

FAST LOW RANK APPROXIMATIONS OF MATRICES AND TENSORS*

S. FRIEDLAND[†], V. MEHRMANN[‡], A. MIEDLAR[§], AND M. NKENGLA[†]

Abstract. In many applications such as data compression, imaging or genomic data analysis, it is important to approximate a given $m \times n$ matrix A by a matrix B of rank at most k which is much smaller than m and n . The best rank k approximation can be determined via the singular value decomposition which, however, has prohibitively high computational complexity and storage requirements for very large m and n .

We present an optimal least squares algorithm for computing a rank k approximation to an $m \times n$ matrix A by reading only a limited number of rows and columns of A . The algorithm has complexity $\mathcal{O}(k^2 \max(m, n))$ and allows to iteratively improve given rank k approximations by reading additional rows and columns of A . We also show how this approach can be extended to tensors and present numerical results.

Key words. Singular value decomposition, CUR decomposition, Rank k approximation, Least squares, Tucker decomposition.

AMS subject classifications. 15A18, 15A69, 65F15, 93E24.

1. Introduction. Let $A = [a_{i,j}] \in \mathbb{R}^{m \times n}$, i.e., A is a real valued $m \times n$ matrix, where m, n are very large, e.g., $m, n \approx 10^6$. We may think of A as a full matrix which describes a noisy image [17], micro-array [1], genetic marker or HAP SNP data [25]. In order to compress or denoise the matrix, often a low rank (say rank k) approximation $B = [b_{i,j}] \in \mathbb{R}^{m \times n}$ of A is determined and used instead of the original data. Such a decomposition can be written as

$$B = \mathbf{x}_1 \mathbf{y}_1^\top + \cdots + \mathbf{x}_k \mathbf{y}_k^\top, \quad \mathbf{x}_i \in \mathbb{R}^m, \mathbf{y}_i \in \mathbb{R}^n, i = 1, \dots, k.$$

It requires only an amount of storage given by $(m + n)k$ instead of mn floating point numbers for the full A . The best rank k approximation of A with respect to

*Received by the editors on December 28, 2010. Accepted for publication on August 21, 2011.
 Handling Editor: James Nagy.

[†]Department of Mathematics, Statistics and Computer Science, University of Illinois at Chicago, Chicago, Illinois 60607-7045, USA (friedlan@uic.edu, nkengla@msn.com). Partially supported by *Berlin Mathematical School*, Berlin, Germany.

[‡]Institut für Mathematik, TU Berlin, Str. des 17. Juni 136, D-10623 Berlin, FRG (mehrman@math.tu-berlin.de). Partially supported by *Deutsche Forschungsgemeinschaft*, through the DFG Research Center MATHEON *Mathematics for Key Technologies* in Berlin.

[§]Institut für Mathematik, TU Berlin, Str. des 17. Juni 136, D-10623 Berlin, FRG (miedlar@math.tu-berlin.de). Supported by *Berlin Mathematical School*, Berlin, Germany.

the Frobenius norm can be determined via the singular value decomposition (SVD) $A = W\Sigma V^T$, with orthogonal matrices $W \in \mathbb{R}^{m \times m}$, $V \in \mathbb{R}^{n \times n}$ and a diagonal matrix $\Sigma = \text{diag}(\sigma_1, \dots, \sigma_r, 0, \dots, 0) \in \mathbb{R}^{m \times n}$ with $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r > 0$. Here r is the rank of A and the best rank k approximation is given by $B = W\tilde{\Sigma}V^T$, where $\tilde{\Sigma}$ is obtained from Σ by replacing $\sigma_{k+1}, \dots, \sigma_r$ with zeros [12].

The standard procedures for computing the SVD have a computational complexity of $\mathcal{O}(\min(m, n)^2 \max(m, n))$ floating point operations (flops), see e.g., [2, 12]. This makes the SVD unsuitable for most large scale applications. If only a low rank approximation is of interest, one can use Arnoldi or Lanczos methods [12, 19] to compute the singular values and singular vectors associated with the k largest singular values. This requires a substantial number of matrix-vector multiplications with the matrix A and thus a complexity of at least $\mathcal{O}(mnk)$ which still may be too large. Due to its high demand in applications, the topic of deriving faster methods to compute rank k approximation has therefore recently received a lot of attention.

Let us briefly review some of the recent literature. In the context of data-sparse approximations for the numerical solution of large linear systems, *pseudo-skeleton approximation* algorithms were suggested [13, 14, 15]. These procedures use block Gaussian elimination with a suitable $k \times k$ block matrix to determine a rank k approximation via a so-called *CUR decomposition*. They have complexity $\mathcal{O}(k^2(m+n))$ once an appropriate block is chosen.

Other methods with complexity $\mathcal{O}(mnk)$ that are also based on *CUR decompositions* and *subspace sampling* were suggested in [26, 27] and later improved and analyzed in [3, 4]. An improvement of the sampling idea was suggested in [9], which uses an iterative algorithm that reads several rows of A at random, determines from these data a rank k approximation and then iteratively improves the approximation by reading additional rows. Each update gives a better rank k approximation to A and the algorithm is terminated if the approximation is not improved any further, or if an allocated number of computational work is exceeded. The complexity of this algorithm, however, is still $\mathcal{O}(mnk)$. The first approaches for sampling methods that go below this complexity were presented in [16, 27].

In this paper, we describe a method that combines several of these ideas to an iterative method that allows on the one hand adaptive improvement of current rank k approximations and on the other hand has small computational complexity. This method is extended to tensors and proofs for the 3-tensor and 4-tensor case are given. The concept has been generalized to n -tensors in [24]. Let us briefly introduce the main idea of our algorithm for matrices. Assume that we read at random p columns and q rows of A , indexed by the sets $J \subset \{1, \dots, n\}$, $I \subset \{1, \dots, m\}$. This information corresponds to $m \times p$ and $q \times n$ matrices C and R respectively. We look for a matrix $F = CUR \in \mathbb{R}^{m \times n}$, with $U \in \mathbb{R}^{p \times q}$ still to be determined. We determine U_{opt} as a

solution to the least square problem of minimizing $\sum_{(i,j) \in \mathcal{S}} (a_{i,j} - (CUR)_{i,j})^2$, i.e.,

$$U_{\text{opt}} = \arg \min_{U \in \mathbb{R}^{p \times q}} \sum_{(i,j) \in \mathcal{S}} (a_{i,j} - (CUR)_{i,j})^2. \quad (1.1)$$

Here $\mathcal{S} \subset \{1, \dots, m\} \times \{1, \dots, n\}$ combines the entries of A which are contained in C and R . Let $A_{I,J}$ be the submatrix of A formed by the rows in I and columns in J . We show that a U_{opt} with minimal Frobenius norm is given by $A_{I,J}^\dagger$, the Moore-Penrose inverse of $A_{I,J}$. See [12] for the definition and the properties of the Moore-Penrose inverse. Assume that $p = q$ and the matrix $A_{I,J}$ is invertible. Then $U_{\text{opt}} = A_{I,J}^{-1}$ is the unique solution to (1.1), and the rows and columns of A indexed by I, J , respectively, coincide with $CA_{I,J}^{-1}R$. If $A_{I,J}$ is not a square matrix, or if $A_{I,J}$ is ill-conditioned, then we modify the CUR approximation to be of the form $C\tilde{U}_{\text{opt}}R$, where $\tilde{U}_{\text{opt}}$ is the Moore-Penrose inverse of a best rank ℓ approximation of $A_{I,J}$. To ensure that $A_{I,J}$ was not poorly chosen, we will pick several choices of $A_{I,J}$ and we will use the one which has the maximal number of significant singular values and among those the one with the maximal product. This concept is an alternative to the approaches to use fiber-crosses [8] or a square submatrix $A_{I,J}$ of maximal absolute determinant [13, 14].

The paper is organized as follows. We present the mathematical formulation of our method in Section 2 and in Section 3 we sketch how to extend our ideas to tensors. In Section 4 we present some computational results obtained with our algorithm on real and synthetic data and we compare our results with the best rank k approximation and with a full least squares approximation.

2. Mathematical formulation. In this section, we describe the mathematical basis of our method.

Let $A = [a_{i,j}]_{i,j=1}^{m,n} \in \mathbb{R}^{m \times n}$ and let $\|A\|_F := (\sum_{i,j=1}^{m,n} a_{i,j}^2)^{1/2}$ denote the Frobenius norm of A . We use the notation $\langle m \rangle := \{1, \dots, m\}$, $\langle n \rangle := \{1, \dots, n\}$ and let

$$I = \{1 \leq \alpha_1 < \dots < \alpha_q \leq m\} \subset \langle m \rangle, \quad J = \{1 < \beta_1 < \dots < \beta_p \leq n\} \subset \langle n \rangle$$

be two nonempty sets of cardinality q and p , respectively. Using the indices in I, J , we consider the submatrices

$$\begin{aligned} A_{I,J} &= [a_{\alpha_k, \beta_l}]_{k,l=1}^{q,p} \in \mathbb{R}^{q \times p}, \\ R &= A_{I, \langle n \rangle} = [a_{\alpha_k, j}]_{k,j=1}^{q,n} \in \mathbb{R}^{q \times n}, \\ C &= A_{\langle m \rangle, J} = [a_{i, \beta_l}]_{i,l=1}^{m,p} \in \mathbb{R}^{m \times p}, \end{aligned} \quad (2.1)$$

with index set $\mathcal{S} := \langle m \rangle \times \langle n \rangle \setminus ((\langle m \rangle \setminus I) \times (\langle n \rangle \setminus J))$, of cardinality $\#\mathcal{S} = mp + qn - pq$. Thus, $C = A_{\langle m \rangle, J}$ and $R = A_{I, \langle n \rangle}$ are composed of the columns J and the rows I of A , respectively.

A commonly used approximation [1, 4, 5, 6] of A based on C, R is of the form CUR , for some $U \in \mathbb{R}^{p \times q}$. If we know all the entries of A then the optimal choice of U is given by U_b satisfying

$$\|A - CU_bR\|_F = \min_{U \in \mathbb{R}^{p \times q}} \|A - CUR\|_F,$$

and it is well known that $U_b = C^\dagger AR^\dagger$.

Let $\mathcal{C}_r(p, q) \subset \mathbb{R}^{p \times q}$ denote the set of all real $p \times q$ matrices of rank at most r . Then, more generally, for $r \leq \min(p, q)$ the optimal rank r choice for U is given by $U_{b,r}$ satisfying

$$\|A - CU_{b,r}R\|_F = \min_{U \in \mathcal{C}_r(p, q)} \|A - CUR\|_F.$$

An explicit formula for $U_{b,r}$ is given in [10]. For $r = \min(\text{rank } C, \text{rank } R) \leq \min(p, q)$, one can choose $U_b = U_{b,r}$. The computational complexity to compute U_b is $\mathcal{O}(mnp^2q^2)$ (or roughly $\mathcal{O}(mn)$ if $p \ll m$ and $q \ll n$), which is again prohibitively large if m and n are very large. Note that the iterative algorithm suggested in [9] to compute a rank r approximation of A is of order $\mathcal{O}(mnr)$.

We have seen that to compute the best approximation to A of the form $F = CUR$ using only the entries of A given by C and R , we have to determine U_{opt} by solving the minimization problem (1.1). The optimal solution can be determined by the least squares solution, via a reformulation as linear system

$$T\mathbf{u} = \mathbf{a}, \quad T \in \mathbb{R}^{(mp+qn-pq) \times pq}, \mathbf{u} \in \mathbb{R}^{pq}, \mathbf{a} \in \mathbb{R}^{mp+qn-pq}, \quad (2.2)$$

where $\mathbf{u}, \mathbf{a}$ are vectors, whose coordinates are the entries of U and those entries of A , which are either in C or R , respectively. Moreover T is a submatrix of the Kronecker product $A \otimes A^\top$. This least squares approach is nice for the analysis of the problem, but it usually cannot be used for an efficient computation.

Instead, in the following two theorems, we give an explicit solution that can even be computed fast. If $A_{I,J}$ is a square and invertible, then the overdetermined system (2.2) has a unique solution but even if $A_{I,J}$ is not square we can make the solution unique by requiring norm minimization in Frobenius norm.

THEOREM 2.1. *Let $A \in \mathbb{R}^{m \times n}$, and let $I \subset \langle m \rangle$, $J \subset \langle n \rangle$ have cardinality q and p , respectively. Let $C = A_{\langle m \rangle, J} \in \mathbb{R}^{m \times p}$, and $R = A_{I, \langle n \rangle} \in \mathbb{R}^{q \times n}$ be as in (2.1) and suppose that $A_{I,J}$ is invertible. Then the overdetermined system (2.2) has a unique solution $U = A_{I,J}^{-1}$, i.e., the rows in I and the columns in J of the matrix $CA_{I,J}^{-1}R$ are equal to the corresponding rows and columns of A , respectively.*

Proof. For any $I \subset \langle m \rangle$, $J \subset \langle n \rangle$, with $\#I = q$, $\#J = p$, and $U \in \mathbb{R}^{p \times q}$ we have the identity

$$(A_{\langle m \rangle, J} U A_{I, \langle n \rangle})_{I, J} = A_{I, J} U A_{I, J}.$$

Hence the part of the system (2.2) corresponding to $(CUR)_{I,J} = A_{I,J}$ reduces to the equation

$$A_{I,J} U A_{I,J} = A_{I,J}.$$

If $A_{I,J}$ is a square matrix and invertible, then the unique solution to this matrix equation is $U = A_{I,J}^{-1}$. Furthermore

$$\begin{aligned} (A_{\langle m \rangle, J} A_{I, J}^{-1} A_{I, \langle n \rangle})_{I, \langle n \rangle} &= A_{I, J} A_{I, J}^{-1} A_{I, \langle n \rangle} = A_{I, \langle n \rangle}, \\ (A_{\langle m \rangle, J} A_{I, J}^{-1} A_{I, \langle n \rangle})_{\langle m \rangle, J} &= A_{\langle m \rangle, J} A_{I, J}^{-1} A_{I, J} = A_{\langle m \rangle, J}. \quad \square \end{aligned}$$

This result can be extended to the general nonsquare case.

THEOREM 2.2. *Let $A \in \mathbb{R}^{m \times n}$, and let $I \subset \langle m \rangle$, $J \subset \langle n \rangle$ have cardinality q and p , respectively. Let $C = A_{\langle m \rangle, J} \in \mathbb{R}^{m \times p}$, and $R = A_{I, \langle n \rangle} \in \mathbb{R}^{q \times n}$ be as in (2.1). Then $U = A_{I, J}^\dagger$ is the minimal solution (in Frobenius norm) of (1.1).*

Proof. Using the SVD of $A_{I, J}$ we may assume w.l.o.g. that $A_{I, J} = \Sigma_r \oplus 0_{(q-r) \times (p-r)}$, where $\Sigma_r = \text{diag}(\sigma_1, \dots, \sigma_r)$ and $r = \text{rank } A_{I, J}$. For $U \in \mathbb{R}^{p \times q}$ denote by $U_r \in \mathbb{R}^{p \times q}$ the matrix obtained from U by replacing the last $p-r$ rows and $q-r$ columns by rows and columns of zeros, respectively. Note that then $CUR = CU_r R$ and $\|U_r\|_F \leq \|U\|_F$, and equality holds if and only if $U = U_r$. Hence, the minimal Frobenius norm least squares solution U of is given by $U = U_r$. Using the fact that the rows $r+1, \dots, q$ and columns $r+1, \dots, p$ of CUR are zero it follows that the minimum in (1.1) is reduced to the minimum on $\mathcal{S}' = \langle m \rangle \times \langle r \rangle \cup \langle r \rangle \times \langle n \rangle$. Then, by Theorem 2.1 the solution to the minimal Frobenius norm least squares problem is given by $\Sigma^\dagger$. $\square$

To make use of these ideas in finite precision arithmetic we have to work with the *numerical rank* rather than with the rank. For $A \in \mathbb{R}^{m \times n}$ the numerical rank $\text{rank}_{\text{num}} A$ of A is defined to be the number of singular values σ_j of A that are above a certain threshold, which is usually taken to be $\epsilon \cdot c(n, m) \cdot \sigma_1$, where ϵ is the machine precision, $c(n, m)$ is a low degree polynomial in the dimensions n, m and σ_1 is the largest singular value of A , see [2, 12] for details.

Assume that we have selected at random sets of q rows I and p columns J , and set $r_{\text{num}} := \text{rank}_{\text{num}} A_{I, J}$. If the rank of $A_{I, J}$ equals its numerical rank, then we can use Theorem 2.2 to choose $U_{\text{opt}} = A_{I, J}^\dagger$. If $\text{rank}_{\text{num}} < \text{rank } A_{I, J}$, then we choose as approximation

$$B = A_{\langle m \rangle, J} A_{I, J, r_{\text{num}}}^\dagger A_{J, \langle n \rangle}, \quad \tilde{U}_{\text{opt}} = A_{I, J, r_{\text{num}}}^\dagger, \quad (2.3)$$

where $A_{I, J, r_{\text{num}}}$ is the best rank r_{num} approximation of $A_{I, J}$.

We summarize the presented procedure in the following algorithm.

ALGORITHM 2.3. **FSVD** Let $A \in \mathbb{R}^{m \times n}$. Fix integers $p, t_{\max} \geq 1$ and set $q = p$.

1. Choose two sets $I \subset \langle m \rangle, J \subset \langle n \rangle$ of cardinality p at random. Determine the numerical rank $k(I, J) = r_{\text{num}} A_{I,J}$ and the product of the first $k(I, J)$ singular values of $A_{I,J}$, which is denoted by $\pi(I, J)$. Repeat this procedure $t_{\max}$ times and keep the rows and columns I, J with the maximal $k := k(I, J)$ and of those, keep that with maximal $\pi(I, J)$.
2. Compute the best rank k approximation of $A_{I,J}$ denoted by $A_{I,J,k}$.
3. Use $B = CA_{I,J,k}^\dagger R$ as rank k approximation of A .

Based on the information that we read, the $\mathcal{S}$ -average error (SAE) of our approximation is

$$\text{Error}_{\text{av}}(B) = \frac{\sum_{(i,j) \in \mathcal{S}} (a_{i,j} - b_{i,j})^2}{\sum_{(i,j) \in \mathcal{S}} a_{i,j}^2} \quad (2.4)$$

and typically this error is expected to be small. But if we are not satisfied with the average approximation error, then we can adaptively improve the approximation by adding more rows and columns to the matrices C, R , i.e., by increasing the index sets I, J and by computing a new $A_{I,J,k}^\dagger$ for the extended sets.

Since B and B_1 are of low rank, we can efficiently, (in complexity $\mathcal{O}(k^2(n+m))$), compute the average distance

$$\frac{\|B - B_1\|_F^2}{\|B\|_F \|B_1\|_F}. \quad (2.5)$$

This is done as follows. Write explicitly each matrix X in (2.5) as a sum of rank one matrices, and recall that $\|X\|_F^2 = \text{tr} XX^\top$, where tr denotes the trace. Finally note that $\text{tr} \mathbf{xy}^\top \mathbf{uv}^\top = (\mathbf{y}^\top \mathbf{u})(\mathbf{v}^\top \mathbf{x})$.

We adaptively increase the sets I, J until either the distance (2.5) has converged, or we have exceeded the allowed amount of computational work. We can increase sets I, J either incrementally by adding a fixed number of new columns and rows, or by certain percentage which, for example, may depend on the size of the achieved reduction in the $\mathcal{S}$ -average error (SAE). However, since the choice of particular increase strategy is problem-dependent, it is usually based on some heuristics.

In this iterative improvement, we employ updating the computation of $A_{I,J,k}^\dagger$. Since updating of the singular value decomposition is usually not possible, we can reduce the computational work (in particular for larger k) by replacing the singular value decomposition with the *ULV* decomposition [29] which can be updated when rows and columns are added.

Since the complexity of computing $A_{I,J,k}^\dagger$ is $\mathcal{O}(\min(p, q)^2 \max(p, q))$, it follows that for small $p, q \ll m, n$ the computational complexity of this method is dominated

by the reading of the data and the computation of the average error (2.4) and thus we obtain a complexity of $\mathcal{O}(k^2 \max(m, n))$.

After presenting the basic concept of our method for matrices, in the next section, we show how these ideas can be extended to tensors.

3. Extensions to tensors. Tensor decompositions are used in many applications to help explain interactions among multi-way data. These applications include chemometrics [11], psychometrics [18, 30], computer image and human motion recognition [31] and image restoration to name a few as well as many other areas using multi-way data analysis [7]. Collecting and storing large datasets of sensor data, social network data, fMRI medical data is easier than ever with commodity, terabyte disks. This data explosion creates deep research challenges that require scalable, tensor-based algorithms. The volume of a n -tensor is the product of the component dimensions $l_1 l_2 \cdots l_n$ and therein lies the curse of dimensionality. In many applications $N = l_1 l_2 \cdots l_n$ is big primarily because n is big. And n is getting bigger because researchers are interested in developing more sophisticated models that capture multiple interactions instead of the simplistic pairwise interactions. In higher dimensions, familiar linear algebra concepts such as rank, become complicated. Determining a closed-form solution for the rank of a general tensor is still an open problem. More pertinently, developing an analog to the Eckart-Young theorem applied in the SVD is an area that continues to garner widespread interest for multilinear tensors. The most common and widely used definition of a tensor rank refers to the minimum number of rank-1 tensors needed to sum to the tensor, which is a very hard problem. Instead, we use as tensor rank the n -rank of a tensor. The following definition involves the matrix representation of tensors, where $\mathcal{A}^{(n)}$ is the unfolding of the tensor $\mathcal{A}$ along its n^{th} mode (dimension).

DEFINITION 3.1. Let $\mathcal{A} \in \mathbb{R}^{l_1 \times \cdots \times l_d}$ and suppose that $\mathcal{A}^{(n)}$ represents the matrix derived by unfolding $\mathcal{A}$ in the n^{th} -mode ($n \in \langle d \rangle$), then the n -rank of $\mathcal{A}$ denoted by $\text{rank}_n(\mathcal{A})$, is

$$\text{rank}_n(\mathcal{A}) = \text{rank}(\mathcal{A}^{(n)}) \quad \text{for } n \in \langle d \rangle.$$

Thus, the $\text{rank}_n(\mathcal{A})$ is the rank of the matrix $\mathcal{A}^{(n)}$. As such, the 1-rank of a matrix (a tensor of dimension 2) is its column rank, while the 2-rank is its row rank. In this section, we discuss how our idea of computing least squares optimal CUR decompositions can be extended to 3 and 4 tensors, see also [21]. For a general n -tensor proof, see [24].

3.1. 3-tensors. Let $\mathcal{A} = [a_{i_1, i_2, i_3}] \in \mathbb{R}^{l_1 \times l_2 \times l_3}$ be a 3-tensor, where the dimensions l_1, l_2, l_3 , are large. For each $j = 1, 2, 3$, we read subtensors of $\mathcal{A}$ denoted by $\mathcal{C}_j = [c_{i_1, j, i_2, j, i_3, j}^{(j)}] \in \mathbb{R}^{l_{1,j} \times l_{2,j} \times l_{3,j}}$. We assume that $\mathcal{C}_j$ has the same number of coor-

ordinates as $\mathcal{A}$ in j -th direction, and a small number of coordinates in the other two directions. That is, $l_{j,j} = l_j$ and the other two indices $l_{s,j}$, $s \in \{1, 2, 3\} \setminus \{j\}$ are of order $\mathcal{O}(k)$, for $j = 1, 2, 3$. Therefore, $\mathcal{C}_j$ corresponds to the j -section of the tensor $\mathcal{A}$. The *small* dimensions of $\mathcal{C}_j$ are $(l_{s_j,j}, l_{t_j,j})$ where $\{s_j, t_j\} = \{1, 2, 3\} \setminus \{j\}$ for $j = 1, 2, 3$. Let $m_j := l_{s_j,j} l_{t_j,j}$ for $j = 1, 2, 3$.

To determine an approximation, we then look for a 6-tensor

$$\mathcal{V} = [v_{q_1, q_2, q_3, q_4, q_5, q_6}] \in \mathbb{R}^{l_{2,1} \times l_{3,1} \times l_{1,2} \times l_{3,2} \times l_{1,3} \times l_{2,3}}$$

and approximate the given tensor $\mathcal{A}$ by a tensor

$$\mathcal{B} = [b_{i_1, i_2, i_3}] := \mathcal{V} \cdot \mathcal{C}_1 \cdot \mathcal{C}_2 \cdot \mathcal{C}_3 \in \mathbb{R}^{\ell_1 \times \ell_2 \times \ell_3},$$

where we contract the 6 indices in $\mathcal{V}$ and the corresponding two indices $\{1, 2, 3\} \setminus \{j\}$ in $\mathcal{C}_j$ for $j = 1, 2, 3$, i.e., our approximation has the entries

$$b_{i_1, i_2, i_3} = \sum_{q_1=1}^{\ell_{2,1}} \sum_{q_2=1}^{\ell_{3,1}} \sum_{q_3=1}^{\ell_{1,2}} \sum_{q_4=1}^{\ell_{3,2}} \sum_{q_5=1}^{\ell_{1,3}} \sum_{q_6=1}^{\ell_{2,3}} v_{q_1, q_2, q_3, q_4, q_5, q_6} c_{i_1, q_1, q_2}^{(1)} c_{q_3, i_2, q_4}^{(2)} c_{q_5, q_6, i_3}^{(3)}.$$

This approximation is equivalent to a so-called *Tucker approximation* [30]. Indeed, if we represent each tensor $\mathcal{C}_j$ by a matrix $C_j \in \mathbb{R}^{m_j \times l_j}$ that has the same number of columns as the range of the j -th index of the tensor $\mathcal{A}$ and as number of rows the product of the ranges of the remaining two *small* indices of $\mathcal{C}_j$, i.e., $C_j = [c_{r, i_j}^{(j)}]_{r, i_j=1}^{m_j \cdot \ell_j}$. Then $c_{r, i_j}^{(j)}$ is equal to the corresponding entry $c_{i_1, i_2, i_3}^{(j)}$, where the value of r corresponds to the double index (i_s, i_t) for $\{s, t\} = \{1, 2, 3\} \setminus \{j\}$.

Now with $\mathcal{U} = [u_{j_1, j_2, j_3}] \in \mathbb{R}^{m_1 \times m_2 \times m_3}$, the equivalent Tucker representation of $\mathcal{B} = [b_{i_1, i_2, i_3}]$ is given by the entries

$$b_{i_1, i_2, i_3} = \sum_{j_1=1}^{m_1} \sum_{j_2=1}^{m_2} \sum_{j_3=1}^{m_3} u_{j_1, j_2, j_3} c_{j_1, i_1}^{(1)} c_{j_2, i_2}^{(2)} c_{j_3, i_3}^{(3)}, \quad (i_1, i_2, i_3) \in \langle \ell_1 \rangle \times \langle \ell_2 \rangle \times \langle \ell_3 \rangle.$$

This formula is expressed commonly as

$$\mathcal{B} = \mathcal{U} \times_1 C_1 \times_2 C_2 \times_3 C_3. \quad (3.1)$$

We now choose three subsets of the rows, columns and tubes of $\mathcal{A}$, with the goal of maximizing the n -rank of the resulting subtensor,

$$I \subset \langle \ell_1 \rangle, \#I = p, \quad J \subset \langle \ell_2 \rangle, \#J = q, \quad K \subset \langle \ell_3 \rangle, \#K = r.$$

Let

$$\begin{aligned} \mathcal{C}_1 &= \mathcal{A}_{\langle \ell_1 \rangle, J, K} := [a_{i, j, k}], \quad i \in \langle \ell_1 \rangle, j \in J, k \in K, \\ \mathcal{C}_2 &= \mathcal{A}_{I, \langle \ell_2 \rangle, K} := [a_{i, j, k}], \quad i \in I, j \in \langle \ell_2 \rangle, k \in K, \\ \mathcal{C}_3 &= \mathcal{A}_{I, J, \langle \ell_3 \rangle} := [a_{i, j, k}], \quad i \in I, j \in J, k \in \langle \ell_3 \rangle, \\ \mathcal{S} &= (\langle \ell_1 \rangle \times J \times K) \cup (I \times \langle \ell_2 \rangle \times K) \cup (I \times J \times \langle \ell_3 \rangle). \end{aligned} \quad (3.2)$$

Then we define $\mathcal{U}_{\text{opt}}$ similarly to (1.1) as solution of a least squares problem, i.e.,

$$\mathcal{U}_{\text{opt}} = \arg \min_{\mathcal{U} \in \mathbb{R}^{m_1 \times m_2 \times m_3}} \sum_{(i,j,k) \in \mathcal{S}} (a_{i,j,k} - (\mathcal{U} \times_1 C_1 \times_2 C_2 \times_3 C_3)_{i,j,k})^2. \quad (3.3)$$

A tensor $\mathcal{A}$ is called *generic* if its entries are not the set of zeros of any finite number of given polynomial equations in the entries of $\mathcal{A}$.

THEOREM 3.2. *For a given generic tensor $\mathcal{A} \in \mathbb{R}^{\ell_1 \times \ell_2 \times \ell_3}$ let $I \subset \langle \ell_1 \rangle$, $J \subset \langle \ell_2 \rangle$, and $K \subset \langle \ell_3 \rangle$ be three sets of integers with the particular cardinalities $\#I = p$, $\#J = p$, $\#K = p^2$ (chosen as such purely for convenience, for a given integer p). Let $\mathcal{C}_1 = \mathcal{A}_{\langle \ell_1 \rangle, J, K} \in \mathbb{R}^{\ell_1 \times p \times p^2}$, $\mathcal{C}_2 = \mathcal{A}_{I, \langle \ell_2 \rangle, K} \in \mathbb{R}^{p \times \ell_2 \times p^2}$, and $\mathcal{C}_3 = \mathcal{A}_{I, J, \langle \ell_3 \rangle} \in \mathbb{R}^{p \times p \times \langle \ell_3 \rangle}$ be the subtensors defined as in (3.2). Assume that $C_j \in \mathbb{R}^{m_j \times \ell_j}$ is the matrix induced by the tensor $\mathcal{C}_j$ for $j = 1, 2, 3$.*

Then the minimum given in (3.3) is zero, i.e., there exists $\mathcal{U} \in \mathbb{R}^{p^3 \times p^3 \times p^2}$ such that the tensor $\mathcal{B}$ given by (3.1) has the same entries as $\mathcal{A}$ for $(i, j, k) \in \mathcal{S}$.

Proof. We can represent the tensor $\mathcal{A} = [a_{i,j,k}]$ as a matrix $E = [e_{s,k}] \in \mathbb{R}^{(\ell_1 \cdot \ell_2) \times \ell_3}$. So $e_{s,k} = a_{i,j,k}$ for the corresponding pair of indices $(i, j) \in \langle \ell_1 \rangle \times \langle \ell_2 \rangle$. Then the set of indices $(i, j) \in I \times J$ corresponds to the set of indices $L \subset \langle \ell_1 \cdot \ell_2 \rangle$, where $\#L = p^2$. Since $\mathcal{A}$ is generic it follows that the submatrix $E_{L,K}$ that has row indices in L and column indices in K is invertible. Hence, the matrix $\tilde{E} = E_{\langle \ell_1 \cdot \ell_2 \rangle, K} E_{L,K}^{-1} E_{L, \langle \ell_3 \rangle}$ has the same entries as E in the places $\langle \ell_1 \cdot \ell_2 \rangle \times K \cup L \times \langle \ell_3 \rangle$. Equivalently, one can represent $\tilde{E}$ as

$$\mathcal{A}_{\langle \ell_1 \rangle, \langle \ell_2 \rangle, K} E_{L,K}^{-1} \mathcal{A}_{I, J, \langle \ell_3 \rangle}. \quad (3.4)$$

For each $k \in K$, consider the matrix

$$F_k := \mathcal{A}_{\langle \ell_1 \rangle, \langle \ell_2 \rangle, k} = [a_{i,j,k}]_{i,j=1}^{\ell_1, \ell_2} \in \mathbb{R}^{\ell_1 \times \ell_2}.$$

Since $\mathcal{A}$ is generic, $\det(F_k)_{I,J} \neq 0$ for each $k \in K$. Hence, F_k can be approximated by $G_k := (F_k)_{\langle \ell_1 \rangle, J} (F_k)_{I, J}^{-1} (F_k)_{I, \langle \ell_2 \rangle}$. Moreover, G_k has the same entries as F_k in the positions $\langle \ell_1 \rangle \times J \cup I \times \langle \ell_2 \rangle$.

Equivalently, we have that

$$\mathcal{A}_{\langle \ell_1 \rangle, J, k} \mathcal{A}_{I, J, k}^{-1} \mathcal{A}_{I, \langle \ell_2 \rangle, k}, \quad (3.5)$$

is an approximation of $\mathcal{A}_{\langle \ell_1 \rangle, \langle \ell_2 \rangle, k}$, which has the same entries as $\mathcal{A}_{\langle \ell_1 \rangle, \langle \ell_2 \rangle, k}$ in the positions $\langle \ell_1 \rangle \times J \cup I \times \langle \ell_2 \rangle$ for $k \in K$. Replacing $\mathcal{A}_{\langle \ell_1 \rangle, \langle \ell_2 \rangle, k}$ in (3.4) with the expression that appears in (3.5) gives rise to the approximation $\mathcal{B}$ of the form (3.1). Furthermore, $\mathcal{B}$ has the same entries as $\mathcal{A}$ for $(i, j, k) \in \mathcal{S}$. $\square$

Note that the Tucker decomposition outlined in Theorem 3.2 is not symmetric with respect to the three coordinates of the 3-tensor $\mathcal{A}$. In certain applications this can be useful, where one coordinate dimension is significantly larger than the other two, see e.g., [28] and the references therein.

Using the construction suggested in the proof of Theorem 3.2, one can derive an algorithm for computing $\mathcal{U}$ which is similar to the approximation algorithm for matrices outlined in Section 2. The complexity of this algorithm would be of order

$$\mathcal{O}(\max(p^3\ell_1, p^3\ell_2, p^2\ell_3) + p^8),$$

namely, the memory needed for the matrices C_1, C_2, C_3 and the tensor $\mathcal{U}$. The numerical properties of this method are currently under investigation.

3.2. 4-tensors. The ideas that we have developed in the previous subsection for 3-tensors can be extended easily to 4-tensors. However, in this case, we obtain a result that is symmetric in all coordinates.

THEOREM 3.3. *For a given generic tensor $\mathcal{A} \in \mathbb{R}^{\ell_1 \times \ell_2 \times \ell_3 \times \ell_4}$ let $I \subset \langle \ell_1 \rangle$, $J \subset \langle \ell_2 \rangle$, $K \subset \langle \ell_3 \rangle$, and $L \subset \langle \ell_4 \rangle$ be four sets of integers all of cardinality p and let*

$$\begin{aligned} \mathcal{C}_1 &= \mathcal{A}_{\langle \ell_1 \rangle, J, K, L} \in \mathbb{R}^{\ell_1 \times p \times p \times p}, \\ \mathcal{C}_2 &= \mathcal{A}_{I, \langle \ell_2 \rangle, K, L} \in \mathbb{R}^{p \times \ell_2 \times p \times p}, \\ \mathcal{C}_3 &= \mathcal{A}_{I, J, \langle \ell_3 \rangle, L} \in \mathbb{R}^{p \times p \times \ell_3 \times p}, \\ \mathcal{C}_4 &= \mathcal{A}_{I, J, K, \langle \ell_4 \rangle} \in \mathbb{R}^{p \times p \times p \times \ell_4} \end{aligned}$$

be four sections of $\mathcal{A}$. Let $C_1 \in \mathbb{R}^{p^3 \times \ell_1}$, $C_2 \in \mathbb{R}^{p^3 \times \ell_2}$, $C_3 \in \mathbb{R}^{p^3 \times \ell_3}$, $C_4 \in \mathbb{R}^{p^3 \times \ell_4}$ be matrix representations corresponding to the four sections $\mathcal{C}_1, \mathcal{C}_2, \mathcal{C}_3, \mathcal{C}_4$ of $\mathcal{A}$.

Then, there exists a tensor $\mathcal{U} \in \mathbb{R}^{p^3 \times p^3 \times p^3 \times p^3}$ such that the entries of the tensor

$$\mathcal{B} = \mathcal{U} \times_1 C_1 \times_2 C_2 \times_3 C_3 \times_4 C_4 \in \mathbb{R}^{\ell_1 \times \ell_2 \times \ell_3 \times \ell_4}.$$

coincide with the entries of $\mathcal{A}$ for $(i, j, k, l) \in \mathcal{S}$, where

$$\mathcal{S} := (\langle \ell_1 \rangle \times J \times K \times L) \cup (I \times \langle \ell_2 \rangle \times K \times L) \cup (I \times J \times \langle \ell_3 \rangle \times L) \cup (I \times J \times K \times \langle \ell_4 \rangle).$$

Proof. The proof can be obtained in a similar way as that of the 3-tensor case. $\square$

Note that usually the tensor $\mathcal{U}$ will not be unique, since we can unfold the given 4-tensor $\mathcal{A}$ to a matrix E in 6 different ways.

4. Experimental results. In this section, we present several experimental results that demonstrates the usefulness of our algorithm when applied to some synthetic

and real image data. The numerical results were run in MATLAB [22] on a PC with an Intel(R) Pentium(R) 4 CPU 3.20 GHz processor with 1 GB RAM memory, machine precision $2.2204 \cdot 10^{-16}$ and operating system SUSE Linux 10.2.

For the given matrix A and the computed rank k approximation B , we present the *total relative error (TRE) of the approximation*, defined as

$$\|A - B\|_F / \|A\|_F,$$

and the $\mathcal{S}$ -average error (SAE) as given in (2.4).

In order to demonstrate the statements made in Theorem 2.2, we compute U_{opt} in two ways. First we solve the least squares problem (2.2) (we denote the solution by U_{opt_1}), then we compute $U_{opt_2} = A_{I,J}^\dagger$ as in Theorem 2.2. As before, we denote by $t_{\max}$ the number of trials to find a well conditioned matrix $A_{I,J}$. We compare these results with $\tilde{U}_{opt}$ defined in (2.3) and with the solution of the complete least squares problem (denoted by CLS), given by $U_b = C^\dagger AR^\dagger$.

First of all, we observe that our algorithm tends to perform better with p and q chosen closer to the matrix rank.

Figure 4.1 portrays the original image of the ‘Tire’ picture from the Image Processing Toolbox of MATLAB [22], given by a matrix $A \in \mathbb{R}^{205 \times 232}$ of rank 205, the image compression given by the SVD (using the MATLAB function `svds` [23]) of rank 30 and the image compression given by $B_b = CU_bR$.



FIG. 4.1. Tire image (a) original, (b) SVD rank-30 approximation, (c) CLS rank-30 approximation, $t_{\max} = 100$.

The corresponding image compressions given by the approximations B_{opt_1} , B_{opt_2} and $\tilde{B}_{opt}$ are displayed respectively in Figure 4.2. Here, $t_{\max} = 100$ and $p = q = 30$. Note that the number of trials $t_{\max}$ is set to the large value of 100 for all simulations in order to be able to compare results for different (small and large) matrices.

In Table 4.1 we present the $\mathcal{S}$ -average and total relative errors of the image data



FIG. 4.2. Tire image compression to a rank-30 image (a) B_{opt_1} , (b) B_{opt_2} , (c) $\tilde{B}_{opt}$, $t_{\max} = 100$.

	rank	SAE	TRE
B_{svd}	30	0.0072	0.0851
B_b	30	0.0162	0.1920
B_{opt_1}	30	$1.6613 \cdot 10^{-26}$	0.8274
B_{opt_2}	30	$3.2886 \cdot 10^{-29}$	0.8274
$\tilde{B}_{opt}$	30	$1.9317 \cdot 10^{-29}$	0.8274

TABLE 4.1

Comparison of rank, $\mathcal{S}$ -average error (SAE) and total relative error (TRE) for the Tire image.

compression. Here, $B_b = CU_bR$, $B_{opt_2} = CU_{opt_2}R$ and $\tilde{B}_{opt} = C\tilde{U}_{opt}R$. Table 4.1 indicates that the less computationally costly FSVD with B_{opt_1} , B_{opt_2} and $\tilde{B}_{opt}$ obtains a smaller $\mathcal{S}$ -average error than the more expensive complete least squares solution CLS and the SVD. On the other hand, CLS and the SVD yield better results in terms of the total relative error. However, it should be noted that CLS is very costly and cannot be applied to very large matrices.

Figure 4.3 shows the results for the compression of the data for the original image of a ‘Camera man’ from the Image Processing Toolbox of MATLAB [22]. This data is a matrix $A \in \mathbb{R}^{256 \times 256}$ of rank 253 and the resulting image compression of rank 69 is derived using the SVD and the complete least square approximation CLS given by $B_b = CU_bR$. Notice that there is no discernible difference to the eye in the first two pictures. Figure 4.4 displays estimates given by the FSVD approximation $B_{opt_2} = CU_{opt_2}R$ and $\tilde{B}_{opt} = C\tilde{U}_{opt}R$, respectively. Here, we chose $t_{\max} = 100$ and $p = q = 80$. Correspondingly, in Table 4.2 we provide the resulting $\mathcal{S}$ -average and total relative errors. Due to the inability of the MATLAB least squares solver for large systems of equations, B_{opt_1} cannot be computed for large data. We see that the FSVD approximation performs well when considering the computational cost compared with the CLS approximation. Generally we want the total relative error (TRE) to be small, however, since we are not able to compute for ‘real world’ applications, we use the



FIG. 4.3. Camera man image (a) original, (b) SVD rank-69 approximation, (c) CLS rank-69 approximation, $t_{\max} = 100$.

SAE as an indicator for our error, since it can be computed effectively.

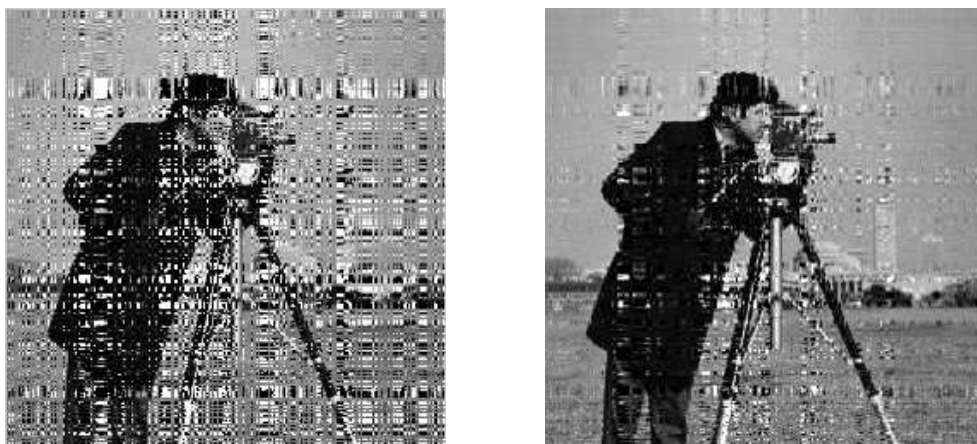


FIG. 4.4. Camera man image compression. FSVD rank-69 approximation with (a) $B_{opt_2} = CU_{opt_2}R$, (b) $\tilde{B}_{opt} = C\tilde{U}_{opt}R$. $t_{\max} = 100$.

	rank	SAE	TRE
B_{svd}	69	0.0020	0.0426
B_b	80	0.0049	0.0954
B_{opt_1}	—	—	—
B_{opt_2}	80	$3.7614 \cdot 10^{-27}$	1.5154
$\tilde{B}_{opt}$	69	$7.0114 \cdot 10^{-4}$	0.2175

TABLE 4.2

Comparison of rank, S -average error (SAE) and total relative error (TRE) for the Camera man image.

Figure 4.5 portrays the plots of total relative errors (TRE) versus number of

selected rows q and columns p for the camera man image. The fact that the error stabilizes as the number of rows and columns increases is quite perceptible. This can be attributed to the fact that this increase implies that the rank of the submatrix $A_{I,J}$ is getting closer to the rank of the original matrix and as such will result in a decreasing total relative error. Although the FSVD converges reasonably well, the SVD and CLS exhibit comparatively faster convergence.

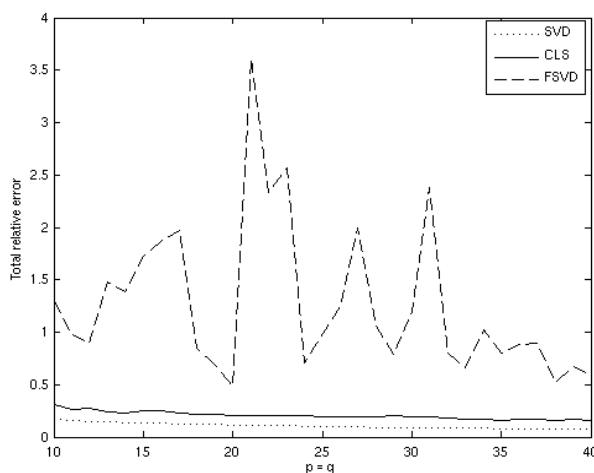


FIG. 4.5. *Camera man: total relative errors (TRE) versus number of selected rows q and columns p , $t_{\max} = 100$.*

Figure 4.6 portrays the plot of the $\mathcal{S}$ -average error (SAE) versus the number of chosen rows q and columns p for the camera man image. The theory suggests that better choices of q and p will result in better approximations and result in smaller average errors. This logic can be substantiated using the basis corresponding to specific singular values from the SVD of $A_{I,J}$. Larger singular values lead to better approximations. In addition, we would expect the graph to exhibit a decrease in error as the q and p values get larger because the resulting submatrix $A_{I,J}$ has a greater probability of being well conditioned. This in turn will provide us with a better Moore-Penrose inverse. Correspondingly, we can observe that the errors converge as p and q approach the matrix rank.

Figures 4.7 and 4.8 exhibit the result of the same analysis applied to an original high resolution image of ‘Canal at Night’ (an image from Max Lyons Digital Image Gallery [20]) given by a matrix $A \in \mathbb{R}^{812 \times 1200}$ of rank 812. Figure 4.7 shows the original and the SVD approximation image, of rank 159, with no ‘eye-recognizable’ difference, while Figure 4.8 shows the approximations given by CLS and FSVD (with $\tilde{B}_{opt} = C\tilde{U}_{opt}R$), respectively. The parameters here are $t_{\max} = 100$ and $p = q = 200$.

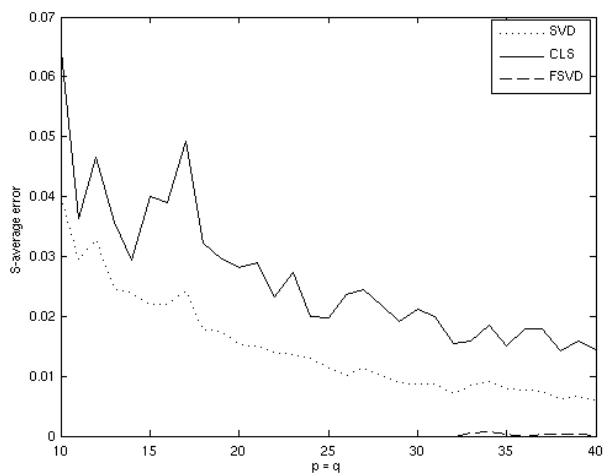


FIG. 4.6. *Camera man*: S -average errors versus (SAE) number of selected rows q and columns p , $t_{\max} = 100$.



FIG. 4.7. *Canal at Night* image (a) original, (b) SVD rank-159 approximation, $t_{\max} = 100$.



FIG. 4.8. *Canal at Night* image compression (a) CLS rank-200 approximation, (b) FSVD rank-159 approximation with $\tilde{B}_{opt}$, $t_{\max} = 100$.

In addition, Table 4.3 presents the corresponding $\mathcal{S}$ -average (SAE) and total relative errors (TRE) of the algorithms when performed with the ranks as stated. The results can be seen to be in accordance with the previous results. Hence, it can be concluded that the algorithm performs similar for images of different resolutions.

	rank	SAE	TRE
B_{svd}	159	0.0091	0.0951
B_b	200	0.0109	0.1492
B_{opt_1}	—	—	—
B_{opt_2}	200	$4.6677 \cdot 10^{-28}$	2.0919
$\tilde{B}_{opt}$	159	0.0031	0.2246

TABLE 4.3

Comparison of rank, $\mathcal{S}$ -average error (SAE) and total relative error (TRE) for the Canal at Night image.

To portray the significance of the cost in ‘real world’ approximations with very large data sets, we examine the algorithms on a uniformly created 2500×2500 random matrix of rank 50. Table 4.4 shows the results.

	Time (seconds)	Error SAE	Error TRE	rank
FSVD	0.123	9.6×10^{-15}	0.0012	42
SVD of full A	258.4	1.16×10^{-39}	6.2×10^{-30}	50
SVD of A with rank = 50	35.21	5.26×10^{-39}	2.19×10^{-29}	50
SVD of A with rank = 42	33.47	6.68×10^{-14}	2.58×10^{-4}	42

TABLE 4.4

Performance of SVD vs FSVD on 2500×2500 matrix of rank 50

As the size of the data matrix grows, the advantage in the computational cost of the FSVD makes a great difference. In addition, it is important to note that the SVD algorithm used was a built-in heuristic from MATLAB, which means it has been optimized for the system. On the other hand, FSVD has not been optimized and runs solely as an iterative process. Moreover, the structure of the FSVD allows to make use of parallelization, or to take advantage of very sparse matrices by optimizing the choice of columns and rows. In general, the choice of p and q is strongly problem-dependent, therefore any additional information about the problem should be used to select the appropriate number of columns and rows. If we are interested in a small number of columns and rows one possibility would be selecting several sets I, J and choosing the one giving the best approximation among all. The above concepts are research topics of future work.

5. Conclusions. We have introduced an iterative updating procedure for the fast computation of rank k approximations of large dense matrices. We have demonstrated the properties of the methods with real and synthetic data and we have shown how the ideas can be extended to tensors.

Acknowledgment. The authors would like to thank the DFG Research Center MATHEON and the Berlin Mathematical School for the financial support. We also thank anonymous referees for valuable comments and suggestions which led to improvements of this paper.

REFERENCES

- [1] O. Alter, P. Brown, and D. Botstein. Processing and modeling genome-wide expression data using singular value decomposition. *Microarrays: Optical Technologies and Informatics, Proc. SPIE*, 4266:171–186, 2001.
- [2] J. Demmel and W. Kahan. Accurate singular values of bidiagonal matrices. *SIAM J. Sci. Statist. Comput.*, 11(5):873–912, 1990.
- [3] A. Deshpande and S. Vempala. Adaptive sampling and fast low-rank matrix approximation. *Electronic Colloquium on Computational Complexity (ECCC) TR06-042*, March 2006.
- [4] P. Drineas, R. Kannan, and M.W. Mahoney. Fast Monte Carlo algorithms for matrices I–III: Computing a compressed approximate matrix decomposition. *SIAM J. Comput.*, 36(1):132–206, 2006.
- [5] P. Drineas, M.W. Mahoney, and S. Muthukrishnan. Polynomial time algorithm for column-row-based relative-error low-rank matrix approximation. *Technical Report TR 2006-04*, DIMACS, March 2006.
- [6] P. Drineas, M.W. Mahoney, and S. Muthukrishnan. Subspace sampling and relative-error matrix approximation: Column-row-based methods. *Proceedings of the 14th Annual European Symposium on Algorithms (ESA)*, 304–314, 2006.
- [7] N. Eriksson. Tree construction using singular value decomposition, 347–358. L. Pachter and B. Sturmfels (editors), *Algebraic Statistics for Computational Biology*, Cambridge University Press, 2005.
- [8] M. Espig, L. Grasedyck, and W. Hackbusch. Black box low tensor-rank approximation using fiber-crosses. *Constr. Approx.*, 30(3):557–597, 2009.
- [9] S. Friedland, M. Kaveh, A. Niknejad, and H. Zare. Fast Monte-Carlo low rank approximations for matrices. *Proc. IEEE SoSE*, 218–223, Los Angeles, 2006.
- [10] S. Friedland and A. Torokhti. Generalized rank-constrained matrix approximations. *SIAM J. Matrix Anal. Appl.*, 29(2):656–659, 2007.
- [11] P.J. Gemperline, K.H. Miller, T.L. West, J.E. Weinstein, J.C. Hamilton, and J.T. Bray. Principal component analysis, trace elements, and blue crab shell disease. *Analytical Chemistry*, 64(9):523A–532A, 1992.
- [12] G.H. Golub and C.F. Van Loan. *Matrix Computations*, 3rd edition. John Hopkins Univ. Press, Baltimore, 1996.
- [13] S.A. Goreinov and E.E. Tyrtyshnikov. The maximum-volume concept in approximation by low-rank matrices. *Contemporary Mathematics*, 280:47–51, 2001.
- [14] S.A. Goreinov, E.E. Tyrtyshnikov, and N.L. Zamarashkin. Pseudo-skeleton approximations of matrices. *Dokl. Akad. Nauk*, 343(2):151–152, 1995.
- [15] S.A. Goreinov, E.E. Tyrtyshnikov, and N.L. Zamarashkin. A theory of pseudoskeleton approximations. *Linear Algebra Appl.*, 261:1–21, 1997.

- [16] S. Har-Peled. Low rank matrix approximation in linear time. Manuscript, 2006, available at <http://valis.cs.uiuc.edu/~sariel/papers/05/lrank/lrank.pdf>.
- [17] J. Kamm and J.G. Nagy. Kronecker product and SVD approximations in image restoration. *Linear Algebra Appl.*, 284:177–192, 1998.
- [18] P.M. Kroonenberg. *Three-Mode Principal Component Analysis: Theory and Applications*. DSWO Press, Leiden, 1983.
- [19] R.B. Lehoucq, D.C. Sorensen, and C. Yang. *ARPACK User's Guide: Solution of Large-Scale Eigenvalue Problems With Implicitly Restarted Arnoldi Methods (Software, Environments, and Tools)*. SIAM Publications, 1998.
- [20] M. Lyons. Max Lyons Digital Image Gallery. Available at <http://www.tawbaware.com/maxlyons/index.html>.
- [21] M.W. Mahoney, M. Maggioni, and P. Drineas. Tensor-CUR decompositions for tensor-based data. *Proceedings of the 12th Annual ACM SIGKDD Conference*, 327–336, 2006.
- [22] MATLAB Image Processing Toolbox Version 5.4 (R2007a). The MathWorks, inc., 24 Prime Park Way, Natick, MA 01760-1500, USA, 2007.
- [23] MATLAB, Version 7.4.0.336 (R2007a). The MathWorks, inc., 24 Prime Park Way, Natick, MA 01760-1500, USA, 2007.
- [24] M. Nkengla. *Low Rank Approximations of Matrices and Tensors*. PhD Thesis, University of Illinois at Chicago, 2010.
- [25] P. Paschou, M. W. Mahoney, A. Javed, J.R. Kidd, A.J. Pakstis, S. Gu, K.K. Kidd, and P. Drineas. Intra- and interpopulation genotype reconstruction from tagging SNPs. *Genome Res.*, 17:96–107, 2007.
- [26] M. Rudelson and R. Vershynin. Sampling from large matrices: An approach through geometric functional analysis. *J. ACM*, 54, 2007.
- [27] T. Sarlos. Improved approximation algorithms for large matrices via random projections. *Proceedings of the 47th Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 143–152, 2006.
- [28] A. Stegeman, J.M.F. Ten Berge, and L. De Lathauwer. Sufficient conditions for uniqueness in Candecomp/Parafac and Indscal with random component matrices. *Psychometrika*, 71:219–229, 2006.
- [29] G.W. Stewart. Updating a rank-revealing ULV decomposition. *SIAM J. Matrix Anal. Appl.*, 14:494–499, 1993.
- [30] L.R. Tucker. Some mathematical notes on three-mode factor analysis. *Psychometrika*, 31(3):279–311, 1966.
- [31] M.A.O. Vasilescu and D. Terzopoulos. Multilinear image analysis for facial recognition. *Proceedings of the International Conference on Pattern Recognition (ICPR'2002)*, 3:511–514, Quebec City, Canada, 2002.