

Low-rank approximation: from matrices to tensors

Konstantin Usevich, CNRS, CRAN (UMR 7039 / Univ. Lorraine),
Nancy

31.05.2024, Marseille, workshop on Low-rank optimization

Overview

Matrix LRA

Tensors

Factorization

(CP) decomposition

CP approximation

Extensions

Overview

Matrix LRA

Tensors

Factorization

(CP) decomposition

CP approximation

Extensions

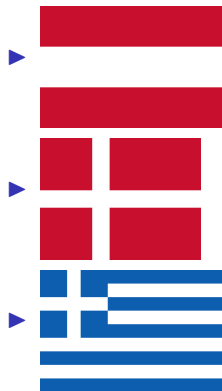
Matrix rank: factorization view

Rank of $\mathbf{X} \in \mathbb{R}^{m \times n}$ (or $\mathbb{C}^{m \times n}$)

- ▶ $\stackrel{\text{def}}{=}$ number of linearly independent columns/ rows in $\mathbf{X}$

How do low-rank matrices look like?

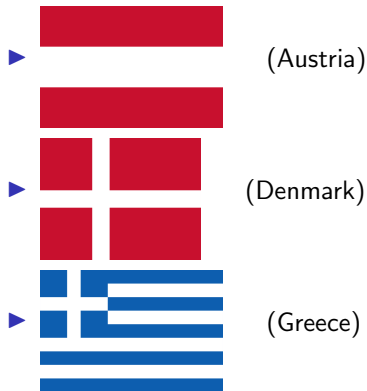
Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$):



Inspired by talks of Alex Townsend

How do low-rank matrices look like?

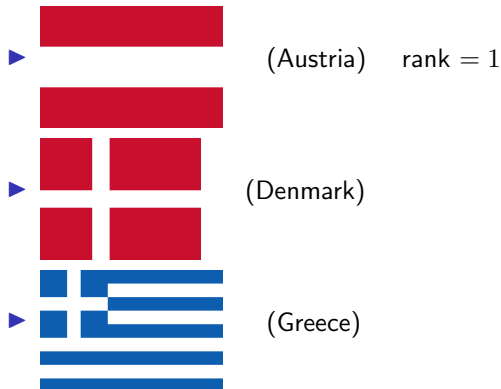
Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$):



Inspired by talks of Alex Townsend

How do low-rank matrices look like?

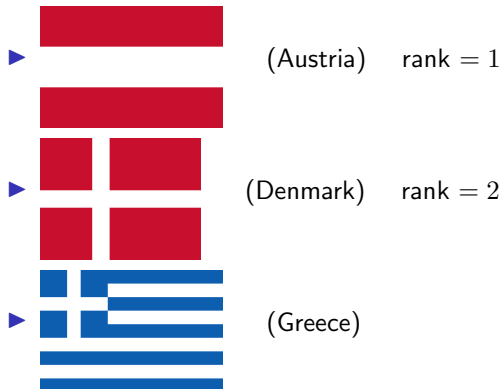
Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$):



Inspired by talks of Alex Townsend

How do low-rank matrices look like?

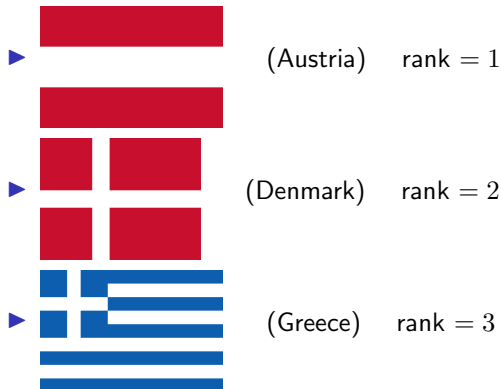
Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$):



Inspired by talks of Alex Townsend

How do low-rank matrices look like?

Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$):



Inspired by talks of Alex Townsend

How do low-rank matrices look like?

Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$), cont'd:



How do low-rank matrices look like?

Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$), cont'd:



(Scotland)



(Wales)

How do low-rank matrices look like?

Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$), cont'd:



(Scotland) rank $\approx \frac{m}{2}$ (symmetry)



(Wales)

How do low-rank matrices look like?

Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$), cont'd:



(Scotland) $\text{rank} \approx \frac{m}{2}$ (symmetry)



(Wales) $\text{rank} \approx \frac{5}{6}m$ (finite support)

How do low-rank matrices look like?

Ranks of flags (as $\mathbf{X} \in \mathbb{R}^{m \times n}$), cont'd:



(Scotland) rank $\approx \frac{m}{2}$ (symmetry)



(Wales) rank $\approx \frac{5}{6}m$ (finite support)

- ▶ random matrix (with a.c. probability distribution):
rank = $\min(m, n)$ a.s. (with probability 1)
- ▶ many interesting matrices are [well approximated by] low-rank
[Townsend, Udell, 2017]

Singular value decomposition (SVD)

a/the (“economy size”) SVD ($m \leq n$):

$$\begin{matrix} & n \\ m & \mathbf{X} \end{matrix} = \mathbf{U} \begin{matrix} \sigma_1 & & \\ & \ddots & \\ & & \sigma_m \end{matrix} \mathbf{V}^T = \sum_{k=1}^m \sigma_k \mathbf{u}_k \mathbf{v}_k^T$$

where

- ▶ $\mathbf{U}^T \mathbf{U} = \mathbf{V}^T \mathbf{V} = \mathbf{I}$ — semi-orthogonal matrices of singular vectors
 $\mathbf{U} = [\mathbf{u}_1 \quad \cdots \quad \mathbf{u}_m], \quad \mathbf{V} = [\mathbf{v}_1 \quad \cdots \quad \mathbf{v}_m]$
- ▶ $\sigma_1 \geq \cdots \geq \sigma_m \geq 0$ — singular values

Singular value decomposition (SVD)

a/the (“economy size”) SVD ($m \leq n$):

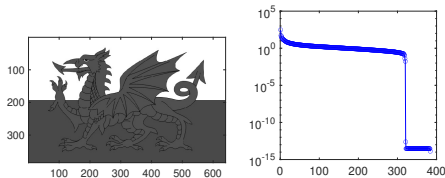
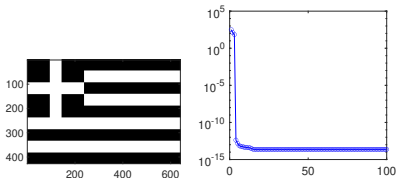
$${}_m \begin{matrix} n \\ \mathbf{X} \end{matrix} = \begin{matrix} \mathbf{U} \end{matrix} \begin{matrix} \sigma_1 & & \\ & \ddots & \\ & & \sigma_m \end{matrix} \begin{matrix} \mathbf{V}^T \end{matrix} = \sum_{k=1}^m \sigma_k \mathbf{u}_k \mathbf{v}_k^T$$

where

- ▶ $\mathbf{U}^T \mathbf{U} = \mathbf{V}^T \mathbf{V} = \mathbf{I}$ — semi-orthogonal matrices of singular vectors

$$\mathbf{U} = [\mathbf{u}_1 \quad \cdots \quad \mathbf{u}_m], \quad \mathbf{V} = [\mathbf{v}_1 \quad \cdots \quad \mathbf{v}_m]$$

- ▶ $\sigma_1 \geq \cdots \geq \sigma_m \geq 0$ — singular values



SVD and low-rank approximation

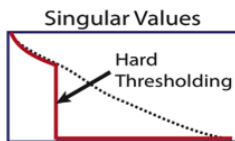
Eckart-Young(-Mirsky(-Schmidt)) theorem: **best rank- r** approximation

$$\min_{\hat{\mathbf{X}}} \|\hat{\mathbf{X}} - \mathbf{X}\| \quad \text{subject to} \quad \text{rank } \hat{\mathbf{X}} \leq r$$

in any unitarily invariant* norm $\|\cdot\|$ is given by

$$\text{tSVD}_r(\mathbf{X}) := \begin{bmatrix} \mathbf{U} & \begin{bmatrix} \sigma_1 & & \\ & \ddots & \\ & & \sigma_r \\ & & & \mathbf{0} \end{bmatrix} & \mathbf{V}^T \end{bmatrix}$$

truncated SVD



* Examples: Frobenius norm $\|\mathbf{X}\|_F^2 := \sum_{i,j=1}^{m,n} X_{ij}^2$, spectral norm...

SVD and low-rank approximation

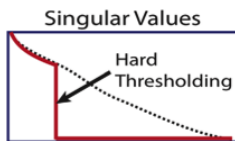
Eckart-Young(-Mirsky(-Schmidt)) theorem: **best rank- r** approximation

$$\min_{\hat{\mathbf{X}}} \|\hat{\mathbf{X}} - \mathbf{X}\| \quad \text{subject to} \quad \text{rank } \hat{\mathbf{X}} \leq r$$

in any unitarily invariant* norm $\|\cdot\|$ is given by

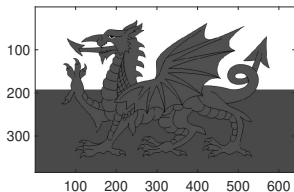
$$\text{tSVD}_r(\mathbf{X}) := \begin{bmatrix} \mathbf{U} & \begin{bmatrix} \sigma_1 & \dots & \sigma_r \\ \mathbf{0} \end{bmatrix} & \mathbf{V}^T \end{bmatrix} = \begin{bmatrix} \mathbf{U}_{1:r,:} & \begin{bmatrix} \sigma_1 & \dots & \sigma_r \\ \mathbf{0} \end{bmatrix} & \mathbf{V}_{1:r,:}^T \end{bmatrix} = \sum_{k=1}^r \sigma_k \mathbf{u}_k \mathbf{v}_k^T$$

truncated SVD

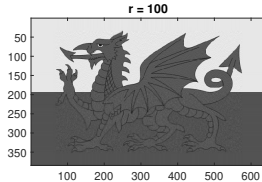
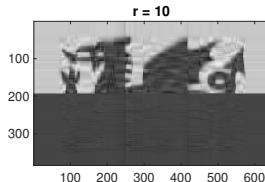
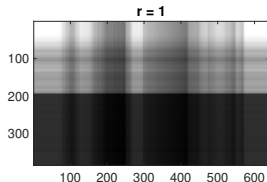
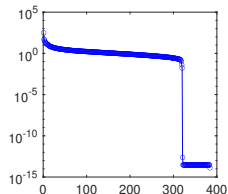


* Examples: Frobenius norm $\|\mathbf{X}\|_F^2 := \sum_{i,j=1}^{m,n} X_{ij}^2$, spectral norm...

Low-rank approximations: example



$$\approx \begin{bmatrix} U_{1:r,:} \\ \sigma_r \\ V_{1:r,:}^T \end{bmatrix}$$



Computational aspects: full and partial SVD

$${}_m \boxed{\overset{n}{\mathbf{X}}} = \boxed{\mathbf{U}} \boxed{\begin{matrix} \sigma_1 & & \\ & \ddots & \\ & & \sigma_m \end{matrix}} \boxed{\mathbf{V}^T} = \sum_{k=1}^m \sigma_k \mathbf{u}_k \mathbf{v}_k^T$$

- ▶ Full SVD ($m \leq n$): $\mathcal{O}(m^2 n)$ ([Golub, Reinsh, 1970])
- ▶ $(\mathbf{u}_k, \sigma_k^2)$ — eigenvalues of $\mathbf{C} = \mathbf{X}\mathbf{X}^T$: $\mathcal{O}(m^3)$ if $m \ll n$
- ▶ Truncated SVD (r first eigenvalues/vectors):
 - ▶ Iterative (e.g., Lanczos/Arnoldi) algorithms (e.g., [Simon, 1984])
“generalization of power method”
 - ▶ Randomized SVD [Halko, Martinsson, Tropp, 2011]

Both roughly $\boxed{\mathcal{O}(Mr)}$, where $M \leq mn$ (see next slide):

Power iteration

Goal: find the top eigenpair $\mathbf{u}_1, \lambda_1$ of $\mathbf{C}$.

▶ Set $\mathbf{u}^{(0)} \in \mathbb{R}^{m \times m}$ random.

▶ Iterate $\mathbf{u}^{(k+1)} = \frac{\mathbf{C}\mathbf{u}^{(k)}}{\|\mathbf{C}\mathbf{u}^{(k)}\|}$

Case $\mathbf{C} = \mathbf{X}\mathbf{X}^T$: cost $\boxed{\approx \mathcal{O}(Mn_{iter})}$, where

▶ full matrices: $M = mn$

▶ sparse matrices: $M = \# \text{non-zero entries}$

▶ structured matrices: (e.g., $M = n \log(n)$ for Hankel)

$M = \text{cost of matrix-vector product (e.g., } \mathbf{X}\mathbf{v})$

▶ used e.g., by Google for PageRank

▶ do not need to store the matrix $\mathbf{X}$

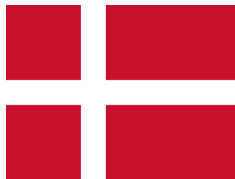
▶ generalizes to rank- r approximation (cost $\mathcal{O}(Mr + mr^2)$)

Cross approximation

Observation. A rank- r matrix is uniquely determined by the $r \times r$ cross if the $r \times r$ submatrix $\mathbf{X}_{\mathcal{I},\mathcal{J}}$ is nonsingular

$$\mathbf{X} = \mathbf{X}_{:, \mathcal{J}} (\mathbf{X}_{\mathcal{I}, \mathcal{J}})^{-1} \mathbf{X}_{\mathcal{I}, :} \quad (*)$$

Example. $r = 2$:

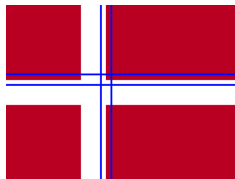


Cross approximation

Observation. A rank- r matrix is uniquely determined by the $r \times r$ cross if the $r \times r$ submatrix $\mathbf{X}_{\mathcal{I},\mathcal{J}}$ is nonsingular

$$\mathbf{X} = \mathbf{X}_{:, \mathcal{J}} (\mathbf{X}_{\mathcal{I}, \mathcal{J}})^{-1} \mathbf{X}_{\mathcal{I}, :} \quad (*)$$

Example. $r = 2$:



Cross approximation

Observation. A rank- r matrix is uniquely determined by the $r \times r$ cross if the $r \times r$ submatrix $\mathbf{X}_{\mathcal{I},\mathcal{J}}$ is nonsingular

$$\mathbf{X} = \mathbf{X}_{:, \mathcal{J}} (\mathbf{X}_{\mathcal{I}, \mathcal{J}})^{-1} \mathbf{X}_{\mathcal{I}, :} \quad (*)$$

Example. $r = 2$:



Cross approximation

Observation. A rank- r matrix is uniquely determined by the $r \times r$ cross if the $r \times r$ submatrix $\mathbf{X}_{\mathcal{I},\mathcal{J}}$ is nonsingular

$$\mathbf{X} = \mathbf{X}_{:, \mathcal{J}} (\mathbf{X}_{\mathcal{I}, \mathcal{J}})^{-1} \mathbf{X}_{\mathcal{I}, :} \quad (*)$$

Example. $r = 2$:



$\Rightarrow (*)$ can be used as an approximation

Cross approximation

If the matrix is huge or expensive to compute:

CUR (cross, or pseudo-skeleton) approximation (for size- r subsets $\mathcal{I}, \mathcal{J}$):

$$\hat{\mathbf{X}}_{cross}(\mathcal{I}, \mathcal{J}) = \mathbf{X}_{:, \mathcal{J}} (\mathbf{X}_{\mathcal{I}, \mathcal{J}})^{-1} \mathbf{X}_{\mathcal{I}, :}$$

[Mahoney, Drineas, 2012]

Advantages:

- ▶ Need a **small portion (cross)** of the matrix ($O(r(m+n))$)
- ▶ Quasi-optimality (thm. in [Goreinov, Tyrtyshnikov, 2001])

$$\|\hat{\mathbf{X}}_{cross}^* - \mathbf{X}\|_{max} \leq (r+1)\sigma_r(\mathbf{X})$$

for $|\det(\mathbf{X}_{\mathcal{I}, \mathcal{J}})| \rightarrow \max$

- ▶ iterative or randomized strategies to select $\mathcal{I}, \mathcal{J}$

Extensions of the basic problem

$$\min_{\substack{\hat{\mathbf{X}}=\mathbf{A}\mathbf{B} \\ \mathbf{A}\in\mathbb{R}^{m\times r}, \mathbf{B}\in\mathbb{R}^{r\times n}}} \|\mathbf{X} - \hat{\mathbf{X}}\|_F$$

- ▶ Other norms (e.g. $\|\mathbf{X}\|_W^2 = \sum_{i,j} W_{ij} X_{ij}^2$), missing data
- ▶ Constraint on the matrix:
structured $\hat{\mathbf{X}}$ — structured low-rank approximation
- ▶ Constraints on the factors $\mathbf{A}$, $\mathbf{B}$ (e.g., nonnegative)

Weighted (unstructured) LRA

$$\min_{\hat{\mathbf{X}}} \|\mathbf{X} - \hat{\mathbf{X}}\|_W^2 \quad \text{subject to} \quad \text{rank}(\hat{\mathbf{X}}) \leq r$$

where $\|\mathbf{X}\|_W^2 = \sum_{i,j=1}^{m,n} \mathbf{X}_{ij}^2 W_{ij}$, $W_{ij} \in (0; +\infty)$ weighted norm

- ▶ $W_{ij} \equiv 1$ (or $\text{rank } W = 1$) $\rightarrow$ solution by SVD
- ▶ In general case, no closed form solution:
Gillis, Glineur, **Low-Rank Matrix Approximation with Weights or Missing Data is NP-hard**, *SIMAX*, 2011.

Extended semi-norms and matrix completion

$$\min_{\hat{\mathbf{X}}} \|\mathbf{X} - \hat{\mathbf{X}}\|_W^2 \text{ subject to } \text{rank}(\hat{\mathbf{X}}) \leq r$$

$$\|\mathbf{X}\|_W^2 := \sum_{i,j=1}^{m,n} \mathbf{X}_{ij}^2 W_{ij}, \quad W_{ij} \in [0; +\infty] \quad \text{weighted extended semi-norm}$$

- **fixed** values: $W_{ij} = +\infty \iff \text{constraint } \mathbf{X}_{ij} = \hat{\mathbf{X}}_{ij}$
- **missing** values: $W_{ij} = 0 \iff \hat{\mathbf{X}}_{ij}$ is not important

Extended semi-norms and matrix completion

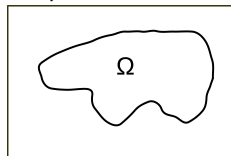
$$\min_{\hat{\mathbf{X}}} \|\mathbf{X} - \hat{\mathbf{X}}\|_W^2 \quad \text{subject to} \quad \text{rank}(\hat{\mathbf{X}}) \leq r$$

$$\|\mathbf{X}\|_W^2 := \sum_{i,j=1}^{m,n} \mathbf{X}_{ij}^2 W_{ij}, \quad W_{ij} \in [0; +\infty] \quad \text{weighted extended semi-norm}$$

- **fixed** values: $W_{ij} = +\infty \iff \text{constraint } \mathbf{X}_{ij} = \hat{\mathbf{X}}_{ij}$
- **missing** values: $W_{ij} = 0 \iff \hat{\mathbf{X}}_{ij}$ is not important

Extreme case: $W_{ij} \in \{0, +\infty\}$ — exact matrix completion

$$\begin{aligned} &\min \text{rank}(\hat{\mathbf{X}}) \\ &\text{subject to } (\hat{\mathbf{X}})_{ij} = X_{ij}, \quad \forall (i, j) \in \Omega \end{aligned}$$

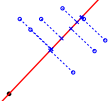


Structured low-rank approximation

[Markovsky, 2008] : **Problem (SLRA)**. Given a **structured** matrix $\mathbf{X} \in \mathcal{S}$

$$\underset{\hat{\mathbf{X}}}{\text{minimize}} \|\mathbf{X} - \hat{\mathbf{X}}\|_W^2 \quad \text{subject to} \quad \hat{\mathbf{X}} \in \mathcal{S} \text{ and } \text{rank } \hat{\mathbf{X}} \leq r$$

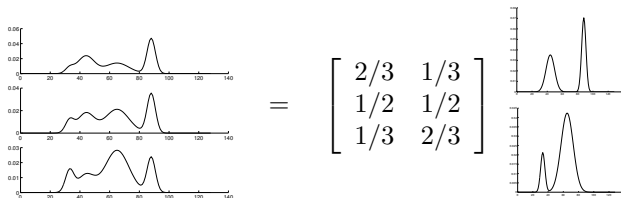
Data $\approx$ low-complexity model

structure $\mathcal{S}$	approximation problem
unstructured	fit by r -dim. subspace 
Hankel	fitting by complex exponentials
block-Hankel	linear system identification model reduction
Sylvester	approx. greatest common divisor
generalized Vandermonde	fit set of points by algebraic hypersurfaces

Source separation $\leftrightarrow$ matrix factorization

Example from spectroscopy:

- ▶ each observed spectrum is a linear combination of “pure spectra”
- ▶ different conditions — different coefficients.

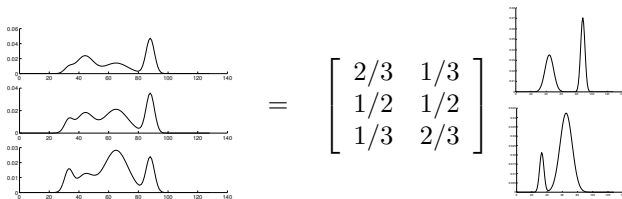


Instantaneous mixture model: $\mathbf{X}(x, \lambda) = \sum_k \mathbf{a}_k(x) \mathbf{s}_k(\lambda)$

Source separation $\leftrightarrow$ matrix factorization

Example from spectroscopy:

- ▶ each observed spectrum is a linear combination of “pure spectra”
- ▶ different conditions — different coefficients.



Instantaneous mixture model: $\mathbf{X}(x, \lambda) = \sum_k \mathbf{a}_k(x) s_k(\lambda)$

$$\boxed{\mathbf{X}} = \begin{array}{c} | \\ \mathbf{a}_1 \\ | \end{array} \begin{array}{c} \text{---} \mathbf{s}_1^T \text{---} \\ + \dots + \\ \text{---} \mathbf{s}_r^T \text{---} \\ | \end{array} = \boxed{\mathbf{A}}^r \boxed{\mathbf{S}^T} \quad \text{matrix factorization}$$

Can we recover $\mathbf{A}$, $\mathbf{S}$ from $\mathbf{X}$?

Matrix factorizations are non-unique

Does not happen for matrices:

$$\boxed{\mathbf{M}} = \boxed{\mathbf{A}} \boxed{\mathbf{S}^T} = \boxed{\mathbf{A}} \boxed{\mathbf{Q}} \boxed{\mathbf{Q}^{-1}} \boxed{\mathbf{S}^T}$$

nonunique (change of basis)

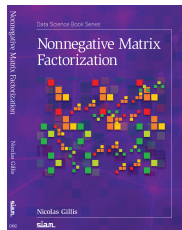
However, **constraints** on factors can guarantee essential uniqueness

$$\mathbf{Q} = \underbrace{\mathbf{\Lambda}}_{\text{diagonal}} \cdot \underbrace{\mathbf{\Pi}}_{\text{permutation}}$$

Examples of constraints:

nonnegativity

$$\mathbf{A} \geq 0, \mathbf{S} \geq 0$$

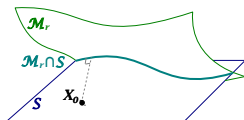


independence
(constraint on $\mathbf{S}$)



Other tools/link to optimization

- $\mathcal{M}_r^{m \times n} = \{\mathbf{X} \in \mathbb{R}^{m \times n} \mid \text{rank}(\mathbf{X}) = r\}$ — smooth manifold



visualisation
of SLRA

→ optimization on manifolds [Absil, Mahoney, Sepulchre, 2008], [Boumal, 2023]

- $\mathcal{M}_{\leq r}^{m \times n} = \{\mathbf{X} \mid \text{rank}(\mathbf{X}) \leq r\}$ — algebraic variety (stratified set)
link to determinantal representations of algebraic varieties
- low rank $\leftrightarrow$ sparsity of singular values

$$\underbrace{\|(\sigma_1, \dots, \sigma_m)\|_0}_{\text{rank}} \leftrightarrow \underbrace{\|(\sigma_1, \dots, \sigma_m)\|_1}_{\text{nuclear norm}}$$

other sparsity-promoting penalties ...

- other norms/divergences/losses....
- dynamical low-rank approximation $\mathbf{X}(t) \in \mathcal{M}_r$, varies over time

Matrix factorization $\leftrightarrow$ neural networks

Matrix factorization: $\mathbf{X} = \mathbf{U}\mathbf{V}^T$, $\mathbf{U} = [\mathbf{u}_1 \ \cdots \ \mathbf{u}_r]$, $\mathbf{V} = [\mathbf{v}_1 \ \cdots \ \mathbf{v}_r]$

Decompose linear map $\mathbf{f} : \mathbb{R}^m \rightarrow \mathbb{R}^n$, $\mathbf{f}(\mathbf{z}) = \mathbf{X}\mathbf{z}$ as

$$\mathbf{f}(\mathbf{z}) = \mathbf{u}_1 \cdot (\mathbf{v}_1^\top \mathbf{z}) + \cdots + \mathbf{u}_r \cdot (\mathbf{v}_r^\top \mathbf{z}),$$

Matrix factorization $\leftrightarrow$ neural networks

Given **nonlinear** map $\mathbf{f} : \mathbb{R}^m \rightarrow \mathbb{R}^n$, decompose it as

$$\mathbf{f}(\mathbf{z}) = \mathbf{u}_1 \mathbf{g}_1(\mathbf{v}_1^\top \mathbf{z}) + \cdots + \mathbf{u}_r \mathbf{g}_r(\mathbf{v}_r^\top \mathbf{z}),$$

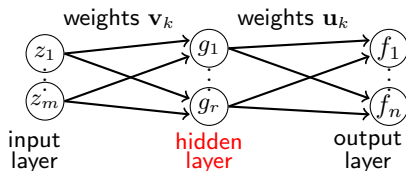
where $g_k(t)$ are **univariate functions** (see. e.g., [Comon,Qi,U., 2017])

Matrix factorization $\leftrightarrow$ neural networks

Given **nonlinear** map $\mathbf{f} : \mathbb{R}^m \rightarrow \mathbb{R}^n$, decompose it as

$$\mathbf{f}(\mathbf{z}) = \mathbf{u}_1 g_1(\mathbf{v}_1^\top \mathbf{z}) + \cdots + \mathbf{u}_r g_r(\mathbf{v}_r^\top \mathbf{z}),$$

where $g_k(t)$ are **univariate functions** (see. e.g., [Comon, Qi, U., 2017])



See for example:

- ▶ One-hidden layer model: ([Marcotte, Gribonval, Peyré, 2024])
- ▶ deep linear networks (products of matrices): [Malgouyres, 2020]
- ▶ deep NMF [Leplat et al, 2024]

Overview

Matrix LRA

Tensors

Factorization

(CP) decomposition

CP approximation

Extensions

Tensors: some notation

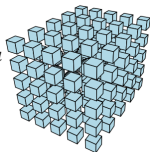
Tensor product of vector spaces (over a field $\mathbb{F} = \mathbb{R}$ or $\mathbb{C}$):

► $\mathcal{T} \in \mathbb{F}^{n_1 \times \dots \times n_d} \stackrel{\text{def}}{=} \mathbb{F}^{n_1} \otimes \mathbb{F}^{n_2} \otimes \dots \otimes \mathbb{F}^{n_d}$

► d -way array $\mathcal{T} = [\mathcal{T}_{i_1, i_2, \dots, i_d}]_{i_1, i_2, \dots, i_d=1}^{n_1, n_2, \dots, n_d} \in \mathbb{F}^{n_1 \times \dots \times n_d}$

► 3-rd order **tensor**: $\mathcal{T} = [\mathcal{T}_{ijk}]_{i,j,k=1}^{I,J,K} \in \mathbb{F}^{I \times J \times K}$

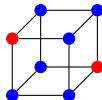
► tensor (outer) product: $\mathcal{T} = \mathbf{a} \otimes \mathbf{b} \otimes \mathbf{c}$: $\mathcal{T}_{ijk} = a_i b_j c_k$



Examples:

► $\mathcal{T} = \begin{bmatrix} 1 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 0 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} 2 \\ -3 \end{bmatrix}$ $\mathcal{T}_{:, :, 1} = \begin{bmatrix} 0 & 2 \\ 0 & 2 \end{bmatrix}$, $\mathcal{T}_{:, :, 2} = \begin{bmatrix} 0 & -3 \\ 0 & -3 \end{bmatrix}$,

► $\mathcal{T} = \mathbf{e}_1 \otimes \mathbf{e}_1 \otimes \mathbf{e}_1 + \mathbf{e}_2 \otimes \mathbf{e}_2 \otimes \mathbf{e}_2$ **diagonal tensor**



$\mathcal{T}_{:, :, 1} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}$, $\mathcal{T}_{:, :, 2} = \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix}$,

Some history

- ▶ late 1800s–early 1900s: algebraic geometry (Sylvester, Terracini, Segre, ...)
- ▶ 1927 (Hitchcock): introduced tensor rank
- ▶ Multiway models in psychometrics: Cattell (1940s), Tucker (1960s), Harshman (1970s)
- ▶ Popular models in chemometrics (1980s)
- ▶ Theory of complexity: Strassen (1980s)
- ▶ Signal processing: Comon (1990s)
- ▶ ...

Some references

Modern references (tensor decompositions):

- ▶ [Kolda, Bader, 2009]: generic entry reference
- ▶ [Comon, 2009, 2014]: focus on CPD and its properties
- ▶ [Landsberg, 2012]: algebraic viewpoint
- ▶ [Grasedyck et al, 2013]: focus on approximation, scientific computing
- ▶ [Sidiropoulos et al, 2017]: more recent overview on uniqueness
- ▶ [Cichocki et al, 2016]: book on tensor networks

Reminder: matrix rank

- $\stackrel{\text{def}}{=}$ smallest r such that $\mathbf{X}$ can be **factorized** as

$$\begin{matrix} & n \\ & \boxed{\mathbf{X}} \\ m & \end{matrix} = \begin{matrix} & r \\ & \boxed{\mathbf{A}} \\ m & \end{matrix}^r \begin{matrix} & n \\ & \boxed{\mathbf{B}} \\ & \end{matrix}$$

- $\stackrel{\text{def}}{=}$ minimal r such that $\mathbf{X}$ can be **decomposed** as

$$\boxed{\mathbf{X}} = \begin{matrix} | \\ \mathbf{a}_1 \\ | \end{matrix} \begin{matrix} \text{---} & \mathbf{b}_1^T & \text{---} \\ \hline \end{matrix} + \cdots + \begin{matrix} | \\ \mathbf{a}_r \\ | \end{matrix} \begin{matrix} \text{---} & \mathbf{b}_r^T & \text{---} \\ \hline \end{matrix}$$

Generalization to tensors leads to **two different versions of rank!**

1. Multilinear rank (Tucker) — **factorization**
 2. Tensor rank (CPD) — **decomposition**
- different decompositions for different purpose
 - most other decompositions are combinations of CP and Tucker

Overview

Matrix LRA

Tensors

Factorization

(CP) decomposition

CP approximation

Extensions

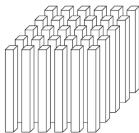
► Matrices: columns and rows:



► Tensors: fibers:

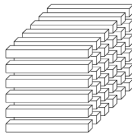
columns

$$\mathcal{T}_{:,j,k}$$



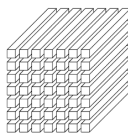
rows

$$\mathcal{T}_{i,:,k}$$



tubes

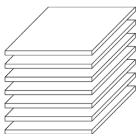
$$\mathcal{T}_{i,j,:}$$



► Tensors: slices:

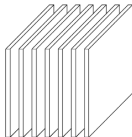
horizontal

$$\mathcal{T}_{i,:,:}$$



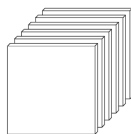
lateral

$$\mathcal{T}_{:,j,:}$$



frontal

$$\mathcal{T}_{:,:,k}$$



Generalizing matrix rank #1: multilinear rank

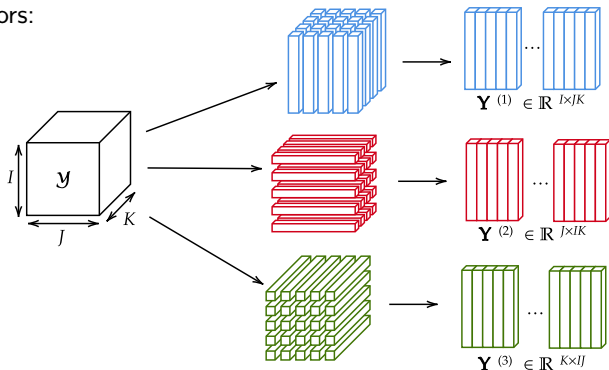
Matrix rank $\stackrel{\text{def}}{=}$ dimension of column or row span: 

For tensors:

Generalizing matrix rank #1: multilinear rank

Matrix rank $\stackrel{\text{def}}{=}$ dimension of column or row span: 

For tensors:

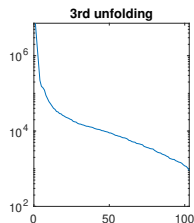
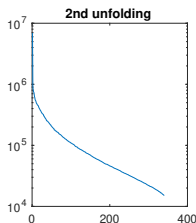
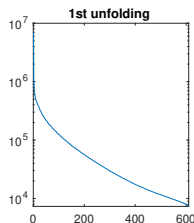
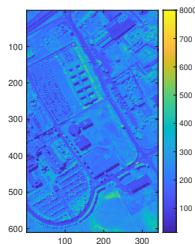


tuple of (different) multilinear ranks $(R_1, \dots, R_d)$, $R_k = \text{rank}(\mathbf{Y}^{(k)})$

SVDs of the unfoldings

$\mathcal{Y} \in \mathbb{R}^{610 \times 340 \times 103}$,
hyperspectral image

singular values of $\mathbf{Y}^{(1)}, \mathbf{Y}^{(2)}, \mathbf{Y}^{(3)}$:



Pavia University hyperspectral dataset

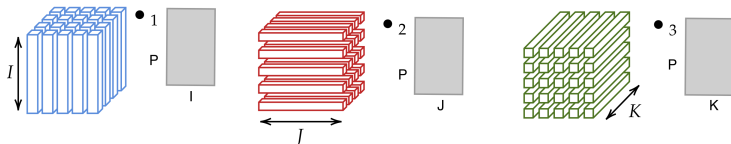
- ▶ singular values: not necessarily same distribution
- ▶ interconnected and have well-behaved geometry
[Hackbusch, Kressner, Uschmajew, 2017],[Krämer, 2019]

Towards factorization: tensor/matrix product

k -th mode contraction :

► with $\mathbf{M} \in \mathbb{F}^{m \times n_k}$: $(\mathcal{Y} \bullet_k \mathbf{M})_{i_1 \dots i_d} \stackrel{\text{def}}{=} \sum_{j=1}^{n_k} \mathcal{Y}_{i_1 \dots i_{k-1} j i_{k+1} \dots i_d} M_{i_k, j}$

For $\mathcal{Y} \in \mathbb{R}^{I \times J \times K}$



For matrices ($d = 2$):

$$\mathbf{Y} \bullet_1 \mathbf{M} = \mathbf{M} \mathbf{Y}, \quad \mathbf{Y} \bullet_2 \mathbf{M} = \mathbf{Y} \mathbf{M}^T$$

For tensors:

$$(\mathcal{Y} \bullet_k \mathbf{M})^{(k)} = \mathbf{M} \mathbf{Y}^{(k)}$$

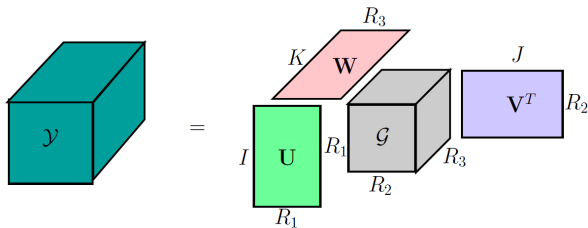
multiplication of k -th unfolding on the left

Multilinear rank and factorization

For $\mathcal{Y} \in \mathbb{R}^{I \times J \times K}$ with ML ranks (R_1, R_2, R_3) :

Tucker factorization with factors $\mathcal{G} \in \mathbb{R}^{R_1 \times R_2 \times R_3}$ (*core tensor*),
 $\mathbf{U} \in \mathbb{R}^{I \times R_1}$, $\mathbf{V} \in \mathbb{R}^{J \times R_2}$, $\mathbf{W} \in \mathbb{R}^{K \times R_3}$

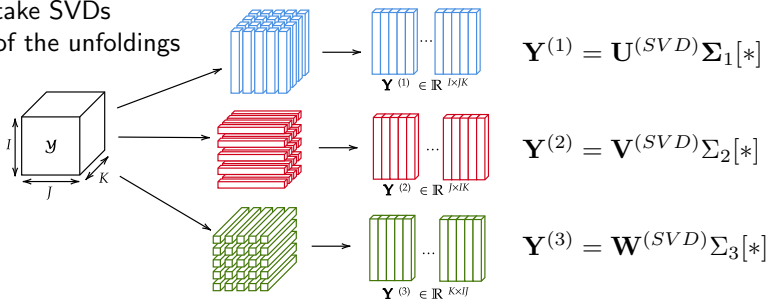
$$\mathcal{Y} = [\mathcal{G}; \mathbf{U}, \mathbf{V}, \mathbf{W}] \stackrel{\text{def}}{=} \mathcal{G} \bullet_1 \mathbf{U} \bullet_2 \mathbf{V} \bullet_3 \mathbf{W}.$$



- ▶ non-unique (as in the matrix case): $\tilde{\mathbf{U}} = \mathbf{U}\mathbf{Q}$
- ▶ in general (random $\mathcal{Y}$),
 $(R_1, R_2, R_3) = (\min(I, JK), \min(J, IK), \min(K, IJ))$

Higher-order SVD

- take SVDs of the unfoldings



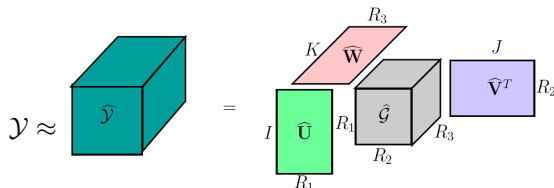
- Compute $\mathcal{G}^{(SVD)} = \mathcal{Y} \bullet_1 (\mathbf{U}^{(SVD)})^\top \bullet_2 (\mathbf{V}^{(SVD)})^\top \bullet_3 (\mathbf{W}^{(SVD)})^\top$

HOSVD

$$\mathcal{Y} = \mathcal{G}^{(SVD)} \bullet_1 \mathbf{U}^{(SVD)} \bullet_2 \mathbf{V}^{(SVD)} \bullet_3 \mathbf{W}^{(SVD)}$$

Best low multilinear approximation

Compute best (R_1, R_2, R_3) -Tucker approximation:



$$\min_{\hat{\mathcal{G}}, \hat{\mathbf{U}}, \hat{\mathbf{V}}, \hat{\mathbf{W}}} \|\mathcal{Y} - \hat{\mathcal{G}} \bullet_1 \hat{\mathbf{U}} \bullet_2 \hat{\mathbf{V}} \bullet_3 \hat{\mathbf{W}}\|_F$$

- ▶ a good (suboptimal) solution is given by truncating **HOSVD**:

$$\hat{\mathbf{U}} = \mathbf{U}_{:,1:R_1}^{(SVD)}, \hat{\mathbf{V}} = \mathbf{V}_{:,1:R_2}^{(SVD)}, \hat{\mathbf{W}} = \mathbf{W}_{:,1:R_3}^{(SVD)}, \hat{\mathcal{G}} = \mathcal{G}_{1:R_1,1:R_2,1:R_3}^{(SVD)}$$

- ▶ HOOI: alternating minimization over $\hat{\mathbf{U}}, \hat{\mathbf{V}}, \hat{\mathbf{W}}$
[De Lathauwer, De Moor, Vandewalle, 2000]
- ▶ optimization over the (R_1, R_2, R_3) -rank manifold
[Kressner, Steinlechner, Vandereycken, 2013], [Kasai, Mishra, 2016]

Tucker approximation: very useful for compression, completion tasks

Overview

Matrix LRA

Tensors

Factorization

(CP) decomposition

CP approximation

Extensions

Tensor rank and CPD

Rank-1 tensor: $\mathcal{T} = \mathbf{a} \otimes \mathbf{b} \otimes \mathbf{c}$: $\mathcal{T}_{ijk} = a_i b_j c_k$

(Canonical) **Polyadic Decomposition** — sum of R rank-one tensors:

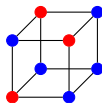
$$\mathcal{T} = \sum_{k=1}^R \mathbf{a}_k \otimes \mathbf{b}_k \otimes \mathbf{c}_k, \quad \boxed{\mathcal{T}} = \mathbf{a}_1 \left| \begin{array}{c} \mathbf{c}_1 / \\ \mathbf{b}_1 \end{array} \right. + \cdots + \mathbf{a}_R \left| \begin{array}{c} \mathbf{c}_R / \\ \mathbf{b}_R \end{array} \right.$$

(CP) **tensor rank**: $\text{rank}(\mathcal{T}) \stackrel{\text{def}}{=} \text{minimal such } R$

Earlier names: **CANDECOMP/PARAFAC**

Example:

$$\mathcal{T} = \mathbf{e}_1 \otimes \mathbf{e}_1 \otimes \mathbf{e}_2 + \mathbf{e}_1 \otimes \mathbf{e}_2 \otimes \mathbf{e}_1 + \mathbf{e}_2 \otimes \mathbf{e}_1 \otimes \mathbf{e}_1$$

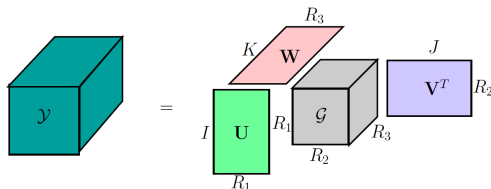


$$\mathcal{T}_{:, :, 1} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \mathcal{T}_{:, :, 2} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}$$

Tensor rank: basic properties

$$\boxed{\mathcal{T}} = \mathbf{a}_1 \left| \overline{\mathbf{b}_1} \right. + \cdots + \mathbf{a}_R \left| \overline{\mathbf{b}_R} \right.$$

- Relation with multilinear ranks:



$$\max(R_1, R_2, R_3) \leq R \leq \min(R_2 R_3, R_1 R_3, R_1 R_2)$$

- NP-hard (to compute exact rank): [Hillar, Lim, 2013]
- But has many nice properties and applications

Simultaneous matrix factorization and CPD

$$\boxed{\mathbf{M}} = \begin{array}{c} | \\ \mathbf{a}_1 \\ | \end{array} \begin{array}{c} \text{---} \mathbf{b}_1^T \text{---} \\ \\ \end{array} + \dots + \begin{array}{c} | \\ \mathbf{a}_R \\ | \end{array} \begin{array}{c} \text{---} \mathbf{b}_R^T \text{---} \\ \\ \end{array}$$

Simultaneous matrix factorization and CPD

$$\begin{array}{rcl}
 \boxed{\mathbf{M}_1} & = & \overset{c_{1,1}}{\mathbf{a}_1} \mathbf{b}_1^\top + \dots + \overset{c_{R,1}}{\mathbf{a}_R} \mathbf{b}_R^\top \\
 \vdots & & \\
 \boxed{\mathbf{M}_N} & = & \overset{c_{1,N}}{\mathbf{a}_1} \mathbf{b}_1^\top + \dots + \overset{c_{R,N}}{\mathbf{a}_R} \mathbf{b}_R^\top
 \end{array}$$

Simultaneous matrix factorization and CPD

$$\begin{array}{l}
 \boxed{\mathbf{M}_1} = \begin{array}{c} c_{1,1} \text{---} \\ | \\ \mathbf{a}_1 \end{array} \text{---} \mathbf{b}_1^\top \text{---} + \dots + \begin{array}{c} c_{R,1} \text{---} \\ | \\ \mathbf{a}_R \end{array} \text{---} \mathbf{b}_R^\top \text{---} \\
 \vdots \\
 \boxed{\mathbf{M}_N} = \begin{array}{c} c_{1,N} \text{---} \\ | \\ \mathbf{a}_1 \end{array} \text{---} \mathbf{b}_1^\top \text{---} + \dots + \begin{array}{c} c_{R,N} \text{---} \\ | \\ \mathbf{a}_R \end{array} \text{---} \mathbf{b}_R^\top \text{---} \\
 \Updownarrow \\
 \begin{array}{c} \dots \\ \boxed{\mathbf{M}_2} \\ \boxed{\mathbf{M}_1} \end{array} = \begin{array}{c} \mathbf{c}_1 \diagdown \\ \mathbf{a}_1 \text{---} \mathbf{b}_1 \end{array} + \dots + \begin{array}{c} \mathbf{c}_R \diagdown \\ \mathbf{a}_R \text{---} \mathbf{b}_R \end{array}
 \end{array}$$

Simultaneous matrix factorization and CPD

fluorescence spectroscopy
emission/excitation matrix

rank = # components in a mixture
 N experiments, different concentrations

$$\begin{array}{c}
 \begin{array}{ccc}
 \begin{array}{c} \text{Heatmap 1} \\ \vdots \\ \text{Heatmap N} \end{array} & = & \begin{array}{c} \mathbf{M}_1 \\ \vdots \\ \mathbf{M}_N \end{array} \\
 & & \begin{array}{c} \begin{array}{c} \mathbf{c}_{1,1} \vdots \\ \mathbf{a}_1 \end{array} \mathbf{b}_1^T + \dots + \begin{array}{c} \mathbf{c}_{R,1} \vdots \\ \mathbf{a}_R \end{array} \mathbf{b}_R^T \\
 & & \begin{array}{c} \mathbf{c}_{1,N} \vdots \\ \mathbf{a}_1 \end{array} \mathbf{b}_1^T + \dots + \begin{array}{c} \mathbf{c}_{R,N} \vdots \\ \mathbf{a}_R \end{array} \mathbf{b}_R^T \\
 \end{array} \\
 \Downarrow \\
 \begin{array}{ccc}
 \begin{array}{c} \mathbf{M}_N \\ \vdots \\ \mathbf{M}_2 \\ \mathbf{M}_1 \end{array} & = & \begin{array}{c} \mathbf{c}_1 \diagdown \\ \mathbf{a}_1 \mid \mathbf{b}_1 \end{array} + \dots + \begin{array}{c} \mathbf{c}_R \diagdown \\ \mathbf{a}_R \mid \mathbf{b}_R \end{array}
 \end{array}
 \end{array}$$

Tensor rank in applications

Application area	tensor rank R
(blind source separation) independent component analysis	# of sources
multiway factor analysis (spectroscopy, chemometrics, ...)	# of components
antenna array processing	# of transmitters
$\vdots$	$\vdots$

$$\begin{array}{c}
 \text{Heatmap 1} \\
 \vdots \\
 \text{Heatmap N}
 \end{array}
 =
 \begin{array}{c}
 \boxed{\mathbf{M}_1} \\
 \vdots \\
 \boxed{\mathbf{M}_N}
 \end{array}
 =
 \begin{array}{c}
 c_{1,1} \begin{array}{c} | \\ \mathbf{a}_1 \end{array} \begin{array}{c} \text{---} \mathbf{b}_1^T \text{---} \\ \text{---} \end{array} \\
 + \dots + \\
 c_{R,1} \begin{array}{c} | \\ \mathbf{a}_R \end{array} \begin{array}{c} \text{---} \mathbf{b}_R^T \text{---} \\ \text{---} \end{array} \\
 \\
 c_{1,N} \begin{array}{c} | \\ \mathbf{a}_1 \end{array} \begin{array}{c} \text{---} \mathbf{b}_1^T \text{---} \\ \text{---} \end{array} \\
 + \dots + \\
 c_{R,N} \begin{array}{c} | \\ \mathbf{a}_R \end{array} \begin{array}{c} \text{---} \mathbf{b}_R^T \text{---} \\ \text{---} \end{array}
 \end{array}$$

Essential uniqueness of a CPD

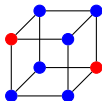
$$\mathcal{T} = \sum_{k=1}^R \mathbf{a}_k \otimes \mathbf{b}_k \otimes \mathbf{c}_k,$$

$$\boxed{\mathcal{T}} = \underset{\mathbf{a}_1}{\text{c}_1 / \overline{\mathbf{b}_1}} + \cdots + \underset{\mathbf{a}_R}{\text{c}_R / \overline{\mathbf{b}_R}}$$

up to **permutations** and **rescaling** ($\mathbf{a}'_k = \alpha \mathbf{a}_k, \mathbf{b}'_k = \beta \mathbf{b}_k, \mathbf{c}'_k = \frac{1}{\alpha\beta} \mathbf{c}_k$)

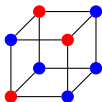
Examples:

- unique: $\mathcal{T} = \mathbf{e}_1 \otimes \mathbf{e}_1 \otimes \mathbf{e}_1 + \mathbf{e}_2 \otimes \mathbf{e}_2 \otimes \mathbf{e}_2$



$$\mathcal{T}_{::,1} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}, \quad \mathcal{T}_{::,2} = \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix},$$

- non-unique: $\mathcal{T} = \mathbf{e}_1 \otimes \mathbf{e}_1 \otimes \mathbf{e}_2 + \mathbf{e}_1 \otimes \mathbf{e}_2 \otimes \mathbf{e}_1 + \mathbf{e}_2 \otimes \mathbf{e}_1 \otimes \mathbf{e}_1$



$$\mathcal{T}_{::,1} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \mathcal{T}_{::,2} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}$$

Third-order tensors: Kruskal sufficient condition

Shorthand notation: $\llbracket \mathbf{A}, \mathbf{B}, \mathbf{C} \rrbracket := \sum_{k=1}^R \mathbf{a}_k \otimes \mathbf{b}_k \otimes \mathbf{c}_k$

factor matrices $\mathbf{A} = [\mathbf{a}_1 \cdots \mathbf{a}_R]$, $\mathbf{B} = [\mathbf{b}_1 \cdots \mathbf{b}_R]$, $\mathbf{C} = [\mathbf{c}_1 \cdots \mathbf{c}_R]$

Theorem [Kruskal, 1978].

The decomposition $\overset{I}{\underset{J}{\text{cube}}} \mathcal{T} = \llbracket \mathbf{A}, \mathbf{B}, \mathbf{C} \rrbracket$ is unique if

$$\text{kr}(\mathbf{A}) + \text{kr}(\mathbf{B}) + \text{kr}(\mathbf{C}) \geq 2r + 2,$$

$\underbrace{\text{kr}(\mathbf{A})}_{\text{Kruskal rank}} \stackrel{\text{def}}{=} \text{maximal } k \text{ such that any } k \text{ columns are linearly independent.}$

Kruskal rank

Example (my favourite tensor):

- ▶ $K = 2$ (2 frontal slices)
- ▶ $\mathbf{A}, \mathbf{B}$ — full column rank, $\mathbf{C}$ — with non-collinear columns

$2R + 2 = 2R + 2 \Rightarrow$ unique decomposition

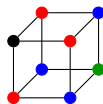
Real vs. complex rank

In general, they are different:

$$\text{rank}_{\mathbb{C}}(\mathcal{T}) \leq \text{rank}_{\mathbb{R}}(\mathcal{T})$$

Example ([Kruskal, 1983], but also [Sylvester, 1851]):

$$\mathcal{T}_{:, :, 1} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \mathcal{T}_{:, :, 2} = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix},$$

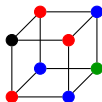


$$\mathcal{T} = \frac{i}{2} \left(\begin{bmatrix} 1 \\ -i \end{bmatrix} \otimes \begin{bmatrix} 1 \\ -i \end{bmatrix} \otimes \begin{bmatrix} 1 \\ -i \end{bmatrix} - \begin{bmatrix} 1 \\ i \end{bmatrix} \otimes \begin{bmatrix} 1 \\ i \end{bmatrix} \otimes \begin{bmatrix} 1 \\ i \end{bmatrix} \right)$$

$$\Rightarrow \text{rank}_{\mathbb{C}}(\mathcal{T}) = \text{rank}_{S, \mathbb{C}}(\mathcal{T}) = 2, \text{ but } \text{rank}_{\mathbb{R}}(\mathcal{T}) = \text{rank}_{S, \mathbb{R}}(\mathcal{T}) = 3$$

Symmetric tensor rank vs. tensor rank

- $S_n^d \stackrel{\text{def}}{=} \text{vector space of symmetric tensors,}$


 S_2^3

Symmetric rank: $\text{rank}_S(\mathcal{T}) \stackrel{\text{def}}{=} \text{minimal } r \text{ such that}$

$$\mathcal{T} = \sum_{k=1}^r \lambda_k \mathbf{a}_k \otimes \cdots \otimes \mathbf{a}_k, \quad \mathbf{a}_k \in \mathbb{F}^n$$

obviously $\text{rank}(\mathcal{T}) \leq \text{rank}_S(\mathcal{T})$

Comon's conjecture: $\forall \mathcal{T} \in S_n^d, \quad \text{rank}(\mathcal{T}) = \text{rank}_S(\mathcal{T})$

- true for
 - matrices $d = 2$
 - ranks smaller than order/dimension [Friedland, 2017], [Zhang, Huang, Qi, 2017]
 - generically for small ranks [Lim, Qi, 2020],
- counterexample by [Shitov, 2017]: $n = 800, d = 3, \text{rank}(\mathcal{T}) = 903$

Maximal and generic ranks



Maximal rank

$r_{max} \stackrel{\text{def}}{=} \text{maximal possible rank}$

- ▶ different for symmetric/non-symmetric
- ▶ different for $\mathbb{R}$ and $\mathbb{C}$

(Very few) known cases

tensor space	dimension	r_{max}	reference
$\mathbb{F}^{n \times n}$	n^2	n	matrices
$\mathbb{C}^{2 \times n \times n}$	$2n^2$	$\lfloor \frac{3n}{2} \rfloor$	[Grigoriev, 1978], [Ja'Ja', 1978]
$\mathbb{F}^{n \times n \times n}$	n^3	$\leq \frac{(n+1)n}{2}$	[Atkinson, Stephens, 1979]

What if we draw tensors randomly?



$$\longrightarrow \text{rank}(\mathcal{T}) = ?$$

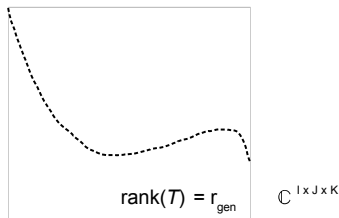
random = from absolutely continuous probability distribution

- ▶ important in practice (noise, numerical errors)
- ▶ often easier to find (than r_{max})

Complex tensors: generic rank

With probability 1 a complex tensor has rank r_{gen} (generic rank)

(i.e. other tensors have measure zero)



- ▶ matrices ($\mathbb{C}^{n \times n}$): $r_{max} = r_{gen} = n$
- ▶ cubic tensors ($\mathbb{C}^{n \times n \times n}$, $n > 3$) [Lickteig, 1985]

$$r_{gen} = \left\lceil \frac{n^3}{3n-2} \right\rceil \approx \frac{n^2}{3}$$

Generic ranks

- symmetric tensors (S_n^d , $d \geq 3$) [Alexander, Hirschowitz, 1996]

$$r_{gen} = \left\lceil \frac{\binom{n+d-1}{n-1}}{n} \right\rceil,$$

with few exceptions: $(d, n) = (3, 5)$ or $d = 4, n \in \{3, 4, 5\}$

- general case ($\mathbb{C}^{I_1 \times \dots \times I_d}$, $d \geq 3$):

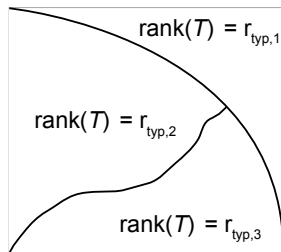
$$r_{gen} \stackrel{?}{=} \left\lceil \frac{I_1 \cdots I_d}{I_1 + \dots + I_d - d + 1} \right\rceil, \quad \text{with few exceptions}$$

- conjectured [Abo, Ottaviani, Peterson, 2009]
- computer proof [Chiantini, Ottaviani, Vannieuwenhoven, 2014]
- [Blekherman, Teitler, 2014]: $r_{max} \leq 2r_{gen}$

Proofs: algebraic geometry (dimensions of secant varieties)

Real tensors: typical ranks

For real tensors, several **typical** ranks may appear with nonzero probability.



Example [Bergqvist, 2013]: $\mathcal{T} \in \mathbb{R}^{2 \times 2 \times 2}$ with i.i.d. Gaussian elements has:

- ▶ rank 2 with probability $\frac{\pi}{4}$
- ▶ rank 3 with probability $1 - \frac{\pi}{4}$

Some numerical consequences

1. noise, numerical errors $\Rightarrow \text{rank}(\mathcal{T}) = r_{gen}$ (or a typical rank in $\mathbb{R}$)
2. **very difficult** to find tensors with **higher ranks**:

If we generate

$$\boxed{\mathcal{T}} = \underset{\mathbf{a}_1}{\text{c}_1 /} \overline{\mathbf{b}_1} + \cdots + \underset{\mathbf{a}_r}{\text{c}_r /} \overline{\mathbf{b}_r}$$

with $\mathbf{a}_k, \mathbf{b}_k, \mathbf{c}_k$ random, has

$$\text{rank}(\mathcal{T}) = \begin{cases} r, & r \leq r_{gen}, \\ r_{gen}, & r_{gen} < r \leq r_{max}. \end{cases}$$

Example. $n \times n \times n$ tensors with $\text{rank} \left[\frac{n^3}{3n-2} \right] < r \leq \frac{(n+1)n}{2}$.

Generic uniqueness (identifiability)

For a fixed rank r : whether “almost all” decompositions are unique

- ▶ Kruskal-type conditions give weak bounds
- ▶ Study the properties secant algebraic variety ($\sigma_r \stackrel{\text{def}}{=} \text{(Zariski) closure of tensors of rank } r$)
- ▶ Generic uniqueness: uniqueness for all tensors in σ_r except a set of Lebesgue measure 0

Most recent results:

1. [Chiantini, Ottaviani, 2012]: CPD of $\mathcal{T} \in \mathbb{C}^{I \times J \times K}$, with $I \geq J \geq K$ is generically unique if $r \leq 2^{\lfloor \log_2 J \rfloor + \lfloor \log_2 K \rfloor - 2}$.
2. [Chiantini, Ottaviani, Vannieuwenhoven, 2014] (computer proof):
complex identifiability holds for all **subgeneric ranks**

$$r < r_{\text{gen}} = \left\lceil \frac{IJK}{I+J+K-2} \right\rceil$$
 no identifiability for $r > r_{\text{gen}}$
3. [Qi, Comon, Lim, 2016], [Chiantini, Ottaviani, Vannieuwenhoven, 2017]:

all identifiability results are valid for **real-valued** tensors

Overview

Matrix LRA

Tensors

Factorization

(CP) decomposition

CP approximation

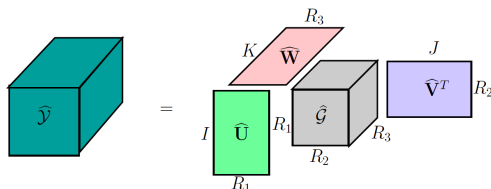
Extensions

CP vs. Tucker

- CPD: $(I + J + K - 2)R$

$$\boxed{\mathcal{T}} \approx \underset{\mathbf{a}_1}{\boxed{\phantom{\mathcal{T}}}} \bigg|_{\mathbf{b}_1}^{\mathbf{c}_1 /} + \cdots + \underset{\mathbf{a}_R}{\boxed{\phantom{\mathcal{T}}}} \bigg|_{\mathbf{b}_R}^{\mathbf{c}_R /}$$

- Tucker: $IR_1 + JR_2 + KR_3 + R_1R_2R_3 - R_1^2 - R_2^2 - R_3^2$

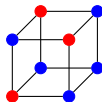


CP approximation is ill-posed

$$\boxed{\mathcal{T}} \approx \mathbf{a}_1 \left| \frac{\mathbf{c}_1}{\mathbf{b}_1} \right| + \cdots + \mathbf{a}_R \left| \frac{\mathbf{c}_R}{\mathbf{b}_R} \right|$$

Best r -rank approximation ($r > 1$) may not exist

$$\mathcal{T} = \mathbf{e}_1 \otimes \mathbf{e}_1 \otimes \mathbf{e}_2 + \mathbf{e}_1 \otimes \mathbf{e}_2 \otimes \mathbf{e}_1 + \mathbf{e}_2 \otimes \mathbf{e}_1 \otimes \mathbf{e}_1$$



$$\mathcal{T}_{:, :, 1} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \mathcal{T}_{:, :, 2} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}$$

$\text{rank } \mathcal{T} = 3$, but approximated by rank-2 tensor to any accuracy

$$\mathcal{T} = \frac{1}{2\varepsilon}(\mathbf{e}_1 + \varepsilon\mathbf{e}_2)^{\otimes 3} - \frac{1}{2\varepsilon}(\mathbf{e}_1 - \varepsilon\mathbf{e}_2)^{\otimes 3} + \mathcal{O}(\varepsilon)$$

Set of rank- $\leq r$ tensors is not closed for $r > 1$

Rank-one tensor approximation

$$\min_{\mathbf{a}, \mathbf{b}, \mathbf{c}} \|\mathcal{T} - \mathbf{a} \otimes \mathbf{b} \otimes \mathbf{c}\|_F^2$$

- ▶ well-posed (minimum exists)
- ▶ block-coordinate descent (ALS, non-symmetric power method) converges globally (to a stationary point) [Uschmajew, 2015]
- ▶ related to the notion of singular vectors/eigenvectors of a tensor [Qi, Luo, 2017]
- ▶ number of stationary points is known [Freidland, Ottaviani, 2014]

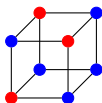
Successive approximation (deflation)

Subtracting rank-one approximation may increase tensor rank:

$$\mathcal{X}_{:, :, 1} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \mathcal{X}_{:, :, 2} = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix}$$

$\text{rank}(\mathcal{X}) = 2$, but $\mathcal{X} - \hat{\mathcal{X}}_1 = \mathcal{T}$, $\text{rank}(\mathcal{T}) = 3$

$$\mathcal{T} = \mathbf{e}_1 \otimes \mathbf{e}_1 \otimes \mathbf{e}_2 + \mathbf{e}_1 \otimes \mathbf{e}_2 \otimes \mathbf{e}_1 + \mathbf{e}_2 \otimes \mathbf{e}_1 \otimes \mathbf{e}_1$$



$$\mathcal{T}_{:, :, 1} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \mathcal{T}_{:, :, 2} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}$$

But when does it work then?

Successive rank-1 approximation

- Orthogonally decomposable tensors [Zhang, Golub, 2001]:

$$\mathcal{T} = \sum_{k=1}^R \mathbf{a}_k \otimes \mathbf{b}_k \otimes \mathbf{c}_k, \quad \mathbf{a}_k \perp \mathbf{a}_j, \quad \mathbf{b}_k \perp \mathbf{b}_j, \quad \mathbf{c}_k \perp \mathbf{c}_j,$$

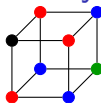
The successive best rank-1 approximation returns the components in the sum

- Cyclic (sequential) rank-one approximation [da Silva, Comon, de Almeida, 2015]:

$$\min \|\mathcal{T} - (\hat{\mathcal{T}}_1 + \dots + \hat{\mathcal{T}}_R)\| \quad \text{subject to } \text{rank}(\hat{\mathcal{T}}_k) = 1.$$

Rank-one approximation of symmetric tensors

- $\mathcal{T} \in S_n^d$ — symmetric tensors,


 S_2^3

- Best non-symmetric approximation can be chosen symmetric (and is unique a.s.) [Friedland, Ottaviani, 2014]
- Best symmetric approximation $\leftrightarrow$ maximization of a polynomial

$$\min_{\lambda, \mathbf{v}} \|\mathcal{T} - \lambda \mathbf{v} \otimes \cdots \otimes \mathbf{v}\|_F^2 \leftrightarrow \max_{\|\mathbf{v}\|_2=1} |\mathcal{T} \bullet_1 \mathbf{v} \cdots \bullet_d \mathbf{v}|$$

- stationary points: eigenvectors of the tensor

$$\mathcal{T} \cdot \mathbf{v}^{d-1} = \mu \mathbf{v}$$

- In total, $\frac{(d-1)^n - 1}{d-2}$ (complex) eigenvectors [Cartwright, Sturmfels, 2013]
- There are cases when the power method diverges [Chen, Saad, 2009]
- Orthogonally decomposable tensors: power method converges, deflation works decomposition [Anandkumar et al, 2013], [Robeva, 2016]

Algorithms for CP approximation

$$\min_{\mathbf{A}, \mathbf{B}, \mathbf{C}} \|\mathcal{Y} - \llbracket \mathbf{A}, \mathbf{B}, \mathbf{C} \rrbracket\|_F^2 + \dots$$

Constraints/regularization:

- ▶ Orthogonality: [Comon, 1993], [Robeva, 2016]
- ▶ Coherence: [Lim, Comon, 2014]
- ▶ Nonnegativity: [Qi, Comon, Lim, 2016]

Algorithms for CP approximation

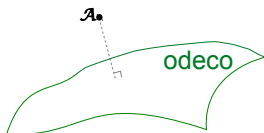
- ▶ Alternating minimization (ALS, AO-ADMM)
(Global) convergence properties in the regularized case [Xu, Yin, 2013]
- ▶ Nonlinear least squares
- ▶ Riemannian optimization
- ▶ Algebraic algorithms
 - ▶ Generalized eigenvalue decomposition (non-symmetric tensors)
 - ▶ Structured matrix approximation (symmetric tensors)

Tensor diagonalization as orthogonal approximation

By duality:

$$\begin{aligned} & \max_{\mathbf{Q} \in \mathcal{O}_n} \|\text{diag } \mathcal{A} \bullet_1 \mathbf{Q}^\top \cdots \bullet_d \mathbf{Q}^\top\|_2^2 \\ &= \|\mathcal{A}\|_F^2 - \underbrace{\min_{\substack{\mathbf{Q} \in \mathcal{O}_n \\ \mathbf{Q} = [\mathbf{u}_1 \cdots \mathbf{u}_n]}} \left\| \mathcal{A} - \sum_{k=1}^n \mu_k \mathbf{u}_k \otimes \cdots \otimes \mathbf{u}_k \right\|_F^2}_{\text{best } n\text{-rank symmetric orthogonal approximation}} \end{aligned}$$

- ▶ Euclidean distance to odeco [Robeva, 2016] variety



- ▶ Best non-symmetric approximation $\neq$ best symmetric approximation [Li, U., Comon, 2019] (a variant of Comon's conjecture is false)

Optimization on orthogonal/unitary group

$$\max_{\mathbf{Q} \in \mathcal{O}_n} f(\mathbf{Q}) \quad \text{or} \quad \max_{\mathbf{U} \in \mathcal{U}_n} f(\mathbf{U}), \quad f \text{ is a low-order polynomial}$$

Algebraic orthogonal tensor decompositions:

- ▶ Deflation (successive rank-one approximation)
[Delfosse, Loubaton, 1995], [Anandkumar, 2013], [Robeva, 2016]
- ▶ EVD of tensor slice(s) [De Lathauwer, 2006], [Kolda, 2015]

Optimization on the manifold:

- ▶ Riemannian optimization (CG, SD, BFGS, RTR) [Absil et al., 2008]
- ▶ Jacobi-type algorithms [Comon, 1993], [Li, U., Comon, 2020]
(convergence)

What do we know about CP approximation?

Given $\mathcal{T} = \llbracket \mathbf{A}, \mathbf{B}, \mathbf{C} \rrbracket + \mathcal{E}$ and

$$(\hat{\mathbf{A}}^*, \hat{\mathbf{B}}^*, \hat{\mathbf{C}}^*) = \arg \min_{\hat{\mathbf{A}}, \hat{\mathbf{B}}, \hat{\mathbf{C}}} \|\mathcal{T} - \llbracket \mathbf{A}, \mathbf{B}, \mathbf{C} \rrbracket\|,$$

- ▶ Approaches using random matrix theory [Goulart, Couillet, Comon, 2022] (mostly rank-1)
- ▶ Perturbation bounds for some algebraic algorithms [Evert, De Lathauwer, 2022] (very small ranks)
- ▶ Uniqueness of best approximation in a small neighborhood [Friedland, Stawiska, 2016] (non-constructive)

Overview

Matrix LRA

Tensors

Factorization

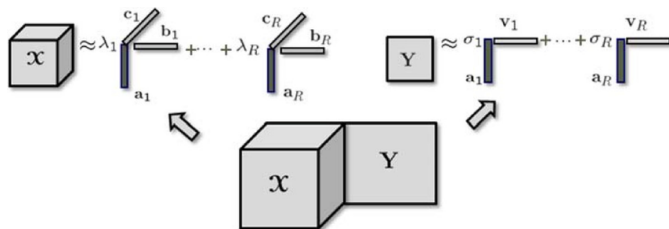
(CP) decomposition

CP approximation

Extensions

Coupled factorizations

Factors of different tensor/matrix decompositions may be shared:



[Acar, Kolda, Dunlavy, 2011]

Joint CPD of symmetric tensors

ICA model: $\mathbf{x} = A\mathbf{s}$, $A = [\mathbf{a}_1 \ \cdots \ \mathbf{a}_r] \in \mathbb{R}^{n \times r}$,

cumulants of $\mathbf{x}$ up to order d :

$$\begin{aligned} \mathcal{C}_{\mathbf{x}}^{(1)} &= c_{1,1}\mathbf{a}_1 + \cdots + c_{1,r}\mathbf{a}_r, \\ \mathcal{C}_{\mathbf{x}}^{(2)} &= c_{2,1}\mathbf{a}_1 \otimes \mathbf{a}_1 + \cdots + c_{2,r}\mathbf{a}_r \otimes \mathbf{a}_r, \\ &\vdots \\ \mathcal{C}_{\mathbf{x}}^{(d)} &= c_{d,1}\mathbf{a}_1 \otimes \cdots \otimes \mathbf{a}_1 + \cdots + c_{d,r}\mathbf{a}_r \otimes \cdots \otimes \mathbf{a}_r, \end{aligned} \tag{1}$$

where $c_{j,k}$ is the j -th cumulant of s_k .

Problem. Given $\mathcal{C}_{\mathbf{x}}^{(j)}$, find $\mathbf{a}_k$ and $c_{j,k}$
(diagonalize all the cumulants simultaneously)

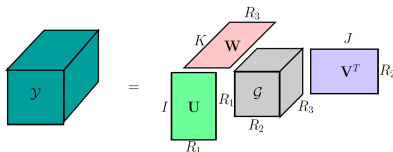
Sum of Tucker: block-term decomposition

[De Lathauwer, 2008]:

For fixed (R_1, R_2, R_3) :

$$\mathcal{T} = \sum_{k=1}^r \mathcal{G}_k \bullet_1 \mathbf{U}_k \bullet_2 \mathbf{V}_k \bullet_3 \mathbf{W}_k$$

each term in the sum has ML-rank (R_1, R_2, R_3) :



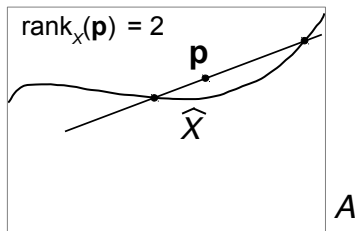
Special case: $(R_1, R_2, R_3) = (L, L, 1)$: very useful in signal processing, see e.g. [Goulart, et al. 2020]

Additive decompositions

X-rank (join) decomposition [Zak, 2004],
[Landsberg, 2012], [Comon, Qi, U., 2017]
(sparse algebraic decomposition)

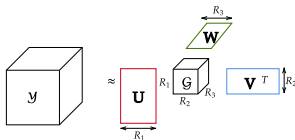
$$\mathbf{p} = \mathbf{x}_1 + \cdots + \mathbf{x}_r, \quad \mathbf{x}_k \in \hat{X}$$

- ▶ $A \stackrel{\text{def}}{=} \text{ambient tensor space (e.g. } A = \mathbb{C}^{I \times J \times K})$
- ▶ $\hat{X} \stackrel{\text{def}}{=} \text{(variety of "simple" terms)}$
(e.g., rank-one $\hat{X} = \{\mathbf{a} \otimes \mathbf{b} \otimes \mathbf{c}\}$)
- ▶ study the properties of secant varieties $\sigma_r(\hat{X})$

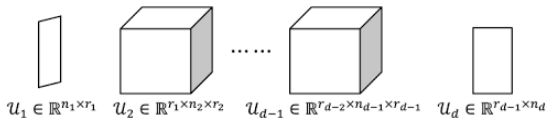


Approximation: higher-order tensors

- Tucker: $O(dIR + R^d)$: curse of dimensionality



- tensor trains, hierarchical Tucker: linear-in- d storage complexity



Decompositions: flexible/multilayer

[Harshman, Lundy, 1996], [Roald et al, 2022]

► CPD (PARAFAC)

$$\mathcal{T}_{::,k} = \underset{I}{\boxed{\mathbf{A}}}^R \cdot \boxed{\mathbf{D}_G^{(k)}} \cdot \underset{R}{\boxed{\mathbf{B}^\top}}^J, \quad k = 1, \dots, K$$

► PARAFAC-2

$$\mathcal{T}_{::,k} = \underset{I}{\boxed{\mathbf{A}}}^R \cdot \boxed{\mathbf{D}_G^{(k)}} \cdot \underset{R}{\boxed{\mathbf{B}_k^\top}}^J, \quad k = 1, \dots, K$$

► ParaTuck-2

$$\mathcal{T}_{::,k} = \underset{I}{\boxed{\mathbf{A}}}^R \cdot \boxed{\mathbf{D}_G^{(k)}} \cdot \underset{R}{\boxed{\mathbf{F}}}^S \cdot \boxed{\mathbf{D}_H^{(k)}} \cdot \underset{S}{\boxed{\mathbf{B}^\top}}^J, \quad k = 1, \dots, K$$

Share the uniqueness features of CPD!

Conclusion

Which tensor decomposition (format) to use?

- ▶ factorization (Tucker, tensor trains) — compression
- ▶ decomposition (CPD, BTD) — identification of components

Challenges (personal choice):

- ▶ Guarantees (approximation bounds) on CP approximation
- ▶ Multi-layer/multilevel tensor decompositions (e.g., ParaTuck-2)

Thank you!

Conclusion

Which tensor decomposition (format) to use?

- ▶ factorization (Tucker, tensor trains) — compression
- ▶ decomposition (CPD, BTD) — identification of components

Challenges (personal choice):

- ▶ Guarantees (approximation bounds) on CP approximation
- ▶ Multi-layer/multilevel tensor decompositions (e.g., ParaTuck-2)

Advertisement (not mine):

- ▶ A number of postdoc/PhD positions in the SiMul team (Nancy)
<https://cran-simul.github.io>
- ▶ Summer school on low-rank approximation (23-29.06.24) in Peyresq
Organized by N. Gillis and J. Cohen

Thank you!