
Compressed Factorization: Fast and Accurate Low-Rank Factorization of Compressively-Sensed Data

Vatsal Sharan^{*1} Kai Sheng Tai^{*1} Peter Bailis¹ Gregory Valiant¹

Abstract

What learning algorithms can be run directly on compressively-sensed data? In this work, we consider the question of accurately and efficiently computing low-rank matrix or tensor factorizations given data compressed via random projections. We examine the approach of first performing factorization in the compressed domain, and then reconstructing the original high-dimensional factors from the recovered (compressed) factors. In both the matrix and tensor settings, we establish conditions under which this natural approach will provably recover the original factors. While it is well-known that random projections preserve a number of geometric properties of a dataset, our work can be viewed as showing that they can also preserve certain solutions of non-convex, NP-Hard problems like non-negative matrix factorization. We support these theoretical results with experiments on synthetic data and demonstrate the practical applicability of compressed factorization on real-world gene expression and EEG time series datasets.

1 Introduction

We consider the setting where we are given data that has been compressed via random projections. This setting frequently arises when data is acquired via compressive measurements (Donoho, 2006; Candès & Wakin, 2008), or when high-dimensional data is projected to lower dimension in order to reduce storage and bandwidth costs (Haupt et al., 2008; Abdulghani et al., 2012). In the former case, the use of compressive measurement enables higher throughput in signal acquisition, more compact sensors, and reduced data storage costs (Duarte et al., 2008; Candès & Wakin, 2008).

^{*}Equal contribution ¹Stanford University, USA. Correspondence to: Vatsal Sharan <vsharan@stanford.edu>, Kai Sheng Tai <kst@cs.stanford.edu>.

In the latter, the use of random projections underlies many sketching algorithms for stream processing and distributed data processing applications (Cormode et al., 2012).

Due to the computational benefits of working directly in the compressed domain, there has been significant interest in understanding which learning tasks can be performed on compressed data. For example, consider the problem of supervised learning on data that is acquired via compressive measurements. Calderbank et al. (2009) show that it is possible to learn a linear classifier directly on the compressively sensed data with small loss in accuracy, hence avoiding the computational cost of first performing sparse recovery for each input prior to classification. The problem of learning from compressed data has also been considered for several other learning tasks, such as linear discriminant analysis (Durrant & Kabán, 2010), PCA (Fowler, 2009; Zhou & Tao, 2011; Ha & Barber, 2015), and regression (Zhou et al., 2009; Maillard & Munos, 2009; Kabán, 2014).

Building off this line of work, we consider the problem of performing low-rank matrix and tensor factorizations directly on compressed data, with the goal of recovering the low-rank factors in the original, uncompressed domain. Our results are thus relevant to a variety of problems in this setting, including sparse PCA, nonnegative matrix factorization (NMF), and Candecomp/Parafac (CP) tensor decomposition. As is standard in compressive sensing, we assume prior knowledge that the underlying factors are *sparse*.

For clarity of exposition, we begin with the matrix factorization setting. Consider a high-dimensional data matrix $M \in \mathbb{R}^{n \times m}$ that has a rank- r factorization $M = WH$, where $W \in \mathbb{R}^{n \times r}$, $H \in \mathbb{R}^{r \times m}$, and W is sparse. We are given the compressed measurements $\tilde{M} = PM$ for a known *measurement matrix* $P \in \mathbb{R}^{d \times n}$, where $d < n$. Our goal is to approximately recover the original factors W and H given the compressed data $\tilde{M}$ as accurately and efficiently as possible. This setting of compressed data with sparse factors arises in a number of important practical domains. For example, gene expression levels in a collection of tissue samples can be clustered using NMF to reveal correlations between particular genes and tissue types (Gao & Church, 2005). Since gene expression levels in each tissue sample are typically sparse, compressive sensing can be used to

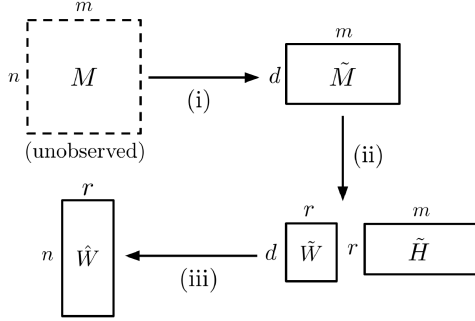


Figure 1: Schematic illustration of compressed matrix factorization. (i) The matrix $\tilde{M}$ is a compressed version of the full data matrix M . (ii) We directly factorize $\tilde{M}$ to obtain matrices $\tilde{W}$ and $\tilde{H}$. (iii) Finally, we approximate the left factor of M via sparse recovery on each column of $\tilde{W}$.

achieve more efficient measurement of the expression levels of large numbers of genes in each sample (Parvaresh et al., 2008). In this setting, each column of the $d \times m$ input matrix $\tilde{M}$ corresponds to the compressed measurements for the m tissue samples, while each column of the matrix W in the desired rank- r factorization corresponds to the pattern of gene expression in each of the r clusters.

We consider the natural approach of performing matrix factorization directly in the compressed domain (Fig. 1): first factorize the compressed matrix $\tilde{M}$ to obtain factors $\tilde{W}$ and $\tilde{H}$, and then approximately recover each column of W from the columns of $\tilde{W}$ using a sparse recovery algorithm that leverages the sparsity of the factors. We refer to this “compressed factorization” method as FACTORIZE-RECOVER. This approach has clear computational benefits over the alternative RECOVER-FACTORIZE method of first recovering the matrix M from the compressed measurements, and then performing low-rank factorization on the recovered matrix. In particular, FACTORIZE-RECOVER requires only r calls to the sparse recovery algorithm, in contrast to $m \gg r$ calls for the alternative. This difference is significant in practice, e.g. when m is the number of samples and r is a small constant. Furthermore, we demonstrate empirically that FACTORIZE-RECOVER also achieves better recovery error in practice on several real-world datasets.

Note that the FACTORIZE-RECOVER approach is guaranteed to work if the factorization of the compressed matrix $\tilde{M}$ yields the factors $\tilde{W} = PW$ and $\tilde{H} = H$, since we assume that the columns of W are sparse and hence can be recovered from the columns of $\tilde{W}$ using sparse recovery. Thus, the success of the FACTORIZE-RECOVER approach depends on finding this particular factorization of $\tilde{M}$. Since matrix factorizations are not unique in general, we ask: *under what conditions is it possible to recover the “correct” factorization $\tilde{M} = (PW)H$ of the compressed data, from which the original factors can be successfully recovered?*

Contributions. In this work, we establish conditions under which FACTORIZE-RECOVER provably succeeds, in both the matrix and tensor factorization domains. We complement our theoretical results with experimental validation that demonstrates both the accuracy of the recovered factors, as well as the computational speedup resulting from FACTORIZE-RECOVER versus the alternative approach of first recovering the data in the original uncompressed domain, and then factorizing the result.

Our main theoretical guarantee for sparse matrix factorizations, formally stated in Section 4.1, provides a simple condition under which the factors of the compressed data are the compressed factors. While the result is intuitive, the proof is delicate, and involves characterizing the likely sparsity of linear combinations of sparse vectors, exploiting graph theoretic properties of expander graphs. The crucial challenge in the proof is that the columns of W get mixed after projection, and we need to argue that they are still the sparsest vectors in any possible factorization after projection. This mixing of the entries, and the need to argue about the uniqueness of factorizations after projection, makes our setup significantly more involved than, for example, standard compressed sensing.

Theorem 1 (informal). *Consider a rank- r matrix $M \in \mathbb{R}^{n \times m}$, where $M = WH$, $W \in \mathbb{R}^{n \times r}$ and $H \in \mathbb{R}^{r \times m}$. Let the columns of W be sparse with the non-zero entries chosen at random. Given the compressed measurements $\tilde{M} = PM$ for a measurement matrix $P \in \mathbb{R}^{d \times n}$, under suitable conditions on P, n, m, d and the sparsity, $\tilde{M} = (PW)H$ is the sparsest rank- r factorization of $\tilde{M}$ with high probability, in which case performing sparse recovery on the columns of (PW) will yield the true factors W .*

While Theorem 1 provides guarantees on the quality of the sparsest rank- r factorization, it does not directly address the algorithmic question of how to find such a factorization efficiently. For some of the settings of interest, such as sparse PCA, efficient algorithms for recovering this sparsest factorization are known, under some mild assumptions on the data (Amini et al., 2009; Zhou & Tao, 2011; Deshpande & Montanari, 2014; Papailiopoulos et al., 2013). In such settings, Theorem 1 guarantees that we can efficiently recover the correct factorization.

For other matrix factorization problems such as NMF, the current algorithmic understanding of how to recover the factorization is incomplete even for uncompressed data, and guarantees for provable recovery require strong assumptions such as separability (Arora et al., 2012). As the original problem (computing NMF of the uncompressed matrix M) is itself NP-hard (Vavasis, 2009), hence one should not expect an analog of Theorem 1 to avoid solving a computationally hard problem and guarantee efficient recovery in general. In practice, however, NMF algorithms are com-

monly observed to yield sparse factorizations on real-world data (Lee & Seung, 1999; Hoyer, 2004) and there is substantial work on explicitly inducing sparsity via regularized NMF variants (Hoyer, 2004; Li et al., 2001; Kim & Park, 2008; Peharz & Pernkopf, 2012). In light of this empirically demonstrated ability to compute sparse NMF, Theorem 1 provides theoretical grounding for why FACTORIZE-RECOVER should yield accurate reconstructions of the original factors.

Our theoretical results assume a noiseless setting, but real-world data is usually noisy and only approximately sparse. Thus, we demonstrate the practical applicability of FACTORIZE-RECOVER through experiments on both synthetic benchmarks as well as several real-world gene expression datasets. We find that performing NMF on compressed data achieves reconstruction accuracy comparable to or better than factorizing the recovered (uncompressed) data at a fraction of the computation time.

In addition to our results on matrix factorization, we show the following analog to Theorem 1 for compressed CP tensor decomposition. The proof in this case follows in a relatively straightforward fashion from the techniques developed for our matrix factorization result.

Proposition 1 (informal). *Consider a rank- r tensor $T \in \mathbb{R}^{n \times m_1 \times m_2}$ with factorization $T = \sum_{i=1}^r A_i \otimes B_i \otimes C_i$, where A is sparse with the non-zero entries chosen at random. Under suitable conditions on P , the dimensions of the tensor, the projection dimension and the sparsity, $\tilde{T} = \sum_{i=1}^r (PA_i) \otimes B_i \otimes C_i$ is the unique factorization of the compressed tensor $\tilde{T}$ with high probability, in which case performing sparse recovery on the columns of (PA) will yield the true factors A .*

As in the case of sparse PCA, there is an efficient algorithm for finding this unique tensor factorization, as tensor decomposition can be computed efficiently when the factors are linearly independent (see e.g. Kolda & Bader (2009)). We empirically validate our approach for tensor decomposition on a real-world EEG dataset, demonstrating that factorizations from compressed measurements can yield interpretable factors that are indicative of the onset of seizures.

2 Related Work

There is an enormous body of algorithmic work on computing matrix and tensor decompositions more efficiently using random projections, usually by speeding up the linear algebraic routines that arise in the computation of these factorizations. This includes work on randomized SVD (Halko et al., 2011; Clarkson & Woodruff, 2013), NMF (Wang & Li, 2010; Tepper & Sapiro, 2016) and CP tensor decomposition (Battaglino et al., 2017). This work is rather different in spirit, as it leverages projections to accelerate certain

components of the algorithms, but still requires repeated accesses to the original uncompressed data. In contrast, our methods apply in the setting where we are only given access to the compressed data.

As mentioned in the introduction, learning from compressed data has been widely studied, yielding strong results for many learning tasks such as linear classification (Calderbank et al., 2009; Durrant & Kabán, 2010), multi-label prediction (Hsu et al., 2009) and regression (Zhou et al., 2009; Maillard & Munos, 2009). In most of these settings, the goal is to obtain a good predictive model in the compressed space itself, instead of recovering the model in the original space. A notable exception to this is previous work on performing PCA and matrix co-factorization on compressed data (Fowler, 2009; Ha & Barber, 2015; Yoo & Choi, 2011); we extend this line of work by considering sparse matrix decompositions like sparse PCA and NMF. To the best of our knowledge, ours is the first work to establish conditions under which sparse matrix factorizations can be recovered directly from compressed data.

Compressive sensing techniques have been extended to reconstruct higher-order signals from compressed data. For example, Kronecker compressed sensing (Duarte & Baraniuk, 2012) can be used to recover a tensor decomposition model known as Tucker decomposition from compressed data (Cai & Cichocki, 2013; 2015). Uniqueness results for reconstructing the tensor are also known in certain regimes (Sidiropoulos & Kyrillidis, 2012). Our work extends the class of models and measurement matrices for which uniqueness results are known and additionally provides algorithmic guarantees for efficient recovery under these conditions.

From a technical perspective, the most relevant work is Spielman et al. (2012), which considers the sparse coding problem. Although their setting differs from ours, the technical cores of both analyses involve characterizing the sparsity patterns of linear combinations of random sparse vectors.

3 Compressed Factorization

In this section, we first establish preliminaries on compressive sensing, followed by a description of the measurement matrices used to compress the input data. Then, we specify the algorithms for compressed matrix and tensor factorization that we study in the remainder of the paper.

Notation. Let $[n]$ denote the set $\{1, 2, \dots, n\}$. For any matrix A , we denote its i th column as A_i . For a matrix $P \in \mathbb{R}^{d \times n}$ such that $d < n$, define:

$$\mathcal{R}_P(w) = \underset{x: Px=w}{\operatorname{argmin}} \|x\|_1 \quad (1)$$

as the sparse recovery operator on $w \in \mathbb{R}^n$. We omit the subscript P when it is clear from context.

Background on Compressive Sensing. In the compress-

sive sensing framework, there is a sparse signal $x \in \mathbb{R}^n$ for which we are given $d \ll n$ linear measurements Px , where $P \in \mathbb{R}^{d \times n}$ is a known measurement matrix. The goal is to recover x using the measurements Px , given the prior knowledge that x is sparse. Seminal results in compressive sensing (Donoho, 2006; Candes & Tao, 2006; Candes, 2008) show that if the original solution is k -sparse, then it can be exactly recovered from $d = O(k \log n)$ measurements by solving a linear program (LP) of the form (1). More efficient recovery algorithms than the LP for solving the problem are also known (Berinde et al., 2008b; Indyk & Ruzic, 2008; Berinde & Indyk, 2009). However, these algorithms typically require more measurements in the compressed domain to achieve the same reconstruction accuracy as the LP formulation (Berinde & Indyk, 2009).

Measurement Matrices. In this work, we consider sparse, binary measurement (or projection) matrices $P \in \{0, 1\}^{d \times n}$ where each column of P has p non-zero entries chosen uniformly and independently at random. For our theoretical results, we set $p = O(\log n)$. Although the first results on compressive sensing only held for dense matrices (Donoho, 2006; Candes, 2008; Candes & Tao, 2006), subsequent work has shown that sparse, binary matrices can also be used for compressive sensing (Berinde et al., 2008a). In particular, Theorem 3 of Berinde et al. (2008a) shows that the recovery procedure in (1) succeeds with high probability for the class of P we consider if the original signal is k -sparse and $d = \Omega(k \log n)$. In practice, sparse binary projection matrices can arise due to physical limitations in sensor design (e.g., where measurements are sparse and can only be performed additively) or in applications of non-adaptive group testing (Indyk et al., 2010).

Low-Rank Matrix Factorization. We assume that each sample is an n -dimensional column vector in uncompressed form. Hence, the uncompressed matrix $M \in \mathbb{R}^{n \times m}$ has m columns corresponding to m samples, and we assume that it has some rank- r factorization: $M = WH$, where $W \in \mathbb{R}^{n \times r}$, $H \in \mathbb{R}^{r \times m}$, and the columns of W are k -sparse. We are given the compressed matrix $\tilde{M} = PM$ corresponding to the d -dimensional projection Pv for each sample $v \in \mathbb{R}^n$. We then compute a low-rank factorization using the following algorithm:

Algorithm 1 Compressed Matrix Factorization

Input: Compressed matrix $\tilde{M} = PM$, projection matrix P

Algorithm: Outputs estimates $(\hat{W}, \hat{H})$ of (W, H)

```

    Compute rank- $r$  factorization of  $\tilde{M}$  to obtain  $\tilde{W}, \tilde{H}$ 
    Set  $\hat{H} \leftarrow \tilde{H}$ 
    for  $1 \leq i \leq r$  do
        // Solve (1) to recover  $\hat{W}_i$  from  $\tilde{W}_i$ 
        Set  $\hat{W}_i \leftarrow \mathcal{R}(\tilde{W}_i)$ 
    end
```

CP Tensor Decomposition. As above, we assume that each sample is n -dimensional and k -sparse. The samples are now indexed by two coordinates $y \in [m_1]$ and $z \in [m_2]$, and hence can be represented by a tensor $T \in \mathbb{R}^{n \times m_1 \times m_2}$. We assume that T has some rank- r factorization $T = \sum_{i=1}^r A_i \otimes B_i \otimes C_i$, where the columns of A are k -sparse. Here $\otimes$ denotes the outer product: if $a \in \mathbb{R}^n, b \in \mathbb{R}^{m_1}, c \in \mathbb{R}^{m_2}$ then $a \otimes b \otimes c \in \mathbb{R}^{n \times m_1 \times m_2}$ and $(a \otimes b \otimes c)_{ijk} = a_i b_j c_k$. This model, CP decomposition, is the most commonly used model of tensor decomposition. For a measurement matrix $P \in \mathbb{R}^{d \times n}$, we are given a projected tensor $\tilde{T} \in \mathbb{R}^{d \times m_1 \times m_2}$ corresponding to a d dimensional projection Pv for each sample v . Algorithm 2 computes a low-rank factorization of T from $\tilde{T}$.

Algorithm 2 Compressed CP Tensor Decomposition

Input: Compressed tensor $\tilde{T}$, projection matrix P

Algorithm: Outputs estimates $(\hat{A}, \hat{B}, \hat{C})$ of (A, B, C)

```

    Compute rank- $r$  TD of  $\tilde{T}$ :  $\tilde{T} = \sum_{i=1}^r \tilde{A}_i \otimes \tilde{B}_i \otimes \tilde{C}_i$ 
    Set  $\hat{B} \leftarrow \tilde{B}, \hat{C} \leftarrow \tilde{C}$ 
    for  $1 \leq i \leq r$  do
        // Solve (1) to recover  $\hat{A}_i$  from  $\tilde{A}_i$ 
        Set  $\hat{A}_i \leftarrow \mathcal{R}(\tilde{A}_i)$ 
    end
```

We now describe our formal results for matrix and tensor factorization.

4 Theoretical Guarantees

In this section, we establish conditions under which FACTORIZE-RECOVER will provably succeed for matrix and tensor decomposition on compressed data.

4.1 Sparse Matrix Factorization

The main idea is to show that with high probability, $\tilde{M} = (PW)H$ is the *sparsest* factorization of $\tilde{M}$ in the following sense: for any other factorization $\tilde{M} = W'H'$, W' has strictly more non-zero entries than (PW) . It follows that the factorization $(PW)H$ is the optimal solution for a sparse matrix factorization of $\tilde{M}$ that penalizes non-zero entries of $\tilde{W}$. To show this uniqueness property, we show that the projection matrices satisfy certain structural conditions with high probability, namely that they correspond to adjacency matrices of *bipartite expander* graphs (Hoory et al., 2006), which we define shortly. We first formally state our theorem:

Theorem 1. Consider a rank- r matrix $M \in \mathbb{R}^{n \times m}$ which has factorization $M = WH$, for $H \in \mathbb{R}^{r \times m}$ and $W \in \mathbb{R}^{n \times r}$. Assume H has full row rank and $W = B \odot Y$, where $B \in \{0, 1\}^{n \times r}$, $Y \in \mathbb{R}^{n \times r}$ and $\odot$ denotes the elementwise product. Let each column of B have k non-zero entries chosen uniformly and independently at random, and each entry of Y be an independent random variable drawn from

any continuous distribution.¹ Assume $k > C$, where C is a fixed constant. Consider the projection matrix $P \in \{0, 1\}^{d \times n}$ where each column of P has $p = O(\log n)$ non-zero entries chosen independently and uniformly at random. Assume $d = \Omega((r + k) \log n)$. Let $\tilde{M} = PM$. Note that $\tilde{M}$ has one possible factorization $\tilde{M} = \tilde{W}H$ where $\tilde{W} = PW$. For some fixed $\beta > 0$, with failure probability at most $(r/n)e^{-\beta k} + (1/n^5)$, $\tilde{M} = \tilde{W}H$ is the sparsest possible factorization in terms of the left factors: for any other rank- r factorization $\tilde{M} = W'H'$, $\|\tilde{W}\|_0 < \|W'\|_0$.

Theorem 1 shows that if the columns of W are k -sparse, then projecting into $\Omega((r + k) \log n)$ dimensions preserves uniqueness, with failure probability at most $(r/n)e^{-\beta k} + (1/n^5)$, for some constant $\beta > 0$. As real-world matrices have been empirically observed to be typically close to low rank, the (r/n) term is usually small for practical applications. Note that the requirement for the projection dimension being at least $\Omega((r + k) \log n)$ is close to optimal, as even being able to uniquely recover a k -sparse n -dimensional vector x from its projection Px requires the projection dimension to be at least $\Omega(k \log n)$; we also cannot hope for uniqueness for projections to dimensions below the rank r . We also remark that the distributional assumptions on P and W are quite mild, as any continuous distribution suffices for the non-zero entries of W , and the condition on the set of non-zero coordinates for P and W being chosen uniformly and independently for each column can be replaced by a deterministic condition that P and W are adjacency matrices of bipartite expander graphs. We provide a proof sketch below, with the full proof deferred to the Appendix.

Proof sketch. We first show a simple Lemma that for any other factorization $\tilde{M} = W'H'$, the column space of W' and $\tilde{W}$ must be the same (Lemma 5 in the Appendix). Using this, for any other factorization $\tilde{M} = W'H'$, the columns of W' must lie in the column space of $\tilde{W}$, and hence our goal will be to prove that the columns of $\tilde{W}$ are the sparsest vectors in the column space of $\tilde{W}$, which implies that for any other factorization $\tilde{M} = W'H'$, $\|\tilde{W}\|_0 < \|W'\|_0$.

The outline of the proof is as follows. It is helpful to think of the matrix $\tilde{W} \in \mathbb{R}^{d \times r}$ as corresponding to the adjacency matrix of an unweighted bipartite graph G with r nodes on the left part U_1 and d nodes on the right part U_2 , and an edge from a node $u \in U_1$ to a node $v \in U_2$ if the corresponding entry of $\tilde{W}$ is non-zero. For any subset S of the columns of $\tilde{W}$, define $N(S)$ to be the subset of the rows of $\tilde{W}$ which have a non-zero entry in at least one of the columns in S . In the graph representation G , $N(S)$ is simply the neighborhood of a subset S of vertices in the left part U_1 . In part (a) we argue that if we take any subset S of the columns of $\tilde{W}$, $|N(S)|$ will be large. This implies

that taking a linear combination of all the S columns will result in a vector with a large number of non-zero entries—unless the non-zero entries cancel in many of the columns. In part (b), by using the properties of the projection matrix P and the fact that the non-zero entries of the original matrix W are drawn from a continuous distribution, we show this happens with zero probability.

The property of the projection matrix that is key to our proof is that it is the adjacency matrix of a bipartite expander graph, defined below.

Definition 1. Consider a bipartite graph R with n nodes on the left part and d nodes on the right part such that every node in the left part has degree p . We call R a $(\gamma n, \alpha)$ expander if every subset of at most $t \leq \gamma n$ nodes in the left part has at least αtp neighbors in the right part.

It is well-known that adjacency matrices of random bipartite graphs have good expansion properties under suitable conditions (Vadhan et al., 2012). For completeness, we show in Lemma 6 in the Appendix that a randomly chosen matrix P with p non-zero entries per column is the adjacency matrix of a $(\gamma n, 4/5)$ expander for $\gamma n = d/(pe^5)$ with failure probability $(1/n^5)$, if $p = O(\log n)$. Note that part (a) is a requirement on the graph G for the matrix W being a bipartite expander. In order to show that G is a bipartite expander, we show that with high probability P is a bipartite expander, and the matrix B corresponding to the non-zero entries of W is also a bipartite expander. G is a cascade of these bipartite expanders, and hence is also a bipartite expander.

For part (b), we need to deal with the fact that the entries of $\tilde{W}$ are no longer independent because the projection step leads to each entry of $\tilde{W}$ being the sum of multiple entries of W . However, the structure of P lets us control the dependencies, as each entry of W appears at most p times in $\tilde{W}$. Note that for a linear combination of any subset of S columns, $|N(S)|$ rows have non-zero entries in at least one of the S columns, and $|N(S)|$ is large by part (a). Since each entry of W appears at most p times in $\tilde{W}$, we can show that with high probability at most $|S|p$ out of the $|N(S)|$ rows with non-zero entries are zeroed out in any linear combination of the S columns. Therefore, if $|N(S)| - |S|p$ is large enough, then any linear combination of S columns has a large number of non-zero entries and is not sparse. This implies that the columns of $\tilde{W}$ are the sparsest columns in its column space. ■

A natural direction of future work is to relax some of the assumptions of Theorem 1, such as requiring independence between the entries of the B and Y matrices, and among the entries of the matrices themselves. It would also be interesting to show a similar uniqueness result under weaker, deterministic conditions on the left factor matrix W and the projection matrix P . Our result is a step in this direction

¹For example, a Gaussian distribution, or absolute value of Gaussian in the NMF setting.

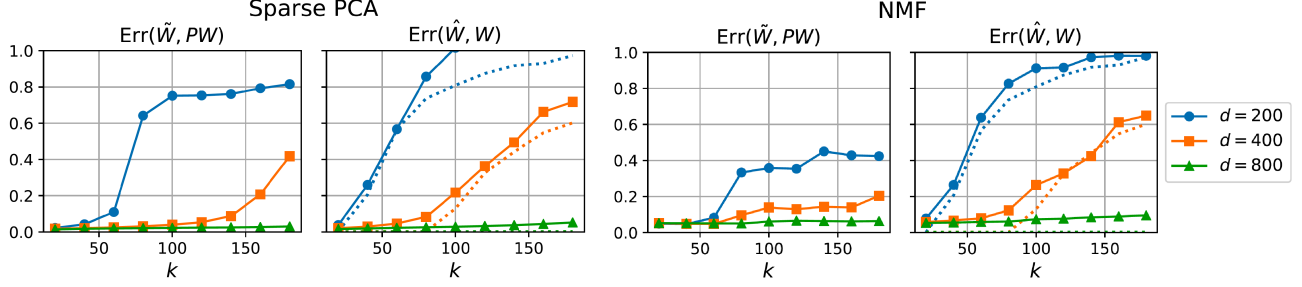


Figure 2: Approximation errors $\text{Err}(X, X_*) := \|X - X_*\|_F / \|X_*\|_F$ for sparse PCA and NMF on synthetic data with varying column sparsity k of W and projection dimension d . The values of d correspond to $10\times$, $5\times$, and $2.5\times$ compression respectively. $\text{Err}(\tilde{W}, PW)$ measures the distance between factors in the compressed domain: low error here is necessary for accurate sparse recovery. $\text{Err}(\hat{W}, W)$ measures the error after sparse recovery: the recovered factors $\hat{W}$ typically incur only slightly higher error than the oracle lower bound (dotted lines) where PW is known exactly.

and shows uniqueness if the non-zero entries of W and P are adjacency matrices of bipartite expanders, but it would be interesting to prove this under more relaxed assumptions.

4.2 Tensor Decomposition

It is easy to show uniqueness for tensor decomposition after random projection since tensor decomposition is unique under mild conditions on the factors (Kruskal, 1977; Kolda & Bader, 2009). Formally:

Proposition 1. Consider a rank- r tensor $T \in \mathbb{R}^{n \times m_1 \times m_2}$ which has factorization $T = \sum_{i=1}^r A_i \otimes B_i \otimes C_i$, for $A \in \mathbb{R}^{n \times r}$, $B \in \mathbb{R}^{m_1 \times r}$ and $C \in \mathbb{R}^{m_2 \times r}$. Assume B and C have full column rank and $A = X \odot Y$, where each column of X has exactly k non-zero entries chosen uniformly and independently at random, and each entry of Y is an independent random variable drawn from any continuous distribution. Assume $k > C$, where C is a fixed constant. Consider a projection matrix $P \in \{0, 1\}^{d \times n}$ with $d = \Omega((r + k) \log n)$ where each column of P has exactly $p = O(\log n)$ non-zero entries chosen independently and uniformly at random. Let $\tilde{T}$ be the projection of T obtained by projecting the first dimension. Note that $\tilde{T}$ has one possible factorization $\tilde{T} = \sum_{i=1}^r (PA_i) \otimes B_i \otimes C_i$. For a fixed $\beta > 0$, with failure probability at most $(r/n)e^{-\beta k} + (1/n^5)$, PA has full column rank, and hence this is a unique factorization of $\tilde{T}$.

Note that efficient algorithms are known for recovering tensors with linearly independent factors (Kolda & Bader, 2009) and hence under the conditions of Proposition 1 we can efficiently find the factorization in the compressed domain from which the original factors can be recovered. In the Appendix, we also show that we can provably recover factorizations in the compressed space using variants of the popular alternating least squares algorithm for tensor decomposition, though these algorithms require stronger assumptions on the tensor such as incoherence.

The proof of Proposition 1 is direct given the results es-

tablished in Theorem 1. We use the fact that tensors have a unique decomposition whenever the underlying factors $(PA), B, C$ are full column rank (Kruskal, 1977). By our assumption, B and C are given to be full rank. The key step is that by the proof of Theorem 1, the columns of PA are the sparsest columns in their column space. Therefore, they must be linearly independent, as otherwise the all zero vector will lie in their column space. Therefore, (PA) has full column rank, and Proposition 1 follows.

5 Experiments

We support our theoretical uniqueness results with experiments on real and synthetic data. On synthetically generated matrices where the ground-truth factorizations are known, we show that standard algorithms for computing sparse PCA and NMF converge to the desired solutions in the compressed space (§5.1). We then demonstrate the practical applicability of compressed factorization with experiments on gene expression data (§5.2) and EEG time series (§5.3).

5.1 Synthetic Data

We provide empirical evidence that standard algorithms for sparse PCA and NMF converge in practice to the desired sparse factorization $\tilde{M} = (PW)H$ —in order to achieve accurate sparse recovery in the subsequent step, it is necessary that the compressed factor $\tilde{W}$ be a good approximation of PW . For sparse PCA, we use alternating minimization with LARS (Zou et al., 2006), and for NMF, we use projected gradient descent (Lin, 2007).² Additionally, we evaluate the quality of the factors obtained after sparse recovery by measuring the approximation error of the recovered factors $\hat{W}$ relative to the true factors W .

We generate synthetic data following the conditions of The-

²For sparse PCA, we report results for the setting of the ℓ_1 regularization parameter that yielded the lowest approximation error. We did not use an ℓ_1 penalty for NMF. We give additional details in the Appendix.

Table 1: Summary of DNA microarray gene expression datasets, along with runtime (seconds) for each stage of the NMF pipeline on compressed data. FACTORIZE-RECOVER runs only r instances of sparse recovery, as opposed to the m instances used by the alternative, RECOVER-FACTORIZE.

Dataset	# Samples	# Features	RECOVER-FAC.		FAC.-RECOVER	
			Recovery	NMF	NMF	Recovery
CNS tumors	266	7,129	76.1	2.7	0.6	5.4
Lung carcinomas	203	12,600	78.8	4.0	0.8	9.3
Leukemia	435	54,675	878.4	39.6	6.9	55.0

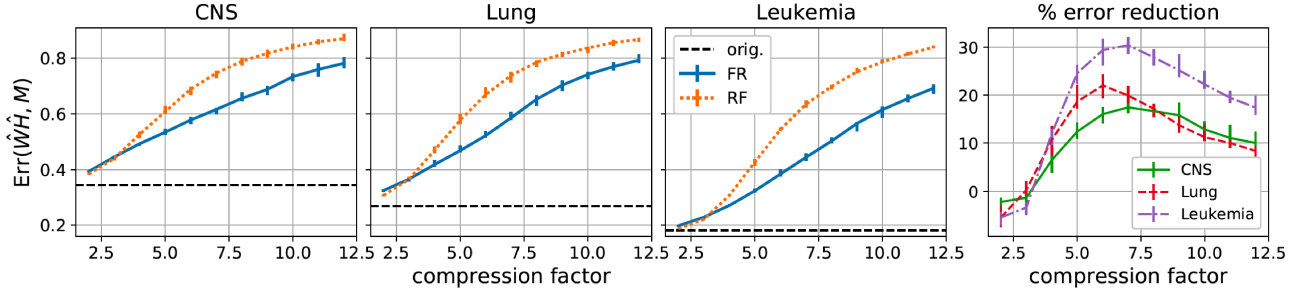


Figure 3: Normalized reconstruction errors $\|\hat{W}\hat{H} - M\|_F / \|M\|_F$ for NMF on gene expression data with varying compression factors n/d . **FR** (blue, solid) is FACTORIZE-RECOVER, **RF** (orange, dotted) is RECOVER-FACTORIZE. The horizontal dashed line is the error when M is decomposed in the original space. Perhaps surprisingly, when $n/d > 3$, we observe a reduction in reconstruction error when compressed data is first factorized. See the text for further discussion.

orem 1. For sparse PCA, we sample matrices $W = B \odot Y$ and H , where each column of $B \in \{0, 1\}^{n \times r}$ has k non-zero entries chosen uniformly at random, $Y_{ij} \stackrel{\text{iid}}{\sim} N(0, 1)$, and $H_{ij} \stackrel{\text{iid}}{\sim} N(0, 1)$. For NMF, an elementwise absolute value function is applied to the values sampled from this distribution. The noisy data matrix is $M = WH + \mathcal{E}$, where the noise term $\mathcal{E}$ is a dense random Gaussian matrix scaled such that $\|\mathcal{E}\|_F / \|WH\|_F = 0.1$. We observe $\tilde{M} = PM$, where P has $p = 5$ non-zero entries per column (in the Appendix, we study the effect of varying p on the error).

Figure 2 shows our results on synthetic data with $m = 2000$, $n = 2000$, and $r = 10$. For small column sparsities k relative to the projection dimension d , the estimated compressed left factors $\tilde{W}$ are good approximations to the desired solutions PW . Encouragingly, we find that the recovered solutions $\hat{W} = \mathcal{R}(\tilde{W})$ are typically only slightly worse in approximation error than $\mathcal{R}(PW)$, the solution recovered when the projection of W is known exactly. Thus, we perform almost as well as the idealized setting where we are given the correct factorization $(PW)H$.

5.2 NMF on Gene Expression Data

NMF is a commonly-used method for clustering gene expression data, yielding interpretable factors in practice (Gao & Church, 2005; Kim & Park, 2007). In the same domain, compressive sensing techniques have emerged as a promising approach for efficiently measuring the (sparse) expression levels of thousands of genes using compact measure-

ment devices (Parvaresh et al., 2008; Dai et al., 2008; Cleary et al., 2017).³ We evaluated our proposed NMF approach on gene expression datasets targeting three disease classes: embryonal central nervous system tumors (Pomeroy et al., 2002), lung carcinomas (Bhattacharjee et al., 2001), and leukemia (Mills et al., 2009) (Table 1). Each dataset is represented as a real-valued matrix where the i th row denotes expression levels for the i th gene across each sample.

Experimental Setup. For all datasets, we fixed rank $r = 10$ following previous clustering analyses in this domain (Gao & Church, 2005; Kim & Park, 2007). For each data matrix $M \in \mathbb{R}^{n \times m}$, we simulated compressed measurements $\tilde{M} \in \mathbb{R}^{d \times m}$ by projecting the feature dimension: $\tilde{M} = PM$. We ran projected gradient descent (Lin, 2007) for 250 iterations, which was sufficient for convergence.

Computation Time. Computation time for NMF on all 3 datasets (Table 1) is dominated by the cost of solving instances of the LP (1). As a result, FACTORIZE-RECOVER achieves much lower runtime as it requires a factor of m/r fewer calls to the sparse recovery procedure. While fast iterative recovery procedures such as SSMP (Berinde & Indyk, 2009) achieve faster recovery times, we found that they require approximately $2\times$ the number of measurements to achieve comparable accuracy to LP-based sparse recovery.

³The measurement matrices for these devices can be modeled as sparse binary matrices since each dimension of the acquired signal corresponds to the measurement of a small set of gene expression levels.

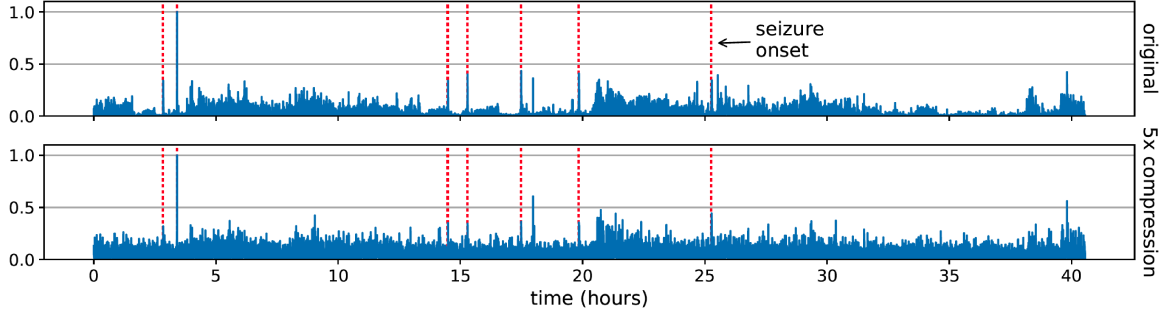


Figure 4: Visualization of a factor from the tensor decomposition of EEG data that correlates with the onset of seizures in a patient (red dotted lines). The factor recovered from a $5\times$ compressed version of the tensor (bottom) retains the peaks that are indicative of seizures.

Reconstruction Error. For a fixed number of measurements d , we observe that the FACTORIZE-RECOVER procedure achieves lower approximation error than the alternative method of recovering prior to factorizing (Figure 3). While this phenomenon is perhaps counter-intuitive, it can be understood as a consequence of the sparsifying effect of NMF. Recall that for NMF, we model each column of the compressed data $\tilde{M}$ as a nonnegative linear combination of the columns of $\tilde{W}$. Due to the nonnegativity constraint on the entries of $\tilde{W}$, we expect the average sparsity of the columns of $\tilde{W}$ to be at least that of the columns of $\tilde{M}$. Therefore, if $\tilde{W}$ is a good approximation of PW , we should expect that the sparse recovery algorithm will recover the columns of W at least as accurately as the columns of M , given a fixed number of measurements.

5.3 Tensor Decomposition on EEG Time Series Data

EEG readings are typically organized as a collection of time series, where each series (or channel) is a measurement of electrical activity in a region of the brain. Order-3 tensors can be derived from this data by computing short-time Fourier transforms (STFTs) for each channel, yielding a tensor where each slice is a time-frequency matrix. We experimented with tensor decomposition on a compressed tensor derived from the CHB-MIT Scalp EEG Database (Shoeb & Guttag, 2010). In the original space, this tensor has dimensions $27804 \times 303 \times 23$ (time $\times$ frequency $\times$ channel), corresponding to 40 hours of data (see the Appendix for further preprocessing details). The tensor was randomly projected along the temporal axis. We then computed a rank-10 non-negative CP decomposition of this tensor using projected Orth-ALS (Sharan & Valiant, 2017).

Reconstruction Error. At projection dimension $d = 1000$, we find that FACTORIZE-RECOVER achieves comparable error to RECOVER-FACTORIZE (normalized Frobenius error of 0.83 vs. 0.82). However, RF is three orders of magnitude slower than FR on this task due to the large number of sparse recovery invocations required (once for each frequency bin/channel pair, or $303 \times 23 = 6969$).

Factor Interpretability. The EEG time series data was recorded from patients suffering from epileptic seizures (Shoeb & Guttag, 2010). We found that the tensor decomposition yields a factor that correlates with the onset of seizures (Figure 4). At $5\times$ compression, the recovered factor qualitatively retains the interpretability of the factor obtained by decomposing the tensor in the original space.

6 Discussion and Conclusion

We briefly discuss our theoretical results on the uniqueness of sparse matrix factorizations in the context of dimensionality reduction via random projections. Such projections are known to preserve geometric properties such as pairwise distances (Kane & Nelson, 2014) and even singular vectors and singular values (Halko et al., 2011). Here, we showed that maximally sparse solutions to certain factorization problems are preserved by sparse binary random projections. Therefore, our results indicate that random projections can also, in a sense, preserve certain solutions of non-convex, NP-Hard problems like NMF (Vavasis, 2009).

To conclude, in this work we analyzed low-rank matrix and tensor decomposition on compressed data. Our main theoretical contribution is a novel uniqueness result for the matrix factorization case that relates sparse solutions in the original and compressed domains. We provided empirical evidence on real and synthetic data that accurate recovery can be achieved in practice. More generally, our results in this setting can be interpreted as the unsupervised analogue to previous work on supervised learning on compressed data. A promising direction for future work in this space is to examine other unsupervised learning tasks which can directly be performed in the compressed domain by leveraging sparsity.

Acknowledgments

We thank our anonymous reviewers for their valuable feedback on earlier versions of this manuscript. This research was supported in part by affiliate members and other support-

ers of the Stanford DAWN project—Ant Financial, Facebook, Google, Intel, Microsoft, NEC, SAP, Teradata, and VMware—as well as Toyota Research Institute, Keysight Technologies, Northrop Grumman, Hitachi, NSF awards AF-1813049 and CCF-1704417, an ONR Young Investigator Award N00014-18-1-2295, and Department of Energy award DE-SC0019205.

References

- Abdulghani, A. M., Casson, A. J., and Rodriguez-Villegas, E. Compressive sensing scalp EEG signals: implementations and practical performance. *Medical & Biological Engineering & Computing*, 2012.
- Amini, M., Usunier, N., and Goutte, C. Learning from multiple partially observed views—an application to multilingual text categorization. In *Advances in Neural Information Processing Systems*, 2009.
- Arora, S., Ge, R., and Moitra, A. Learning topic models—going beyond SVD. In *Foundations of Computer Science (FOCS), 2012 IEEE 53rd Annual Symposium on*, pp. 1–10. IEEE, 2012.
- Battaglino, C., Ballard, G., and Kolda, T. G. A practical randomized CP tensor decomposition. *arXiv preprint arXiv:1701.06600*, 2017.
- Berinde, R. and Indyk, P. Sequential sparse matching pursuit. In *Allerton Conference on Communication, Control, and Computing*, 2009.
- Berinde, R., Gilbert, A. C., Indyk, P., Karloff, H., and Strauss, M. J. Combining geometry and combinatorics: A unified approach to sparse signal recovery. In *Allerton Conference on Communication, Control, and Computing*, 2008a.
- Berinde, R., Indyk, P., and Ruzic, M. Practical near-optimal sparse recovery in the l_1 norm. In *Allerton*, 2008b.
- Bhattacharjee, A., Richards, W. G., Staunton, J., Li, C., Monti, S., Vasa, P., Ladd, C., Beheshti, J., Bueno, R., Gillette, M., et al. Classification of human lung carcinomas by mRNA expression profiling reveals distinct adenocarcinoma subclasses. *Proceedings of the National Academy of Sciences*, 2001.
- Caiafa, C. F. and Cichocki, A. Multidimensional compressed sensing and their applications. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 2013.
- Caiafa, C. F. and Cichocki, A. Stable, robust, and super fast reconstruction of tensors using multi-way projections. *IEEE Transactions on Signal Processing*, 2015.
- Calderbank, R., Jafarpour, S., and Schapire, R. Compressed learning: Universal sparse dimensionality reduction and learning in the measurement domain. *preprint*, 2009.
- Candes, E. J. The restricted isometry property and its implications for compressed sensing. *Comptes Rendus Mathématique*, 2008.
- Candes, E. J. and Tao, T. Near-optimal signal recovery from random projections: Universal encoding strategies? *IEEE Transactions on Information Theory*, 2006.
- Candès, E. J. and Wakin, M. B. An introduction to compressive sampling. *IEEE Signal Processing Magazine*, 2008.
- Clarkson, K. L. and Woodruff, D. P. Low rank approximation and regression in input sparsity time. In *Proceedings of the forty-fifth annual ACM symposium on Theory of computing*. ACM, 2013.
- Cleary, B., Cong, L., Cheung, A., Lander, E. S., and Regev, A. Efficient Generation of Transcriptomic Profiles by Random Composite Measurements. *Cell*, 2017.
- Cormode, G., Garofalakis, M., Haas, P. J., and Jermaine, C. Synopses for massive data: Samples, histograms, wavelets, sketches. *Foundations and Trends in Databases*, 4(1–3):1–294, 2012.
- Dai, W., Sheikh, M. A., Milenkovic, O., and Baraniuk, R. G. Compressive sensing DNA microarrays. *EURASIP journal on bioinformatics and systems biology*, 2008.
- Deshpande, Y. and Montanari, A. Sparse PCA via covariance thresholding. In *Advances in Neural Information Processing Systems*, pp. 334–342, 2014.
- Donoho, D. L. Compressed sensing. *IEEE Transactions on information theory*, 52(4):1289–1306, 2006.
- Duarte, M. F. and Baraniuk, R. G. Kronecker compressive sensing. *IEEE Transactions on Image Processing*, 2012.
- Duarte, M. F., Davenport, M. A., Takhar, D., Laska, J. N., Sun, T., Kelly, K. F., and Baraniuk, R. G. Single-pixel imaging via compressive sampling. *IEEE Signal Processing Magazine*, 25(2):83–91, 2008.
- Durrant, R. J. and Kabán, A. Compressed Fisher linear discriminant analysis: Classification of randomly projected data. In *KDD*, 2010.
- Fowler, J. E. Compressive-projection principal component analysis. *IEEE Transactions on Image Processing*, 18(10):2230–2242, 2009.
- Gao, Y. and Church, G. Improving molecular cancer class discovery through sparse non-negative matrix factorization. *Bioinformatics*, 2005.

- Ha, W. and Barber, R. F. Robust PCA with compressed data. In *Advances in Neural Information Processing Systems*, pp. 1936–1944, 2015.
- Halko, N., Martinsson, P.-G., and Tropp, J. A. Finding structure with randomness: Probabilistic algorithms for constructing approximate matrix decompositions. *SIAM review*, 53(2):217–288, 2011.
- Haupt, J., Bajwa, W. U., Rabbat, M., and Nowak, R. Compressed sensing for networked data. *IEEE Signal Processing Magazine*, 2008.
- Hoory, S., Linial, N., and Wigderson, A. Expander graphs and their applications. *Bulletin of the American Mathematical Society*, 2006.
- Hoyer, P. O. Non-negative matrix factorization with sparseness constraints. *JMLR*, 2004.
- Hsu, D. J., Kakade, S. M., Langford, J., and Zhang, T. Multi-label prediction via compressed sensing. In *Advances in neural information processing systems*, pp. 772–780, 2009.
- Indyk, P. and Ruzic, M. Near-optimal sparse recovery in the l_1 norm. In *FOCS*, 2008.
- Indyk, P., Ngo, H. Q., and Rudra, A. Efficiently decodable non-adaptive group testing. In *Proceedings of the twenty-first annual ACM-SIAM symposium on Discrete Algorithms*, 2010.
- Kabán, A. New bounds on compressive linear least squares regression. In *Artificial Intelligence and Statistics*, pp. 448–456, 2014.
- Kane, D. M. and Nelson, J. Sparsifier Johnson-Lindenstrauss transforms. *JACM*, 2014.
- Kim, H. and Park, H. Sparse non-negative matrix factorizations via alternating non-negativity-constrained least squares for microarray data analysis. *Bioinformatics*, 2007.
- Kim, J. and Park, H. Sparse nonnegative matrix factorization for clustering. Technical report, Georgia Institute of Technology, 2008.
- Kolda, T. G. and Bader, B. W. Tensor decompositions and applications. *SIAM Review*, 2009.
- Kruskal, J. B. Three-way arrays: Rank and uniqueness of trilinear decompositions, with application to arithmetic complexity and statistics. *Linear Algebra and its Applications*, 1977.
- Lee, D. D. and Seung, H. S. Learning the parts of objects by non-negative matrix factorization. *Nature*, 1999.
- Li, S. Z., Hou, X. W., Zhang, H. J., and Cheng, Q. S. Learning spatially localized, parts-based representation. In *Computer Vision and Pattern Recognition, 2001. CVPR 2001. Proceedings of the 2001 IEEE Computer Society Conference on*, volume 1, pp. I–I. IEEE, 2001.
- Lin, C.-J. Projected gradient methods for nonnegative matrix factorization. *Neural Computation*, 2007.
- Maillard, O. and Munos, R. Compressed least-squares regression. In *Advances in Neural Information Processing Systems*, pp. 1213–1221, 2009.
- Mills, K. I., Kohlmann, A., Williams, P. M., Wieczorek, L., Liu, W.-m., Li, R., Wei, W., Bowen, D. T., Loeffler, H., Hernandez, J. M., et al. Microarray-based classifiers and prognosis models identify subgroups with distinct clinical outcomes and high risk of AML transformation of myelodysplastic syndrome. *Blood*, 2009.
- Papailiopoulos, D., Dimakis, A., and Korokythakis, S. Sparse pca through low-rank approximations. In *International Conference on Machine Learning*, pp. 747–755, 2013.
- Parvaresh, F., Vikalo, H., Misra, S., and Hassibi, B. Recovering sparse signals using sparse measurement matrices in compressed DNA microarrays. *IEEE Journal of Selected Topics in Signal Processing*, 2008.
- Peharz, R. and Pernkopf, F. Sparse nonnegative matrix factorization with l_0 -constraints. *Neurocomputing*, 80: 38–46, 2012.
- Pomeroy, S. L., Tamayo, P., Gaasenbeek, M., Sturla, L. M., Angelo, M., McLaughlin, M. E., Kim, J. Y., Goumnerova, L. C., Black, P. M., Lau, C., et al. Prediction of central nervous system embryonal tumour outcome based on gene expression. *Nature*, 2002.
- Sharan, V. and Valiant, G. Orthogonalized ALS: A theoretically principled tensor decomposition algorithm for practical use. *ICML*, 2017.
- Shoeb, A. H. and Guttig, J. V. Application of machine learning to epileptic seizure detection. In *ICML*, 2010.
- Sidiropoulos, N. D. and Kyriillidis, A. Multi-way compressed sensing for sparse low-rank tensors. *IEEE Signal Processing Letters*, 2012.
- Spielman, D. A., Wang, H., and Wright, J. Exact recovery of sparsely-used dictionaries. In *Conference on Learning Theory*, pp. 37–1, 2012.
- Tepper, M. and Sapiro, G. Compressed nonnegative matrix factorization is fast and accurate. *IEEE Transactions on Signal Processing*, 64(9):2269–2283, 2016.

- Vadhan, S. P. et al. Pseudorandomness. *Foundations and Trends® in Theoretical Computer Science*, 7(1–3):1–336, 2012.
- Vavasis, S. A. On the complexity of nonnegative matrix factorization. *SIAM Journal on Optimization*, 20(3):1364–1377, 2009.
- Wang, F. and Li, P. Efficient nonnegative matrix factorization with random projections. In *Proceedings of the 2010 SIAM International Conference on Data Mining*. SIAM, 2010.
- Yoo, J. and Choi, S. Matrix co-factorization on compressed sensing. In *IJCAI*, volume 22, pp. 1595, 2011.
- Zhou, S., Lafferty, J., and Wasserman, L. Compressed and privacy-sensitive sparse regression. *IEEE Transactions on Information Theory*, 55(2):846–866, 2009.
- Zhou, T. and Tao, D. Godec: Randomized low-rank & sparse matrix decomposition in noisy case. In *International conference on machine learning*. Omnipress, 2011.
- Zou, H., Hastie, T., and Tibshirani, R. Sparse principal component analysis. *Journal of computational and graphical statistics*, 15(2):265–286, 2006.