we apply the greedy algorithm (Algorithm C.1) on the singular vectors of $\mathbf{D} = \mathbf{L} + \mathbf{S}$ to choose column and row submatrices. We primarily consider how the incoherence

Algorithm D.1. Initialization by one step of alternating projections

-
- 1: **Input:** $\mathbf{D} = \mathbf{L} + \mathbf{S}$: matrix to be split; r : rank of $\mathbf{L}$; η : thresholding parameter.
 - 2: **Output:** $\mathbf{L}_0, \mathbf{S}_0$
 - 3: $\zeta_0 = \eta \cdot \sigma_1(\mathbf{D})$
 - 4: $\mathbf{S}_0 = \mathcal{T}_{\zeta_0}(\mathbf{D})$ $\triangleright \mathcal{T}_{\zeta}$: hard thresholding operation with thresholding value ζ
 - 5: $\mathbf{L}_0 = \mathcal{D}_r(\mathbf{D} - \mathbf{S}_0)$ $\triangleright \mathcal{D}_r$: truncated rank r SVD operator
-

of $\mathbf{C} = \mathbf{L}(:, J)$ relates to that of $\mathbf{L}$. Our main result is that, if we first initialize $\mathbf{D}$ by one step of alternating projections [39] with a judiciously chosen hard thresholding parameter, then the resulting approximation $\mathbf{L}_0$ to $\mathbf{L}$ is good enough to ensure that the incoherence does not inflate significantly. To achieve this, we study conditions under which the hypothesis of Theorem B.4 holds, namely such that $\|\mathbf{L}(:, J) - \mathbf{L}_0(:, J)\|_2 \leq \frac{1}{8} \sqrt{\frac{|J|}{rn}} \sigma_{\min}(\mathbf{L})$, where $\mathbf{L}_0$ is the output of Algorithm D.1. To proceed, we first state the initialization required.

Here, $\mathcal{T}_{\zeta_0}(\mathbf{D})$ is the hard thresholding operator whose output is (entrywise) $\mathbf{D}_{ij}$ if $\mathbf{D}_{ij} \geq \zeta_0$ and 0 otherwise. Given a matrix $\mathbf{A}$, $\mathcal{D}_r(\mathbf{A})$ is its projection onto the set of rank r matrices (e.g., $\mathcal{D}_r(\mathbf{A}) = \mathbf{U}_{\mathbf{A},r} \Sigma_{\mathbf{A},r} \mathbf{V}_{\mathbf{A},r}^T$).

Theorem D.1. *Let $\mathbf{L} \in \mathbb{R}^{m \times n}$ be a rank r $\{\mu_1(\mathbf{L}), \mu_2(\mathbf{L})\}$ -incoherent matrix and $\mathbf{S} \in \mathbb{R}^{m \times n}$ be an α -sparse matrix. Let $\mu := \max\{\mu_1(\mathbf{L}), \mu_2(\mathbf{L})\}$. If the sparsity level satisfies $\alpha \leq \frac{1}{256\mu^{3/2}r^2\kappa(\mathbf{L})}$ and the thresholding parameter in Algorithm D.1 obeys $\frac{\mu r \sigma_{\max}(\mathbf{L})}{\sqrt{mn} \sigma_{\max}(\mathbf{D})} \leq \eta \leq \frac{3\mu r \sigma_{\max}(\mathbf{L})}{\sqrt{mn} \sigma_{\max}(\mathbf{D})}$, then the outputs of Algorithm D.1 satisfy*

$$\begin{aligned} \|\mathbf{L} - \mathbf{L}_0\|_2 &\leq 8\alpha\mu r \sigma_{\max}(\mathbf{L}), \\ \max_i \|(\mathbf{L}_0 - \mathbf{L})^T \mathbf{e}_i\|_2 &\leq \frac{32\alpha\mu^{1.5}r^{1.5}}{\sqrt{m}} \sigma_{\max}(\mathbf{L}), \\ \max_j \|(\mathbf{L}_0 - \mathbf{L}) \mathbf{e}_j\|_2 &\leq \frac{32\alpha\mu^{1.5}r^{1.5}}{\sqrt{n}} \sigma_{\max}(\mathbf{L}), \\ \|\mathbf{S} - \mathbf{S}_0\|_{\infty} &\leq \frac{\mu r}{n} \sigma_{\max}(\mathbf{L}), \quad \text{and} \quad \text{supp}(\mathbf{S}_0) \subseteq \text{supp}(\mathbf{S}). \end{aligned}$$

Note that since $\|\mathbf{A}(I, :)\|_F^2 \leq \sum_{i \in I} \|\mathbf{A}^T \mathbf{e}_i\|_2^2$ and $\|\mathbf{A}(:, J)\|_F^2 \leq \sum_{j \in J} \|\mathbf{A} \mathbf{e}_j\|_2^2$, Theorem D.1 implies

$$\begin{aligned} \|\mathbf{L}(I, :) - \mathbf{L}_0(I, :)\|_F &\leq \frac{32\alpha\mu^{1.5}r^{1.5}\sqrt{|I|}}{\sqrt{m}} \sigma_{\max}(\mathbf{L}) \leq \frac{1}{8} \sqrt{\frac{|I|}{rm}} \sigma_{\min}(\mathbf{L}), \\ \|\mathbf{L}(:, J) - \mathbf{L}_0(:, J)\|_F &\leq \frac{32\alpha\mu^{1.5}r^{1.5}\sqrt{|J|}}{\sqrt{n}} \sigma_{\max}(\mathbf{L}) \leq \frac{1}{8} \sqrt{\frac{|J|}{rn}} \sigma_{\min}(\mathbf{L}), \\ \|\mathbf{L} - \mathbf{L}_0\|_F &\leq \sqrt{2r} \|\mathbf{L} - \mathbf{L}_0\|_2 \leq \frac{1}{2} \sigma_{\min}(\mathbf{L}), \end{aligned}$$

where the second parts of the inequalities follow from bound of α .

Remark D.2. Note that the conclusion of Theorem D.1 implies that $\mathbf{L}_0$ satisfies the necessary requirement for Theorem C.1. Consequently, we find that the output of Algorithm C.1 applied to $\mathbf{L}_0$ yields a column submatrix $\mathbf{C} = \mathbf{L}(:, J)$ with parameter $\beta \leq 4\kappa(\mathbf{L})\sqrt{r}$, and thus incoherence $\mu_2(\mathbf{C}) \leq 16\kappa(\mathbf{L})^2 r$ according to Theorem 3.1. This bound for β is a factor of $\sqrt{r}$ better than simply applying the greedy algorithm directly to $\mathbf{V}_L$ due to good initialization to find $\mathbf{L}_0$. This method also has the benefit of being tractable as it does not require knowledge of $\mathbf{V}_L$.

Appendix E. Proof of Theorem D.1. As shown in prior art [10, 12, 11, 39], the convergence analysis of the alternating projections-based RPCA algorithms can be reduced to the case of symmetric matrices, since nonsymmetric matrix recovery problems can be cast as problems with respect to symmetric augmented matrices. For more details about how to reduce the general RPCA problems to symmetric cases, we refer the interested reader to [10, section 3.2]. Firstly, we shall present two technical lemmata in the symmetric setting.

Lemma E.1 ([39, Lemma 4]). Let $\mathbf{S} \in \mathbb{R}^{n \times n}$ be an α -sparse symmetric matrix. Then, the following inequality holds:

$$\|\mathbf{S}\|_2 \leq \alpha n \|\mathbf{S}\|_\infty.$$

Lemma E.2 ([39, Lemma 5]). Let $\mathbf{S} \in \mathbb{R}^{n \times n}$ be an α -sparse symmetric matrix. Let $\mathbf{W} \in \mathbb{R}^{n \times r}$ be an orthogonal matrix with μ -incoherence, i.e., $\|\mathbf{W}^T \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu r}{n}}$ for all i . Then

$$\|\mathbf{W}^T \mathbf{S}^a \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu r}{n}} (\alpha n \|\mathbf{S}\|_\infty)^a$$

holds for all i and $a \geq 0$.

We are now ready to prove Theorem D.1.

Proof of Theorem D.1. As discussed above, we only need to prove the theorem under symmetric setting, and it can be generalized to nonsymmetric matrices. The proof consists of four parts. The first three parts have been shown in the proof of [12, Theorem 2], but we still include them here for the completeness.

Part 1. First, note that

$$\|\mathbf{L}\|_\infty = \max_{ij} |\mathbf{e}_i^T \mathbf{W}_L \Sigma_L \mathbf{V}_L^T \mathbf{e}_j| \leq \max_{ij} \|\mathbf{W}_L^T \mathbf{e}_i\|_2 \|\Sigma_L\|_2 \|\mathbf{V}_L^T \mathbf{e}_j\|_2 \leq \frac{\mu r}{n} \sigma_1(\mathbf{L}),$$

where the last inequality follows from μ -incoherence of $\mathbf{L}$. Since $\eta \geq \frac{\mu r \sigma_1(\mathbf{L})}{n \sigma_1(\mathbf{D})}$, we get

$$(E.1) \quad \|\mathbf{L}\|_\infty \leq \eta \sigma_1(\mathbf{D}) =: \zeta_0.$$

Considering the entries of $\mathbf{S}_0$, we notice

$$[\mathbf{S}_0]_{ij} = [\mathcal{T}_{\zeta_0}(\mathbf{S} + \mathbf{L})]_{ij} = \begin{cases} \mathcal{T}_{\zeta_0}([\mathbf{S} + \mathbf{L}]_{ij}), & (i, j) \in \text{supp}(\mathbf{S}), \\ \mathcal{T}_{\zeta_0}([\mathbf{L}]_{ij}), & (i, j) \notin \text{supp}(\mathbf{S}), \end{cases}$$

since $[\mathbf{S}]_{ij} = 0$ outside of its support. Together with (E.1), we have $[\mathbf{S}_0]_{ij} = 0$ for all $(i, j) \notin \text{supp}(\mathbf{S})$, which implies $\text{supp}(\mathbf{S}_0) \subseteq \text{supp}(\mathbf{S})$. Furthermore, we have

$$[\mathbf{S} - \mathbf{S}_0]_{ij} = \begin{cases} 0, & (i, j) \notin \text{supp}(\mathbf{S}), \\ [\mathbf{L}]_{ij}, & (i, j) \in \text{supp}(\mathbf{S}_0), \\ [\mathbf{S}]_{ij} & (i, j) \in \text{supp}(\mathbf{S}) \setminus \text{supp}(\mathbf{S}_0), \end{cases} \leq \begin{cases} 0, \\ \|\mathbf{L}\|_\infty, \\ \|\mathbf{L}\|_\infty + \zeta_0 \end{cases} \leq \begin{cases} 0, \\ \frac{\mu r}{n} \sigma_1(\mathbf{L}), \\ \frac{4\mu r}{n} \sigma_1(\mathbf{L}), \end{cases}$$

where the last inequality follows from $\eta \leq \frac{3\mu r \sigma_1(\mathbf{L})}{n \sigma_1(\mathbf{D})}$, which implies $\zeta_0 \leq \frac{3\mu r}{n} \sigma_1(\mathbf{L})$. Overall, we get

$$(E.2) \quad \text{supp}(\mathbf{S}_0) \subseteq \text{supp}(\mathbf{S}) \quad \text{and} \quad \|\mathbf{S} - \mathbf{S}_0\|_\infty \leq \frac{4\mu r}{n} \sigma_1(\mathbf{L}).$$

Moreover, this implies that $\mathbf{S} - \mathbf{S}_0$ is also α -sparse.

Part 2. Since $\mathbf{L}_0 = \mathcal{D}_r(\mathbf{S} - \mathbf{S}_0)$ is the best rank r approximation of $\mathbf{D} - \mathbf{S}_0$, we have

$$\begin{aligned} \|\mathbf{L} - \mathbf{L}_0\|_2 &\leq \|\mathbf{L} - (\mathbf{D} - \mathbf{S}_0)\|_2 + \|(\mathbf{D} - \mathbf{S}_0) - \mathbf{L}_0\|_2 \\ &\leq 2 \|\mathbf{L} - (\mathbf{D} - \mathbf{S}_0)\|_2 \\ &= 2 \|\mathbf{L} - (\mathbf{L} + \mathbf{S} - \mathbf{S}_0)\|_2 \\ &= 2 \|\mathbf{S} - \mathbf{S}_0\|_2 \\ &\leq 2\alpha n \|\mathbf{S} - \mathbf{S}_0\|_\infty, \end{aligned}$$

where the last inequality uses Lemma E.1. By applying (E.2), we have

$$(E.3) \quad \|\mathbf{L} - \mathbf{L}_0\|_2 \leq 8\alpha\mu r \sigma_1(\mathbf{L}).$$

Part 3. Let λ_i denote the i th eigenvalue of $\mathbf{D} - \mathbf{S}_0$ ordered as $|\lambda_1| \geq |\lambda_2| \geq \dots \geq |\lambda_n|$. Since $\mathbf{D} - \mathbf{S}_0 = \mathbf{L} + (\mathbf{S} - \mathbf{S}_0)$, by Weyl's inequality, we have

$$(E.4) \quad |\sigma_i(\mathbf{L}) - |\lambda_i|| \leq \|\mathbf{S} - \mathbf{S}_0\|_2 \leq \alpha n \|\mathbf{S} - \mathbf{S}_0\|_\infty \leq \frac{\sigma_r(\mathbf{L})}{8}$$

for all i , where the last inequality follows from the fact that $\alpha \leq \frac{1}{32\mu r \kappa(\mathbf{L})}$. Therefore, we have

$$(E.5) \quad \frac{7}{8} \sigma_i(\mathbf{L}) \leq |\lambda_i| \leq \frac{9}{8} \sigma_i(\mathbf{L}), \quad 1 \leq i \leq r,$$

and

$$(E.6) \quad \frac{\|\mathbf{S} - \mathbf{S}_0\|_2}{|\lambda_r|} \leq \frac{\frac{\sigma_r(\mathbf{L})}{8}}{\frac{7\sigma_r(\mathbf{L})}{8}} = \frac{1}{7}.$$

Part 4. Let $\mathbf{D} - \mathbf{S}_0 = [\mathbf{W}_0, \ddot{\mathbf{W}}_0] \begin{bmatrix} \mathbf{\Lambda} & \mathbf{0} \\ \mathbf{0} & \ddot{\mathbf{\Lambda}} \end{bmatrix} [\mathbf{W}_0, \ddot{\mathbf{W}}_0]^T = \mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T + \ddot{\mathbf{W}}_0 \ddot{\mathbf{\Lambda}} \ddot{\mathbf{W}}_0^T$ be the eigenvalue decomposition of $\mathbf{D} - \mathbf{S}_0$, where the eigenvalues are sorted by magnitude and $\mathbf{\Lambda}$ consists of the r largest ones while $\ddot{\mathbf{\Lambda}}$ contains the rest. Correspondingly, $\mathbf{W}_0$ contains the first r eigenvectors,

and $\tilde{\mathbf{W}}_0$ has the rest. By the symmetric setting, we have $\mathbf{L}_0 = \mathcal{D}_r(\mathbf{D} - \mathbf{S}_0) = \mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T$. Denote $\mathbf{E} = \mathbf{S} - \mathbf{S}_0$ and $\mathbf{W}(:, i)$ to be the i th eigenvector of $\mathbf{D} - \mathbf{S}_0 = \mathbf{L} + \mathbf{E}$. So we can see

$$(\mathbf{L} + \mathbf{E})\mathbf{W}(:, i) = \lambda_i \mathbf{W}(:, i)$$

for $1 \leq i \leq r$. Hence,

$$(E.7) \quad \mathbf{W}(:, i) = \left(I - \frac{\mathbf{E}}{\lambda_i}\right)^{-1} \frac{\mathbf{L}}{\lambda_i} \mathbf{W}(:, i) = \sum_{j=0}^{\infty} \left(\frac{\mathbf{E}}{\lambda_i}\right)^j \frac{\mathbf{L}}{\lambda_i} \mathbf{W}(:, i)$$

for all $1 \leq i \leq r$. Note that the expansion in the last equality is valid since (E.6) implies $\frac{\|\mathbf{E}\|_2}{|\lambda_i|} < \frac{1}{7} \leq 1$ for all $1 \leq i \leq r$.

We will first prove the *row version* of the inequality. By applying (E.7), we get

$$\begin{aligned} & \max_i \|(\mathbf{L}_0 - \mathbf{L})^T \mathbf{e}_i\|_2 \\ &= \max_i \|(\mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T - \mathbf{L})^T \mathbf{e}_i\|_2 \\ &= \max_i \left\| \left(\mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T \mathbf{L} - \mathbf{L} + \sum_{a+b>0} \mathbf{E}^a \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{L} \mathbf{E}^b \right)^T \mathbf{e}_i \right\|_2 \\ &\leq \max_i \|(\mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T \mathbf{L} - \mathbf{L})^T \mathbf{e}_i\|_2 + \sum_{a+b>0} \max_i \left\| \left(\mathbf{E}^a \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{L} \mathbf{E}^b \right)^T \mathbf{e}_i \right\|_2 \\ &:= \mathbf{Y}_0 + \sum_{a+b>0} \mathbf{Y}_{ab}. \end{aligned}$$

We will first bound $\mathbf{Y}_0$. By the incoherence property of $\mathbf{L}$, i.e., $\max_i \|\mathbf{W} \mathbf{W}^T \mathbf{e}_i\|_2 \leq \sqrt{\frac{\mu r}{n}}$, we have

$$\begin{aligned} \mathbf{Y}_0 &= \max_i \|(\mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T \mathbf{L} - \mathbf{L})^T \mathbf{e}_i\|_2 \\ &= \max_i \|(\mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T \mathbf{L} - \mathbf{L})^T \mathbf{W} \mathbf{W}^T \mathbf{e}_i\|_2 \\ &\leq \max_i \|\mathbf{W} \mathbf{W}^T \mathbf{e}_i\|_2 \|\mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T \mathbf{L} - \mathbf{L}\|_2 \\ &\leq \sqrt{\frac{\mu r}{n}} \|\mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T \mathbf{L} - \mathbf{L}\|_2, \end{aligned}$$

where the second equation follows from the fact that $\mathbf{L} = \mathbf{W} \mathbf{W}^T \mathbf{L}$.

Note that $\mathbf{L} = \mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T + \tilde{\mathbf{W}}_0 \tilde{\mathbf{\Lambda}} \tilde{\mathbf{W}}_0^T - \mathbf{E}$, so we get

$$\begin{aligned} & \|\mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T \mathbf{L} - \mathbf{L}\|_2 \\ &= \|(\mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T + \tilde{\mathbf{W}}_0 \tilde{\mathbf{\Lambda}} \tilde{\mathbf{W}}_0^T - \mathbf{E}) \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T (\mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T + \tilde{\mathbf{W}}_0 \tilde{\mathbf{\Lambda}} \tilde{\mathbf{W}}_0^T - \mathbf{E}) - \mathbf{L}\|_2 \\ &= \|\mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T - \mathbf{L} - \mathbf{W}_0 \mathbf{W}_0^T \mathbf{E} - \mathbf{E} \mathbf{W}_0 \mathbf{W}_0^T - \mathbf{E} \mathbf{W}_0 \mathbf{\Lambda}^{-1} \mathbf{W}_0^T \mathbf{E}\|_2 \\ &\leq \|\mathbf{E} - \tilde{\mathbf{W}}_0 \tilde{\mathbf{\Lambda}} \tilde{\mathbf{W}}_0^T\|_2 + 2\|\mathbf{E}\|_2 + \frac{\|\mathbf{E}\|_2^2}{|\lambda_r|} \\ &\leq \|\tilde{\mathbf{W}}_0 \tilde{\mathbf{\Lambda}} \tilde{\mathbf{W}}_0^T\|_2 + 4\|\mathbf{E}\|_2 \\ &\leq |\lambda_{r+1}| + 4\|\mathbf{E}\|_2 \\ &\leq 5\|\mathbf{E}\|_2, \end{aligned}$$

where the first and fourth inequality follow from (E.4) and (E.6), and $|\lambda_{r+1}| \leq \|\mathbf{E}\|_2$ since $\sigma_{r+1}(\mathbf{L}) = 0$. Together, we have

$$(E.8) \quad \mathbf{Y}_0 \leq 5\sqrt{\frac{\mu r}{n}} \|\mathbf{E}\|_2 \leq 5\alpha\sqrt{\mu r n} \|\mathbf{E}\|_\infty,$$

where the last inequality follows from Lemma E.1.

Next, we will bound $\mathbf{Y}_{ab}$. Note that

$$\begin{aligned} \mathbf{Y}_{ab} &= \max_i \left\| (\mathbf{E}^a \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{L} \mathbf{E}^b)^T \mathbf{e}_i \right\|_2 \\ &= \max_i \left\| \mathbf{E}^b \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{L} (\mathbf{W} \mathbf{W}^T \mathbf{E}^a \mathbf{e}_i) \right\|_2 \\ &\leq \max_i \left\| \mathbf{W}^T \mathbf{E}^a \mathbf{e}_i \right\|_2 \left\| \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{L} \right\|_2 \|\mathbf{E}\|_2^b \\ &\leq \sqrt{\frac{\mu r}{n}} (\alpha n \|\mathbf{E}\|_\infty)^{a+b} \left\| \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{L} \right\|_2 \\ &\leq \alpha\sqrt{\mu r n} \|\mathbf{E}\|_\infty \left(\frac{\sigma_r(\mathbf{L})}{8} \right)^{a+b-1} \left\| \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{L} \right\|_2, \end{aligned}$$

where the second inequality uses Lemmata E.1 and E.2. Moreover, by applying $\mathbf{L} = \mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T + \ddot{\mathbf{W}}_0 \ddot{\mathbf{\Lambda}} \ddot{\mathbf{W}}_0^T - \mathbf{E}$, we get

$$\begin{aligned} &\left\| \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{L} \right\|_2 \\ &= \left\| (\mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T + \ddot{\mathbf{W}}_0 \ddot{\mathbf{\Lambda}} \ddot{\mathbf{W}}_0^T - \mathbf{E}) \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T (\mathbf{W}_0 \mathbf{\Lambda} \mathbf{W}_0^T + \ddot{\mathbf{W}}_0 \ddot{\mathbf{\Lambda}} \ddot{\mathbf{W}}_0^T - \mathbf{E}) \right\|_2 \\ &= \left\| \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b-1)} \mathbf{W}_0^T - \mathbf{E} \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b)} \mathbf{W}_0^T - \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b)} \mathbf{W}_0^T \mathbf{E} \right. \\ &\quad \left. + \mathbf{E} \mathbf{L} \mathbf{W}_0 \mathbf{\Lambda}^{-(a+b+1)} \mathbf{W}_0^T \mathbf{E} \right\|_2 \\ &\leq |\lambda_r|^{-(a+b-1)} + |\lambda_r|^{-(a+b)} \|\mathbf{E}\|_2 + |\lambda_r|^{-(a+b)} \|\mathbf{E}\|_2 + |\lambda_r|^{-(a+b+1)} \|\mathbf{E}\|_2^2 \\ &= |\lambda_r|^{-(a+b-1)} \left(1 + \frac{2\|\mathbf{E}\|_2}{|\lambda_r|} + \left(\frac{\|\mathbf{E}\|_2}{|\lambda_r|} \right)^2 \right) \\ &= |\lambda_r|^{-(a+b-1)} \left(1 + \frac{\|\mathbf{E}\|_2}{|\lambda_r|} \right)^2 \\ &\leq 2|\lambda_r|^{-(a+b-1)} \\ &\leq 2 \left(\frac{7}{8} \sigma_r(\mathbf{L}) \right)^{-(a+b-1)}, \end{aligned}$$

where the second inequality follows from (E.6) and the last inequality follows from (E.5). Hence,

$$\begin{aligned}
\sum_{a+b>0} Y_{ab} &\leq \sum_{a+b>0} 2\alpha\sqrt{\mu rn} \|E\|_{\infty} \left(\frac{\frac{1}{8}\sigma_r^L}{\frac{7}{8}\sigma_r^L}\right)^{a+b-1} \\
&\leq 2\alpha\sqrt{\mu rn} \|E\|_{\infty} \sum_{a+b>0} \left(\frac{1}{7}\right)^{a+b-1} \\
&\leq 2\alpha\sqrt{\mu rn} \|E\|_{\infty} \left(\frac{1}{1-\frac{1}{7}}\right)^2 \\
(E.9) \quad &\leq 3\alpha\sqrt{\mu rn} \|E\|_{\infty}.
\end{aligned}$$

Combining (E.8) and (E.9), we have the *row version* of the inequality:

$$\begin{aligned}
\max_i \|(L_0 - L)^T e_i\|_2 &\leq Y_0 + \sum_{a+b>0} Y_{ab} \\
&\leq 5\alpha\sqrt{\mu rn} \|E\|_{\infty} + 3\alpha\sqrt{\mu rn} \|E\|_{\infty} \\
&\leq \frac{32\alpha\mu^{1.5}r^{1.5}}{\sqrt{n}} \sigma_{\max}(L),
\end{aligned}$$

where the last step uses (E.2).

The *column version* of the inequality, i.e.,

$$\max_j \|(L_0 - L)e_j\|_2 \leq \frac{32\alpha\mu^{1.5}r^{1.5}}{\sqrt{n}} \sigma_{\max}(L),$$

can be proved similarly. This finishes the proof. ■

REFERENCES

- [1] A. ALDROUBI, K. HAMM, A. B. KOKU, AND A. SEKMEN, *CUR decompositions, similarity matrices, and subspace clustering*, Front. Appl. Math. Stat., 4 (2019), p. 65.
- [2] A. ALDROUBI, A. SEKMEN, A. B. KOKU, AND A. F. CAKMAK, *Similarity matrix framework for data from union of subspaces*, Appl. Comput. Harmon. Anal., 45 (2018), pp. 425–435.
- [3] H. AVRON AND C. BOUTSIDIS, *Faster subset selection for matrices and applications*, SIAM J. Matrix Anal. Appl., 34 (2013), pp. 1464–1499.
- [4] J. J. BARTHOLDI, III, *A good submatrix is hard to find*, Oper. Res. Lett., 1 (1982), pp. 190–193.
- [5] A. BHASKARA, A. ROSTAMIZADEH, J. ALTSCHULER, M. ZADIMOGHADDAM, T. FU, AND V. MIRROKNI, *Greedy column subset selection: New bounds and distributed algorithms*, in Proceedings of the 33rd International Conference on Machine Learning, 2016, pp. 2539–2548.
- [6] C. BOUTSIDIS, P. DRINEAS, AND M. MAGDON-ISMAIL, *Near-optimal column-based matrix reconstruction*, SIAM J. Comput., 43 (2014), pp. 687–717.
- [7] C. BOUTSIDIS AND D. P. WOODRUFF, *Optimal CUR matrix decompositions*, SIAM J. Comput., 46 (2017), pp. 543–589.
- [8] T. BOUWMANS, *Subspace learning for background modeling: A survey*, Recent Patents Comput. Sci., 2 (2009), pp. 223–234.
- [9] T. BOUWMANS, S. JAVED, H. ZHANG, Z. LIN, AND R. OTAZO, *On the applications of robust PCA in image and video processing*, Proc. IEEE, 106 (2018), pp. 1427–1457.
- [10] H. CAI, *Accelerating Truncated Singular-Value Decomposition: A Fast and Provable Method for Robust Principal Component Analysis*, PhD thesis, University of Iowa, 2018.

- [11] H. CAI, J.-F. CAI, T. WANG, AND G. YIN, *Accelerated structured alternating projections for robust spectrally sparse signal recovery*, IEEE Trans. Signal Process., 69 (2021), pp. 809–821.
- [12] H. CAI, J.-F. CAI, AND K. WEI, *Accelerated alternating projections for robust principal component analysis*, J. Mach. Learn. Res., 20 (2019), pp. 685–717.
- [13] H. CAI, K. HAMM, L. HUANG, J. LI, AND T. WANG, *Rapid robust principal component analysis: CUR accelerated inexact low rank estimation*, IEEE Signal Process. Lett., 28 (2021), pp. 116–120.
- [14] E. J. CANDÈS, X. LI, Y. MA, AND J. WRIGHT, *Robust principal component analysis?*, J. ACM, 58 (2011), pp. 1–37.
- [15] V. CHANDRASEKARAN, S. SANGHAVI, P. A. PARRILO, AND A. S. WILLSKY, *Rank-sparsity incoherence for matrix decomposition*, SIAM J. Optim., 21 (2011), pp. 572–596.
- [16] J. CHIU AND L. DEMANET, *Sublinear randomized algorithms for skeleton decompositions*, SIAM J. Matrix Anal. Appl., 34 (2013), pp. 1361–1383.
- [17] A. ÇIVRIL, *Column subset selection problem is UG-hard*, J. Comput. System Sci., 80 (2014), pp. 849–859.
- [18] M. DEREZIŃSKI AND M. K. WARMUTH, *Reverse iterative volume sampling for linear regression*, J. Mach. Learn. Res., 19 (2018), pp. 853–891.
- [19] A. DESHPANDE AND L. RADEMACHER, *Efficient volume sampling for row/column subset selection*, in Proceedings of the 51st Annual IEEE Symposium on Foundations of Computer Science, IEEE, 2010, pp. 329–338.
- [20] P. DRINEAS, R. KANNAN, AND M. W. MAHONEY, *Fast Monte Carlo algorithms for matrices III: Computing a compressed approximate matrix decomposition*, SIAM J. Comput., 36 (2006), pp. 184–206.
- [21] P. DRINEAS, M. MAGDON-ISMAIL, M. W. MAHONEY, AND D. P. WOODRUFF, *Fast approximation of matrix coherence and statistical leverage*, J. Mach. Learn. Res., 13 (2012), pp. 3475–3506.
- [22] P. DRINEAS, M. W. MAHONEY, AND S. MUTHUKRISHNAN, *Relative-error CUR matrix decompositions*, SIAM J. Matrix Anal. Appl., 30 (2008), pp. 844–881.
- [23] E. ELHAMIFAR AND R. VIDAL, *Sparse subspace clustering: Algorithm, theory, and applications*, IEEE Trans. Pattern Anal. Mach. Intell., 35 (2013), pp. 2765–2781.
- [24] A. GEORGHIADES, P. BELHUMEUR, AND D. KRIEGMAN, *From few to many: Illumination cone models for face recognition under variable lighting and pose*, IEEE Trans. Pattern Anal. Mach. Intelligence, 23 (2001), pp. 643–660.
- [25] S. A. GOREINOV, I. V. OSELEDETS, D. V. SAVOSTYANOV, E. E. TYRTYSHNIKOV, AND N. L. ZAMARASHKIN, *How to find a good submatrix*, in Matrix Methods: Theory, Algorithms and Applications: Dedicated to the Memory of Gene Golub, World Scientific, River Edge, NJ, 2010, pp. 247–256.
- [26] S. A. GOREINOV, E. E. TYRTYSHNIKOV, AND N. L. ZAMARASHKIN, *A theory of pseudoskeleton approximations*, Linear Algebra Appl., 261 (1997), pp. 1–21.
- [27] S. A. GOREINOV, N. L. ZAMARASHKIN, AND E. E. TYRTYSHNIKOV, *Pseudo-skeleton approximations*, Dokl. Akad. Nauk, 343 (1995), pp. 151–152.
- [28] S. A. GOREINOV, N. L. ZAMARASHKIN, AND E. E. TYRTYSHNIKOV, *Pseudo-skeleton approximations by matrices of maximal volume*, Math. Notes, 62 (1997), pp. 515–519.
- [29] D. GROSS AND V. NESME, *Note on Sampling Without Replacing from a Finite Collection of Matrices*, arXiv preprint, [arXiv:1001.2738](https://arxiv.org/abs/1001.2738) [CS.IT], 2010.
- [30] M. GU AND S. C. EISENSTAT, *Downdating the singular value decomposition*, SIAM J. Matrix Anal. Appl., 16 (1995), pp. 793–810.
- [31] B. D. HAEFFELE, C. YOU, AND R. VIDAL, *A critique of self-expressive deep subspace clustering*, International Conference on Learning Representations, 2021, <https://openreview.net/pdf?id=FOyUZ26emy>.
- [32] K. HAMM AND L. HUANG, *Perturbations of CUR decompositions*, SIAM J. Matrix Anal. Appl., 42 (2021), pp. 351–375.
- [33] K. HAMM AND L. HUANG, *Perspectives on CUR decompositions*, Appl. Comput. Harmon. Anal., 48 (2020), pp. 1088–1099.
- [34] W.-D. JANG, C. LEE, AND C.-S. KIM, *Primary object segmentation in videos via alternate convex optimization of foreground and background distributions*, in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2016, pp. 696–704.
- [35] X. LI AND Y. PANG, *Deterministic column-based matrix decomposition*, IEEE Trans. Knowledge and Data Engrg., 22 (2010), pp. 145–149.

- [36] G. LIU, Z. LIN, S. YAN, J. SUN, Y. YU, AND Y. MA, *Robust recovery of subspace structures by low-rank representation*, IEEE Trans. Pattern Anal. Mach. Intell., 35 (2012), pp. 171–184.
- [37] M. W. MAHONEY AND P. DRINEAS, *CUR matrix decompositions for improved data analysis*, Proc. Natl. Acad. Sci. USA, 106 (2009), pp. 697–702.
- [38] A. MIKHALEV AND I. V. OSELEDETS, *Rectangular maximum-volume submatrices and their applications*, Linear Algebra Appl., 538 (2018), pp. 187–211.
- [39] P. NETRAPALLI, U. NIRANJAN, S. SANGHAVI, A. ANANDKUMAR, AND P. JAIN, *Non-convex robust PCA*, Adv. Neural Inf. Process. Syst., 27 (2014), pp. 1107–1115.
- [40] I. OSELEDETS AND E. TYRTYSHNIKOV, *Tt-cross approximation for multidimensional arrays*, Linear Algebra Appl., 432 (2010), pp. 70–88.
- [41] A. OSINSKY, *Rectangular Maximum Volume and Projective Volume Search Algorithms*, arXiv preprint, [arXiv:1809.02334](https://arxiv.org/abs/1809.02334) [math.NA], 2018.
- [42] A. OSINSKY AND N. L. ZAMARASHKIN, *Pseudo-skeleton approximations with better accuracy estimates*, Linear Algebra Appl., 537 (2018), pp. 221–249.
- [43] Y. SHITOV, *Column subset selection is NP-complete*, Linear Algebra Appl., 610 (2021), pp. 52–58.
- [44] D. C. SORENSEN AND M. EMBREE, *A DEIM induced CUR factorization*, SIAM J. Sci. Comput., 38 (2016), pp. A1454–A1482.
- [45] G. W. STEWART, *On the perturbation of pseudo-inverses, projections and linear least squares problems*, SIAM Rev., 19 (1977), pp. 634–662.
- [46] S. WANG AND Z. ZHANG, *Improving CUR matrix decomposition and the Nystrom approximation via adaptive sampling*, J. Mach. Learn. Res., 14 (2013), pp. 2729–2769.
- [47] J. A. TROPP, *Column subset selection, matrix factorization, and eigenvalue optimization*, in Proceedings of the 20th Annual ACM-SIAM Symposium on Discrete Algorithms, SIAM, 2009, pp. 978–986.
- [48] J. A. TROPP, *Improved analysis of the subsampled randomized Hadamard transform*, Adv. Adapt. Data Anal., 3 (2011), pp. 115–126.
- [49] R. VIDAL, *Subspace clustering*, IEEE Signal Process Mag., 28 (2011), pp. 52–68.
- [50] S. VORONIN AND P.-G. MARTINSSON, *Efficient algorithms for CUR and interpolative matrix decompositions*, Adv. Comput. Math., 43 (2017), pp. 495–516.
- [51] J. WRIGHT, A. Y. YANG, A. GANESH, S. S. SASTRY, AND Y. MA, *Robust face recognition via sparse representation*, IEEE Trans. Pattern Anal. Mach. Intell. Trans. Inform., 31 (2008), pp. 210–227.
- [52] H. XU, C. CARAMANIS, AND S. SANGHAVI, *Robust PCA via outlier pursuit*, IEEE Trans. Inform. Theory, 58 (2012), pp. 3047–3064.
- [53] X. YI, D. PARK, Y. CHEN, AND C. CARAMANIS, *Fast algorithms for robust PCA via gradient descent*, in Proceedings of the 30th International Conference on Neural Information Processing Systems, 2016, pp. 4152–4160.