

Mathematics of Data: From Theory to Computation

Prof. Volkan Cevher
volkan.cevher@epfl.ch

Supplementary Material: Linear Algebra

Laboratory for Information and Inference Systems (LIONS)
École Polytechnique Fédérale de Lausanne (EPFL)

EE-556 (Fall 2024)



License Information for Mathematics of Data Slides

- ▶ This work is released under a [Creative Commons License](#) with the following terms:
- ▶ **Attribution**
 - ▶ The licensor permits others to copy, distribute, display, and perform the work. In return, licensees must give the original authors credit.
- ▶ **Non-Commercial**
 - ▶ The licensor permits others to copy, distribute, display, and perform the work. In return, licensees may not use the work for commercial purposes – unless they get the licensor's permission.
- ▶ **Share Alike**
 - ▶ The licensor permits others to distribute derivative works only under a license identical to the one that governs the licensor's work.
- ▶ [Full Text of the License](#)

Outline

- ▶ Review of linear algebra

1. Vectors
2. Matrices
3. *Tensors

*: PhD material

Preliminaries

- We use the following standard notation in the Mathematics of Data lectures
 - ▶ **Scalars** are denoted by lowercase letters (e.g., k)
 - ▶ **Vectors** by lowercase boldface letters (e.g., $\mathbf{x}$)
 - ▶ **Matrices** by uppercase boldface letters (e.g., $\mathbf{A}$)
 - ▶ **Component** of a *vector* $\mathbf{x}$, *matrix* $\mathbf{A}$ as x_i , a_{ij} & $A_{i,j,k}, \dots$ respectively
 - ▶ **Sets** by uppercase calligraphic letters (e.g., $\mathcal{S}$)
- We focus on the **field of real** numbers ($\mathbb{R}$)
- Most results here can be **generalized** to the **field of complex** numbers ($\mathbb{C}$)

Vectors

1. Vector spaces
2. Vector norms
3. Inner products
4. Dual norms

Vectors

Definition

A vector is an array of numbers arranged by rows or columns.

Vector spaces

Definition

A vector space or *linear space* over the field $\mathbb{R}$ consists of

- (a) a **set** of vectors $\mathcal{V}$
- (b) an **addition** operation: $\mathcal{V} \times \mathcal{V} \rightarrow \mathcal{V}$
- (c) a **scalar multiplication** operation: $\mathbb{R} \times \mathcal{V} \rightarrow \mathcal{V}$
- (d) a **distinguished** element $\mathbf{0} \in \mathcal{V}$

and satisfies the following properties:

- | | |
|--|--|
| 1. $\mathbf{x} + \mathbf{y} = \mathbf{y} + \mathbf{x}, \quad \forall \mathbf{x}, \mathbf{y} \in \mathcal{V}$ | <i>commutative under addition</i> |
| 2. $(\mathbf{x} + \mathbf{y}) + \mathbf{z} = \mathbf{x} + (\mathbf{y} + \mathbf{z}), \forall \mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathcal{V}$ | <i>associative under addition</i> |
| 3. $\mathbf{0} + \mathbf{x} = \mathbf{x}, \forall \mathbf{x} \in \mathcal{V}$ | <i>$\mathbf{0}$ being additive identity</i> |
| 4. $\forall \mathbf{x} \in \mathcal{V} \quad \exists (-\mathbf{x}) \in \mathcal{V}$ such that $\mathbf{x} + (-\mathbf{x}) = \mathbf{0}$ | <i>$-\mathbf{x}$ being additive inverse</i> |
| 5. $(\alpha\beta)\mathbf{x} = \alpha(\beta\mathbf{x}), \quad \forall \alpha, \beta \in \mathbb{R} \quad \forall \mathbf{x} \in \mathcal{V}$ | <i>associative under scalar multiplication</i> |
| 6. $\alpha(\mathbf{x} + \mathbf{y}) = \alpha\mathbf{x} + \alpha\mathbf{y}, \quad \forall \alpha \in \mathbb{R} \quad \forall \mathbf{x}, \mathbf{y} \in \mathcal{V}$ | <i>distributive</i> |
| 7. $1\mathbf{x} = \mathbf{x}, \forall \mathbf{x} \in \mathcal{V}$ | <i>1 being multiplicative identity</i> |

Vector spaces contd.

Example (Vector space)

1. $\mathcal{V}_1 = \{\mathbf{0}\}$ for $\mathbf{0} \in \mathbb{R}^p$
2. $\mathcal{V}_2 = \mathbb{R}^p$
3. $\mathcal{V}_3 = \sum_{i=1}^k \alpha_i \mathbf{x}_i$ for $\alpha_i \in \mathbb{R}$ and $\mathbf{x}_i \in \mathbb{R}^p$

It is straight forward to show that $\mathcal{V}_1$, $\mathcal{V}_2$, and $\mathcal{V}_3$ satisfy properties 1–7 shown before.

Definition (Subspace)

A **subspace** is a vector space that is a *subset* of another vector space.

Example (Subspace)

$\mathcal{V}_1$, $\mathcal{V}_2$, and $\mathcal{V}_3$ in the example above are subspaces of $\mathbb{R}^p$.

Vector spaces contd.

Definition (Span)

The **span** of a set of vectors, $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k\}$, is the set of all possible **linear combinations** of these vectors; i.e.,

$$\text{span}\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k\} = \{\alpha_1 \mathbf{x}_1 + \alpha_2 \mathbf{x}_2 + \dots + \alpha_k \mathbf{x}_k \mid \alpha_1, \alpha_2, \dots, \alpha_k \in \mathbb{R}\}.$$

Definition (Linear independence)

A set of vectors, $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k\}$, is **linearly independent** if

$$\alpha_1 \mathbf{x}_1 + \alpha_2 \mathbf{x}_2 + \dots + \alpha_k \mathbf{x}_k = \mathbf{0} \Rightarrow \alpha_1 = \alpha_2 = \dots = \alpha_k = 0.$$

Definition (Basis)

The **basis** of a vector space, $\mathcal{V}$, is a set of vectors $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k\}$ that satisfy

(a) $\mathcal{V} = \text{span}\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k\}$, (b) $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k\}$ are linearly independent.

Definition (Dimension)

The **dimension** of a vector space, $\mathcal{V}$, (denoted $\text{dim}(\mathcal{V})$) is the number of vectors in the basis of $\mathcal{V}$.

Vector norms

Definition (Vector norm)

A norm of a vector in $\mathbb{R}^p$ is a function $\|\cdot\| : \mathbb{R}^p \rightarrow \mathbb{R}$ such that for all vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^p$ and scalar $\lambda \in \mathbb{R}$

- (a) $\|\mathbf{x}\| \geq 0$ for all $\mathbf{x} \in \mathbb{R}^p$ *nonnegativity*
- (b) $\|\mathbf{x}\| = 0$ if and only if $\mathbf{x} = \mathbf{0}$ *definitiveness*
- (c) $\|\lambda\mathbf{x}\| = |\lambda|\|\mathbf{x}\|$ *homogeneity*
- (d) $\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|$ *triangle inequality*

The ℓ_q -norms

For $\mathbf{x} \in \mathbb{R}^p$, the ℓ_q -norm is defined as $\|\mathbf{x}\|_q := \left(\sum_{i=1}^p |x_i|^q\right)^{1/q}$ for $q \in [1, \infty]$.

Example

- (1) ℓ_2 -norm: $\|\mathbf{x}\|_2 := \sqrt{\sum_{i=1}^p x_i^2}$ (Euclidean norm)
- (2) ℓ_1 -norm: $\|\mathbf{x}\|_1 := \sum_{i=1}^p |x_i|$ (Manhattan norm)
- (3) ℓ_∞ -norm: $\|\mathbf{x}\|_\infty := \max_{i=1, \dots, p} |x_i|$ (Chebyshev norm)

Vector norms contd.

Definition (Quasi-norm)

A **quasi-norm** satisfies all the norm properties except (d) triangle inequality, which is replaced by $\|\mathbf{x} + \mathbf{y}\| \leq c(\|\mathbf{x}\| + \|\mathbf{y}\|)$ for a constant $c \geq 1$.

Definition (Semi(pseudo)-norm)

A **semi(pseudo)-norm** satisfies all the norm properties except (b) definiteness.

Example

- ▶ The ℓ_q -norm is in fact a quasi norm when $q \in (0, 1)$, with $c = 2^{1/q} - 1$.
- ▶ The **total variation norm** (TV-norm) defined (in 1D): $\|\mathbf{x}\|_{\text{TV}} := \sum_{i=1}^{p-1} |x_{i+1} - x_i|$ is a **semi-norm** since it fails to satisfy (b);
e.g., any $\mathbf{x} = c(1, 1, \dots, 1)^T$ for $c \neq 0$ will have $\|\mathbf{x}\|_{\text{TV}} = 0$ even though $\mathbf{x} \neq \mathbf{0}$.

Definition (ℓ_0 -“norm”)

$$\|\mathbf{x}\|_0 = \lim_{q \rightarrow 0} \|\mathbf{x}\|_q^q = |\{i : x_i \neq 0\}|$$

- Observations:**
- The ℓ_0 -“norm” counts the non-zero components of $\mathbf{x}$. Hence, it is **not** a norm.
 - It does not satisfy the property (c) $\Rightarrow$ it is also neither a **quasi-** nor a **semi-norm**.

Vector norms contd.

Problem (s -sparse approximation)

Find $\arg \min_{\mathbf{x} \in \mathbb{R}^p} \|\mathbf{x} - \mathbf{y}\|_2$ subject to: $\|\mathbf{x}\|_0 \leq s$.

Vector norms contd.

Problem (s -sparse approximation)

Find $\arg \min_{\mathbf{x} \in \mathbb{R}^p} \|\mathbf{x} - \mathbf{y}\|_2$ subject to: $\|\mathbf{x}\|_0 \leq s$.

Solution

Define $\hat{\mathbf{y}} \in \arg \min_{\mathbf{x} \in \mathbb{R}^p: \|\mathbf{x}\|_0 \leq s} \|\mathbf{x} - \mathbf{y}\|_2^2$ and let $\hat{\mathcal{S}} = \text{supp}(\hat{\mathbf{y}})$.

We now consider an optimization over sets

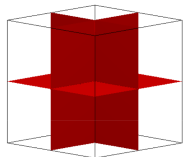
$$\begin{aligned} \hat{\mathcal{S}} &\in \arg \min_{\mathcal{S}: |\mathcal{S}| \leq s} \|\mathbf{y}_{\mathcal{S}} - \mathbf{y}\|_2^2. \\ &\in \arg \max_{\mathcal{S}: |\mathcal{S}| \leq s} \left\{ \|\mathbf{y}\|_2^2 - \|\mathbf{y}_{\mathcal{S}} - \mathbf{y}\|_2^2 \right\} \\ &\in \arg \max_{\mathcal{S}: |\mathcal{S}| \leq s} \left\{ \|\mathbf{y}_{\mathcal{S}}\|_2^2 \right\} = \arg \max_{\mathcal{S}: |\mathcal{S}| \leq s} \sum_{i \in \mathcal{S}} \|y_i\|^2 \quad (\equiv \text{modular approximation problem}). \end{aligned}$$

Thus, the **best s -sparse approximation** of a vector is a vector with the s **largest components** of the vector in **magnitude**.

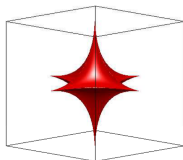
Vector norms contd.

Norm balls

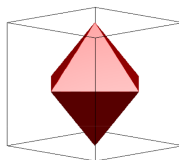
Radius r ball in ℓ_q -norm: $\mathcal{B}_q(r) = \{\mathbf{x} \in \mathbb{R}^p : \|\mathbf{x}\|_q \leq r\}$



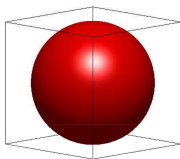
$\|\mathbf{x}\|_0 \leq 2$



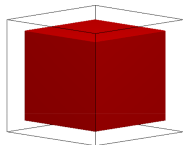
$\ell_{0.5}$ -quasi norm ball



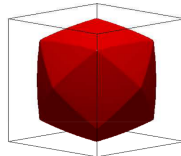
ℓ_1 -norm ball



ℓ_2 -norm ball



ℓ_∞ -norm ball



TV-semi norm ball

Table: Example norm balls in $\mathbb{R}^3$

Inner products

Definition (Inner product)

The **inner product** of any two vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^p$ (denoted by $\langle \cdot, \cdot \rangle$) is defined as $\langle \mathbf{x}, \mathbf{y} \rangle = \mathbf{x}^T \mathbf{y} = \sum_i^p x_i y_i$.

The inner product satisfies the following properties:

1. $\langle \mathbf{x}, \mathbf{y} \rangle = \langle \mathbf{y}, \mathbf{x} \rangle, \forall \mathbf{x}, \mathbf{y} \in \mathbb{R}^p$ *symmetry*
2. $\langle (\alpha \mathbf{x} + \beta \mathbf{y}), \mathbf{z} \rangle = \langle \alpha \mathbf{x}, \mathbf{z} \rangle + \langle \beta \mathbf{y}, \mathbf{z} \rangle, \forall \alpha, \beta \in \mathbb{R}, \forall \mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbb{R}^p$ *linearity*
3. $\langle \mathbf{x}, \mathbf{x} \rangle \geq 0, \forall \mathbf{x} \in \mathbb{R}^p$ *positive definiteness*

Important relations involving the inner product:

- ▶ **Hölder's inequality:** $|\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\|_q \|\mathbf{y}\|_r$, where $r > 1$ and $\frac{1}{q} + \frac{1}{r} = 1$
- ▶ **Cauchy-Schwarz** is a special case of Hölder's inequality ($q = r = 2$)

Definition (Inner product space)

An **inner product space** is a **vector space** endowed with an **inner product**.

Vector norms contd.

Definition (Dual norm)

Let $\|\cdot\|$ be a norm in $\mathbb{R}^p$, then the **dual norm** denoted by $\|\cdot\|^*$ is defined:

$$\|\mathbf{x}\|^* = \sup_{\|\mathbf{y}\| \leq 1} \mathbf{x}^T \mathbf{y}, \quad \text{for all } \mathbf{x}, \mathbf{y} \in \mathbb{R}^p$$

- ▶ The **dual** of the *dual norm* is the **original (primal) norm**, i.e., $\|\mathbf{x}\|^{**} = \|\mathbf{x}\|$.
- ▶ Hölder's inequality $\Rightarrow \|\cdot\|_q$ is a **dual norm** of $\|\cdot\|_r$ when $\frac{1}{q} + \frac{1}{r} = 1$.

Example 1

- i) $\|\cdot\|_2$ is **dual** of $\|\cdot\|_2$ (i.e. $\|\cdot\|_2$ is *self-dual*): $\sup\{\mathbf{z}^T \mathbf{x} \mid \|\mathbf{x}\|_2 \leq 1\} = \|\mathbf{z}\|_2$.
- ii) $\|\cdot\|_1$ is **dual** of $\|\cdot\|_\infty$, (and *vice versa*): $\sup\{\mathbf{z}^T \mathbf{x} \mid \|\mathbf{x}\|_\infty \leq 1\} = \|\mathbf{z}\|_1$.

Example 2

What is the **dual norm** of $\|\cdot\|_q$ for $q = 1 + 1/\log(p)$ for $p > 1$?

Vector norms contd.

Definition (Dual norm)

Let $\|\cdot\|$ be a norm in $\mathbb{R}^p$, then the **dual norm** denoted by $\|\cdot\|^*$ is defined:

$$\|\mathbf{x}\|^* = \sup_{\|\mathbf{y}\| \leq 1} \mathbf{x}^T \mathbf{y}, \quad \text{for all } \mathbf{x}, \mathbf{y} \in \mathbb{R}^p$$

- ▶ The **dual** of the *dual norm* is the **original (primal) norm**, i.e., $\|\mathbf{x}\|^{**} = \|\mathbf{x}\|$.
- ▶ Hölder's inequality $\Rightarrow \|\cdot\|_q$ is a **dual norm** of $\|\cdot\|_r$ when $\frac{1}{q} + \frac{1}{r} = 1$.

Example 1

- i) $\|\cdot\|_2$ is **dual** of $\|\cdot\|_2$ (i.e. $\|\cdot\|_2$ is *self-dual*): $\sup\{\mathbf{z}^T \mathbf{x} \mid \|\mathbf{x}\|_2 \leq 1\} = \|\mathbf{z}\|_2$.
- ii) $\|\cdot\|_1$ is **dual** of $\|\cdot\|_\infty$, (and *vice versa*): $\sup\{\mathbf{z}^T \mathbf{x} \mid \|\mathbf{x}\|_\infty \leq 1\} = \|\mathbf{z}\|_1$.

Example 2

What is the **dual norm** of $\|\cdot\|_q$ for $q = 1 + 1/\log(p)$ for $p > 1$?

Solution

By Hölder's inequality, $\|\cdot\|_r$ is the **dual norm** of $\|\cdot\|_q$ if $\frac{1}{q} + \frac{1}{r} = 1$. Therefore, $r = 1 + \log(p)$ is the dual.

Metrics

- A **metric** on a set is a function that satisfies the minimal properties of a distance.

Definition (Metric)

Let $\mathcal{X}$ be a set, then a function $d(\cdot, \cdot) : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ is a metric if $\forall \mathbf{x}, \mathbf{y} \in \mathcal{X}$:

- (a) $d(\mathbf{x}, \mathbf{y}) \geq 0$ for all $\mathbf{x}$ and $\mathbf{y}$ (*nonnegativity*)
- (b) $d(\mathbf{x}, \mathbf{y}) = 0$ if and only if $\mathbf{x} = \mathbf{y}$ (*definiteness*)
- (c) $d(\mathbf{x}, \mathbf{y}) = d(\mathbf{y}, \mathbf{x})$ (*symmetry*)
- (d) $d(\mathbf{x}, \mathbf{y}) \leq d(\mathbf{x}, \mathbf{z}) + d(\mathbf{z}, \mathbf{y})$ (*triangle inequality*)

Observations:

- A **pseudo-metric** satisfies (a), (c) and (d) but not necessarily (b)
- A **metric space** $(\mathcal{X}, d)$ is a set $\mathcal{X}$ with a metric d defined on $\mathcal{X}$
- **Norms** induce **metrics** while **pseudo-norms** induce **pseudo-metrics**

Example

- ▶ Euclidean distance: $d_E(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|_2$
- ▶ Bregman distance: $d_B(\cdot, \cdot)$ cf. Lecture 3

Matrices

1. Special matrix types
2. Basic matrix definitions
3. Matrix decompositions
4. Complexity of matrix operations
5. Matrix norms

Matrices

Definition

A matrix is a rectangular array of numbers arranged by rows and columns.

- In the sequel, we describe a set of **special matrices** to get started.

Special matrices

Definition (Identity matrix)

The *identity* matrix (denoted $\mathbf{I} \in \mathbb{R}^{p \times p}$) is a **square** matrix of zero entries except on the *main diagonal*, which has ones on it. For compatible matrices $\mathbf{A}$ and $\mathbf{B}$, it satisfies:

$$\mathbf{IA} = \mathbf{A} \text{ and } \mathbf{BI} = \mathbf{B}.$$

Definition (Orthogonal (or Unitary) matrix)

A matrix $\mathbf{A} \in \mathbb{R}^{p \times p}$ is **orthogonal** or **unitary** if $\mathbf{A}^T \mathbf{A} = \mathbf{A} \mathbf{A}^T = \mathbf{I}$.

Definition (Triangular matrix)

A matrix $\mathbf{A} \in \mathbb{R}^{p \times p}$ is **lower triangular** if all its entries above the *main diagonal* are zero, i.e., $a_{ij} = 0$ for $j > i$; while it is **upper triangular** if $\mathbf{A}^T$ is lower triangular.

Definition (Permutation matrix)

A matrix $\mathbf{P} \in \mathbb{R}^{n \times p}$ is **permutation** if it has only one 1 in each row and each column and satisfies $\mathbf{PP}^T = \mathbf{I}$.

Special matrices contd.

Definition (Incidence matrix)

An incidence matrix shows the relationship between two sets $\mathcal{X}$ and $\mathcal{Y}$. The i -th row corresponding to entry $x_i \in \mathcal{X}$ and the j -th column corresponding to entry $y_j \in \mathcal{Y}$ of an incidence matrix is 1 if x_i and y_j are related and 0 if they are not.

Definition (Adjacency matrix)

An adjacency matrix is a symmetric square matrix with $\{0, 1\}$ entries where 1 or 0 at the (i, j) -th location indicates the i -th and the j -th vertices of a graph are adjacent (i.e., share an edge) or not.

- The diagonal entries of adjacency matrices take different values depending on different conventions.

Definition (Stochastic matrix)

A matrix $\mathbf{P} \in \mathbb{R}^{n \times p}$ is **stochastic** (also known as **transition** or **probability**) matrix if $\sum_j p_{ij} = 1$ for $0 \leq p_{ij} \leq 1$; while $\mathbf{A}$ is **doubly stochastic** if $\sum_i p_{ij} = \sum_j p_{ij} = 1$.

Definition (Gaussian matrix)

A matrix $\mathbf{A} \in \mathbb{R}^{p \times p}$ is **Gaussian** if its entries $a_{lk} \sim \mathcal{N}(\mu, \sigma^2)$ for $l, k \in [p]$. That is, its entries are independent and identically distributed (i.i.d.) with mean μ & variance σ^2 according to the Gaussian distribution.

Special matrices contd.

Definition (Fourier matrix)

A matrix $\mathbf{F} \in \mathbb{C}^{p \times p}$ is **Fourier matrix** if its entries

$$f_{lk} = \frac{1}{\sqrt{p}} e^{i2\pi lk/p}, \quad \text{for } l, k \in [p], \quad i = \sqrt{-1}.$$

Definition (Discrete Cosine Transform matrix)

A matrix $\mathbf{A} \in \mathbb{R}^{p \times p}$ is **Discrete Cosine Transform (DCT) matrix** if its entries

$$a_{lk} = \sqrt{\frac{2}{p}} \cos \left(\frac{\pi}{p} (l-1) \left(k - \frac{1}{2} \right) \right); 1 \leq l \leq p, 1 \leq k \leq p.$$

- ▶ The **Fourier and DCT matrices** are both **orthogonal**, i.e., $\mathbf{F}^H \mathbf{F} = \mathbf{F} \mathbf{F}^H = \mathbf{I}$, where $\mathbf{F}^H = \text{complex-conjugate}(\mathbf{F}^T)$.
- ▶ Both matrices are rarely stored since they have an implicit **fast matrix-vector multiplication algorithm**.

Special matrices contd.

Definition (Hadamard matrix [4])

Let the indices $l, k \in [2^n]$ be defined as $l = \sum_{j=1}^n l_j 2^{j-1} + 1$, $k = \sum_{j=1}^n k_j 2^{j-1} + 1$. A matrix

$\mathbf{H} = \mathbf{H}_n \in \mathbb{R}^{2^n \times 2^n}$ is a **Hadamard matrix** (or **Hadamard transform**) if

$$h_{lk} = \frac{1}{2^{n/2}} (-1)^{\sum_{j=1}^n k_j l_j}.$$

- ▶ The **Hadamard matrix** is **orthogonal** and **self-adjoint**, i.e., $\mathbf{H}_n = \mathbf{H}_n^T$.
- ▶ The **Hadamard matrix** is rarely stored since it has a **fast matrix-vector multiplication algorithm** that uses the **recursive identity**:

$$\mathbf{H}_n = \frac{1}{\sqrt{2}} \begin{pmatrix} \mathbf{H}_{n-1} & \mathbf{H}_{n-1} \\ \mathbf{H}_{n-1} & -\mathbf{H}_{n-1} \end{pmatrix}, \quad \mathbf{H}_0 = 1.$$

Special matrices contd.

Definition (Toeplitz matrix [2])

Let a $\mathbf{t} = (t_1, t_2, \dots, t_{2p-1})$ be fixed or drawn from a probability distribution $\mathcal{P}(\mathbf{t})$. Then $\mathbf{T} \in \mathbb{R}^{p \times p}$ is **Toeplitz matrix** if

$$\mathbf{T} = \begin{pmatrix} t_1 & t_2 & t_3 & \cdots & t_{p-1} & t_p \\ t_{p+1} & t_1 & t_2 & \cdots & t_{p-2} & t_{p-1} \\ t_{p+2} & t_{p+1} & t_1 & \cdots & t_{p-3} & t_{p-2} \\ \vdots & \vdots & \vdots & \ddots & \ddots & \vdots \\ t_{2p-2} & t_{2p-3} & \cdots & \cdots & t_1 & t_2 \\ t_{2p-1} & t_{2p-2} & t_{2p-3} & \cdots & t_{p+1} & t_1 \end{pmatrix}.$$

Definition (Circulant matrix [8])

Let a $\mathbf{c} = (c_1, c_2, \dots, c_p)$ be fixed or drawn from a probability distribution $\mathcal{P}(\mathbf{c})$, then $\mathbf{C} \in \mathbb{R}^{p \times p}$ is **Circulant matrix** if

$$\mathbf{C} = \begin{pmatrix} c_1 & c_p & \cdots & c_3 & c_2 \\ c_2 & c_1 & \cdots & c_4 & c_3 \\ \vdots & \vdots & \ddots & \ddots & \vdots \\ c_p & c_{p-1} & \cdots & c_2 & c_1 \end{pmatrix}.$$

Special matrices contd.

Partial Fourier, Partial Toeplitz, Partial Circulant, ...

A **partial** Fourier, Toeplitz or Circulant matrix refers to a matrix consisting of a **subset of the rows** of a Fourier, Toeplitz or Circulant matrix, respectively.

- ▶ Fourier, Hadamard, Toeplitz and Circulant matrices are **structured** matrices. In addition, Toeplitz and Circulant matrices are **banded**.
- ▶ These matrices also have lower degrees-of-freedom as compared to a general matrix in $\mathbb{R}^{p \times p}$. Hence, computations revolving around these matrices are typically cheaper than the computation we need for a general matrix.
- ▶ Incident and adjacency matrices are often used in graph theory. They have important decompositional and computational properties.

Basic matrix definitions

Definition (Nullspace of a matrix)

The **nullspace** of a matrix, $\mathbf{A} \in \mathbb{R}^{n \times p}$, (denoted by $\text{null}(\mathbf{A})$) is defined as

$$\text{null}(\mathbf{A}) = \{\mathbf{x} \in \mathbb{R}^p \mid \mathbf{A}\mathbf{x} = \mathbf{0}\}$$

- ▶ $\text{null}(\mathbf{A})$ is the set of vectors mapped to **zero** by $\mathbf{A}$.
- ▶ $\text{null}(\mathbf{A})$ is the set of vectors **orthogonal** to the rows of $\mathbf{A}$.

Definition (Range of a matrix)

The **range** of a matrix, $\mathbf{A} \in \mathbb{R}^{n \times p}$, (denoted by $\text{range}(\mathbf{A})$) is defined as

$$\text{range}(\mathbf{A}) = \{\mathbf{A}\mathbf{x} \mid \mathbf{x} \in \mathbb{R}^p\} \subseteq \mathbb{R}^n$$

- ▶ $\text{range}(\mathbf{A})$ is the **span** of the columns (or the **column space**) of $\mathbf{A}$.

Definition (Rank of a matrix)

The **rank** of a matrix, $\mathbf{A} \in \mathbb{R}^{n \times p}$, (denoted by $\text{rank}(\mathbf{A})$) is defined as

$$\text{rank}(\mathbf{A}) = \text{dim}(\text{range}(\mathbf{A}))$$

- ▶ $\text{rank}(\mathbf{A})$ is the maximum number of **independent** columns (or rows) of $\mathbf{A}$, $\Rightarrow \text{rank}(\mathbf{A}) \leq \min(n, p)$.
- ▶ $\text{rank}(\mathbf{A}) = \text{rank}(\mathbf{A}^T)$; **and** $\text{rank}(\mathbf{A}) + \text{dim}(\text{null}(\mathbf{A})) = p$.

Matrix definitions contd.

Definition (Eigenvalues & Eigenvectors)

The vector $\mathbf{x}$ is an **eigenvector** of a *square* matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ if $\mathbf{Ax} = \lambda\mathbf{x}$ where $\lambda \in \mathbb{R}$ is called an **eigenvalue** of $\mathbf{A}$.

- ▶ $\mathbf{A}$ scales its eigenvectors by its eigenvalues.

Definition (Singular values & singular vectors)

For $\mathbf{A} \in \mathbb{R}^{n \times p}$ and *unit* vectors $\mathbf{u} \in \mathbb{R}^n$ and $\mathbf{v} \in \mathbb{R}^p$ if

$$\mathbf{Av} = \sigma\mathbf{u} \quad \text{and} \quad \mathbf{A}^T\mathbf{u} = \sigma\mathbf{v}$$

then $\sigma \in \mathbb{R}$ ($\sigma \geq 0$) is a **singular value** of $\mathbf{A}$; $\mathbf{v}$ and $\mathbf{u}$ are the **right singular vector** and the **left singular vector** respectively of $\mathbf{A}$.

Definition (Symmetric matrix)

A matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ is **symmetric** if $\mathbf{A} = \mathbf{A}^T$.

Lemma

The eigenvalues of a symmetric $\mathbf{A}$ are real.

Proof.

Assume $\mathbf{Ax} = \lambda\mathbf{x}$, $\mathbf{x} \in \mathbb{C}^p$, $\mathbf{x} \neq \mathbf{0}$, then $\bar{\mathbf{x}}^T\mathbf{Ax} = \bar{\mathbf{x}}^T(\mathbf{Ax}) = \bar{\mathbf{x}}^T(\lambda\mathbf{x}) = \lambda \sum_{i=1}^n |x_i|^2$

but $\bar{\mathbf{x}}^T\mathbf{Ax} = \overline{(\mathbf{Ax})}^T\mathbf{x} = \overline{(\lambda\mathbf{x})}^T\mathbf{x} = \bar{\lambda} \sum_{i=1}^n |x_i|^2 \Rightarrow \lambda = \bar{\lambda}$ i.e. $\lambda \in \mathbb{R}$

□

Matrix definitions contd.

Definition (Positive semidefinite & positive definite matrices)

A **symmetric** matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ is **positive semidefinite** (denoted $\mathbf{A} \succeq 0$) if $\mathbf{x}^T \mathbf{A} \mathbf{x} \geq 0$ for all $\mathbf{x} \neq 0$; while it is **positive definite** (denoted $\mathbf{A} \succ 0$) if $\mathbf{x}^T \mathbf{A} \mathbf{x} > 0$

- ▶ $\mathbf{A} \succeq 0$ iff all its **eigenvalues** are **nonnegative** i.e. $\lambda_{\min}(\mathbf{A}) \geq 0$.
- ▶ Similarly, $\mathbf{A} \succ 0$ iff all its **eigenvalues** are **positive** i.e. $\lambda_{\min}(\mathbf{A}) > 0$.
- ▶ $\mathbf{A}$ is **negative semidefinite** if $-\mathbf{A} \succeq 0$; while $\mathbf{A}$ is **negative definite** if $-\mathbf{A} \succ 0$.
- ▶ **Semidefinite ordering** of two *symmetric* matrices, $\mathbf{A}$ and $\mathbf{B}$: $\mathbf{A} \succeq \mathbf{B}$ if $\mathbf{A} - \mathbf{B} \succeq 0$.

Example (Matrix inequalities)

1. If $\mathbf{A} \succeq 0$ and $\mathbf{B} \succeq 0$, then $\mathbf{A} + \mathbf{B} \succeq 0$
2. If $\mathbf{A} \succeq \mathbf{B}$ and $\mathbf{C} \succeq \mathbf{D}$, then $\mathbf{A} + \mathbf{C} \succeq \mathbf{B} + \mathbf{D}$
3. If $\mathbf{B} \preceq 0$ then $\mathbf{A} + \mathbf{B} \preceq \mathbf{A}$
4. If $\mathbf{A} \succeq 0$ and $\alpha \geq 0$, then $\alpha \mathbf{A} \succeq 0$
5. If $\mathbf{A} \succ 0$, then $\mathbf{A}^2 \succ 0$
6. If $\mathbf{A} \succ 0$, then $\mathbf{A}^{-1} \succ 0$

Matrix decompositions

Definition (Eigenvalue decomposition)

The **eigenvalue decomposition** of a **square** matrix, $\mathbf{A} \in \mathbb{R}^{n \times n}$, is given by:

$$\mathbf{A} = \mathbf{X}\mathbf{\Lambda}\mathbf{X}^{-1}$$

- ▶ the columns of $\mathbf{X} \in \mathbb{R}^{n \times n}$, i.e. $\mathbf{x}_i$, are **eigenvectors** of $\mathbf{A}$
- ▶ $\mathbf{\Lambda} = \mathbf{diag}(\lambda_1, \lambda_2, \dots, \lambda_n)$ where λ_i (also denoted $\lambda_i(\mathbf{A})$) are **eigenvalues** of $\mathbf{A}$
- ▶ A matrix that admits this decomposition is therefore called **diagonalizable** matrix

Eigendecomposition of symmetric matrices

If $\mathbf{A} \in \mathbb{R}^{n \times n}$ is **symmetric**, the decomposition becomes $\mathbf{A} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^T$
where $\mathbf{U} \in \mathbb{R}^{n \times n}$ is **unitary** (or **orthonormal**), i.e. $\mathbf{U}^T\mathbf{U} = \mathbf{I}$ and λ_i are **real**

If we order $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$, $\lambda_i(\mathbf{A})$ becomes the i^{th} largest eigenvalue of $\mathbf{A}$.

Definition (Determinant of a matrix)

The **determinant** of a **square** matrix $\mathbf{A} \in \mathbb{R}^{p \times p}$, denoted by $\det(\mathbf{A})$, is given by:

$$\det(\mathbf{A}) = \prod_{i=1}^p \lambda_i$$

where λ_i are *eigenvalues* of $\mathbf{A}$.

Matrix decompositions contd

Definition (Singular value decomposition)

The **singular value decomposition** (SVD) of a matrix, $\mathbf{A} \in \mathbb{R}^{n \times p}$, is given by:

$$\mathbf{A} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T$$

- ▶ $\text{rank}(\mathbf{A}) = r \leq \min(n, p)$ and σ_i is the i^{th} **singular value** of $\mathbf{A}$
 - ▶ $\mathbf{u}_i$ and $\mathbf{v}_i$ are the i^{th} **left** and **right singular vectors** of $\mathbf{A}$ respectively
 - ▶ $\mathbf{U} \in \mathbb{R}^{n \times r}$ and $\mathbf{V} \in \mathbb{R}^{p \times r}$ are **unitary** matrices (i.e. $\mathbf{U}^T \mathbf{U} = \mathbf{I}$)
 - ▶ $\mathbf{\Sigma} = \text{diag}(\sigma_1, \sigma_2, \dots, \sigma_r)$ where $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r \geq 0$
-
- ▶ $\mathbf{v}_i$ are **eigenvectors** of $\mathbf{A}^T \mathbf{A}$; $\sigma_i = \sqrt{\lambda_i(\mathbf{A}^T \mathbf{A})}$ (and $\lambda_i(\mathbf{A}^T \mathbf{A}) = 0$ for $i > r$) since $\mathbf{A}^T \mathbf{A} = (\mathbf{U}\mathbf{\Sigma}\mathbf{V}^T)^T (\mathbf{U}\mathbf{\Sigma}\mathbf{V}^T) = (\mathbf{V}\mathbf{\Sigma}^2\mathbf{V}^T)$
 - ▶ $\mathbf{u}_i$ are **eigenvectors** of $\mathbf{A}\mathbf{A}^T$; $\sigma_i = \sqrt{\lambda_i(\mathbf{A}\mathbf{A}^T)}$ (and $\lambda_i(\mathbf{A}\mathbf{A}^T) = 0$ for $i > r$) since $\mathbf{A}\mathbf{A}^T = (\mathbf{U}\mathbf{\Sigma}\mathbf{V}^T) (\mathbf{U}\mathbf{\Sigma}\mathbf{V}^T)^T = (\mathbf{U}\mathbf{\Sigma}^2\mathbf{U}^T)$

Matrix decompositions contd

Definition (LU)

The **LU factorization** of a **nonsingular square** matrix, $\mathbf{A} \in \mathbb{R}^{p \times p}$, is given by:

$$\mathbf{A} = \mathbf{P}\mathbf{L}\mathbf{U}$$

where $\mathbf{P}$ is a **permutation matrix**¹, $\mathbf{L}$ is **lower triangular** and $\mathbf{U}$ is **upper triangular**.

Definition (QR)

The **QR factorization** of any matrix, $\mathbf{A} \in \mathbb{R}^{n \times p}$, is given by:

$$\mathbf{A} = \mathbf{Q}\mathbf{R}$$

where $\mathbf{Q} \in \mathbb{R}^{n \times n}$ is an **orthonormal** matrix, i.e. $\mathbf{Q}^T \mathbf{Q} = \mathbf{I}$, and $\mathbf{R} \in \mathbb{R}^{n \times p}$ is **upper triangular**.

Definition (Cholesky)

The **Cholesky factorization** of a **positive definite and symmetric** matrix, $\mathbf{A} \in \mathbb{R}^{p \times p}$, is given by:

$$\mathbf{A} = \mathbf{L}\mathbf{L}^T$$

where $\mathbf{L}$ is a **lower triangular** matrix with **positive** entries on the *diagonal*.

¹ A matrix $\mathbf{P} \in \mathbb{R}^{p \times p}$ is **permutation** if it has only one 1 in each row and each column.

Complexity of matrix operations

Definition (floating-point operation)

A **floating-point operation** (flop) is one addition, subtraction, multiplication, or division of two floating-point numbers.

Table: Complexity examples: vector are in $\mathbb{R}^p$, matrices in $\mathbb{R}^{n \times p}$, $\mathbb{R}^{p \times m}$ or $\mathbb{R}^{p \times p}$ [5]

Operation	Complexity	Remarks
vector addition	p flops	
vector inner product	$2p - 1$ flops	or $\approx 2p$ for p large
matrix-vector product	$n(2p - 1)$ flops	or $\approx 2np$ for p large $2m$ if $\mathbf{A}$ is sparse with m nonzeros
matrix-matrix product	$mn(2p - 1)$ flops	or $\approx 2mnp$ for p large much less if $\mathbf{A}$ is sparse ¹
LU decomposition	$\frac{2}{3}p^3 + 2p^2$ flops	or $\frac{2}{3}p^3$ for p large much less if $\mathbf{A}$ is sparse ¹
Cholesky decomposition	$\frac{1}{3}p^3 + 2p^2$ flops	or $\frac{1}{3}p^3$ for p large much less if $\mathbf{A}$ is sparse ¹
SVD	$C_1 n^2 p + C_2 p^3$ flops	$C_1 = 4$, $C_2 = 22$ for R-SVD algo.
Determinant	complexity of SVD	

¹ Complexity depends on p , no. of nonzeros in $\mathbf{A}$ and the sparsity pattern.

Computing eigenvalues and eigenvectors

- There are various algorithms to compute eigenpairs of matrices [10].
- One can choose an algorithm depending on the setting.
- Difference considerations include computational complexity, number of eigenvalues or eigenvectors needed.

Power Method

Starting with an initial vector $\mathbf{x}^0$, $\mathbf{x}^{k+1} = \frac{\mathbf{A}\mathbf{x}^k}{\|\mathbf{A}\mathbf{x}^k\|_2}$ converges to the leading eigenvector of the matrix $\mathbf{A}$ under certain conditions. Moreover, $\lambda^k = \frac{\mathbf{x}^{k*} \mathbf{A} \mathbf{x}^k}{\mathbf{x}^{k*} \mathbf{x}^k}$ converges to the leading eigenvalue, i.e., the one with largest absolute value.

- Power method only uses matrix-vector multiplications and normalizations.
- Useful when $\mathbf{A}$ is a large matrix with sparse entries as it does not require singular value matrix decomposition
- Applied in the original PageRank algorithm of Google.

Inverse Power Method

Knowing an upper bound α on the largest eigenvalue of A , apply power method to $\mathbf{A} - \alpha\mathbf{I}$, i.e., iterate $\mathbf{x}^{k+1} = \frac{(\mathbf{A} - \alpha\mathbf{I})\mathbf{x}^k}{\|(\mathbf{A} - \alpha\mathbf{I})\mathbf{x}^k\|_2}$. Then, $\lambda^k = \frac{\mathbf{x}^{k*} \mathbf{A} \mathbf{x}^k}{\mathbf{x}^{k*} \mathbf{x}^k}$ converges to the smallest eigenvalue of $\mathbf{A}$.

Computing eigenvalues and eigenvectors

Shifted Power Method

A variant of the power method is the shifted power method. Here, we choose a scalar s and apply the power method to $\mathbf{A} - s\mathbf{I}$. The parameter s shifts the eigenvalue λ of $\mathbf{A}$ to $\lambda - s$ of $\mathbf{A} - s\mathbf{I}$. If the shift is chosen appropriately, the algorithm can converge faster to the leading eigenvalue.

Remark: ○ Large-scale problems need newer methods that control storage in addition to arithmetic costs.

Randomized Shifted Power Method

When the storage is the overriding concern, we can run the shifted power method with a random starting vector.

Costs (Randomized shifted power method for symmetric matrices)

Let $\mathbf{M} \in \mathbb{S}_n$ a symmetric matrix. For each $\epsilon \in (0, 1]$ and $\delta \in (0, 1]$, the shifted power method computes a unit vector $\mathbf{u} \in \mathbb{F}^n$ that satisfies:

$$\mathbf{u}^* \mathbf{M} \mathbf{u} \leq \lambda_{\min}(\mathbf{M}) + \epsilon \|\mathbf{M}\| \text{ with probability at least } 1 - \delta$$

after $q \geq \frac{1}{2} + \epsilon^{-1} \log(n/\delta^2)$ iterations. The arithmetic cost is $\mathcal{O}(q)$ matrix-vector multiplies with $\mathbf{M}$ and $\mathcal{O}(qn)$ extra operations. The working storage is about $2n$ numbers.

Computing eigenvalues and eigenvectors

- In case storage is not the issue, we can use the randomized Lanczos method.

Lanczos algorithm goal

Given a symmetric matrix $\mathbf{A} \in \mathbb{S}_n$ with eigenvalues $\lambda_1 > \lambda_2 > \dots > \lambda_n$, and the associated eigenvectors $\mathbf{u}_1, \dots, \mathbf{u}_n$, Lanczos algorithms finds approximations for the k largest eigenvalues of $\mathbf{A}$ and its associated eigenvectors, where $k \ll n$.

Randomized Lanczos algorithm

- Select a random vector $\mathbf{v}$, construct a Krylov subspace, $\mathcal{K}(\mathbf{A}, \mathbf{v}, k) = \text{span}\{\mathbf{v}, \mathbf{A}\mathbf{v}, \mathbf{A}^2\mathbf{v}, \dots, \mathbf{A}^{k-1}\mathbf{v}\}$.
- Project $\mathbf{A}$ into this Krylov subspace, $\mathbf{T} = \text{proj}_{\mathcal{K}} \mathbf{A}$
- Use the eigenvalues and vectors of $\mathbf{T}$ as approximations to those of $\mathbf{A}$.

In matrix form:

$$\mathbf{K}_k = [\mathbf{v} \ \mathbf{A}\mathbf{v} \ \dots \ \mathbf{A}^{k-1}\mathbf{v}] \in \mathbb{R}^{n \times k}$$

$$\mathbf{Q}_k = [\mathbf{q}_1 \ \mathbf{q}_2 \ \dots \ \mathbf{q}_k] \leftarrow \text{qr}(\mathbf{K}_k)$$

$$\mathbf{T}_k = \mathbf{Q}_k^T \mathbf{A} \mathbf{Q}_k \in \mathbb{R}^{k \times k}$$

Computing eigenvalues and eigenvectors

Costs (Randomized Lanczos algorithm)

Let $M \in \mathbb{S}_n$. For $\epsilon \in (0, 1]$ and $\delta \in (0, 0.5]$, the randomized Lanczos method computes a unit vector $u \in \mathbb{F}^n$ that satisfies:

$$u^* M u \leq \lambda_{\min}(M) + \frac{\epsilon}{8} \|M\| \text{ with probability at least } 1 - 2\delta$$

after $q \geq \frac{1}{2} + \epsilon^{-1/2} \log(n/\delta^2)$ iterations. The arithmetic cost is at most q matrix-vector multiplies with M and $\mathcal{O}(qn)$ extra operations. The working storage is $\mathcal{O}(qn)$.

Storage Optimal (double loop) Randomized Lanczos [12]

More efficiently one can run the for loop in the Lanczos algorithm twice:

- ▶ The first time to find the weight vector u .
- ▶ The second to generate the eigenvector v using the weights found in u .

This method focuses on regenerating the Lanczos vectors instead of storing them to construct the approximate eigenvector. It takes the weighted average in the end.

With this we can change the storing cost to $\mathcal{O}(q + n)$ instead of $\mathcal{O}(qn)$ as previously stated.

Linear operators

- Matrices are often given in an **implicit** form.
- It is convenient to think of them as *linear operators*.

Proposition (Linear operators & matrices)

Any **linear operator** in finite dimensional spaces can be represented as a **matrix**.

Example

Given matrices $\mathbf{A}$, $\mathbf{B}$ and $\mathbf{X}$ with compatible dimensions and the *linear operator* $\mathcal{M} : \mathbb{R}^{n \times p} \rightarrow \mathbb{R}^{np}$, a linear operator can define the following implicit mapping

$$\mathcal{M}(\mathbf{X}) := (\mathbf{B}^T \otimes \mathbf{A}) \text{vec}(\mathbf{X}) = \text{vec}(\mathbf{A}\mathbf{X}\mathbf{B}),$$

where $\otimes$ is the Kronecker product and $\text{vec} : \mathbb{R}^{n \times p} \rightarrow \mathbb{R}^{np}$ is yet another linear operator that vectorizes its entries.

Note: Clearly, it is more efficient to compute $\text{vec}(\mathbf{A}\mathbf{X}\mathbf{B})$ than to perform the *matrix multiplication* $(\mathbf{B}^T \otimes \mathbf{A}) \text{vec}(\mathbf{X})$.

Matrix norms

Similar to [vector norms](#), **matrix norms** are a [metric](#) over matrices:

Definition (Matrix norm)

A norm of an $n \times p$ matrix is a map $\|\cdot\| : \mathbb{R}^{n \times p} \rightarrow \mathbb{R}$ such that for all matrices $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times p}$ and scalar $\lambda \in \mathbb{R}$

- (a) $\|\mathbf{A}\| \geq 0$ for all $\mathbf{A} \in \mathbb{R}^{n \times p}$ *nonnegativity*
- (b) $\|\mathbf{A}\| = 0$ if and only if $\mathbf{A} = \mathbf{0}$ *definitiveness*
- (c) $\|\lambda \mathbf{A}\| = |\lambda| \|\mathbf{A}\|$ *homogeneity*
- (d) $\|\mathbf{A} + \mathbf{B}\| \leq \|\mathbf{A}\| + \|\mathbf{B}\|$ *triangle inequality*

Definition (Matrix inner product)

Matrix inner product is defined as follows

$$\langle \mathbf{A}, \mathbf{B} \rangle = \text{trace}(\mathbf{A}\mathbf{B}^T).$$

Matrix norms contd.

- Similar to vector ℓ_q -norms, we have Schatten q -norms for matrices.

Definition (Schatten q -norms)

$\|\mathbf{A}\|_q := \left(\sum_{i=1}^p (\sigma(\mathbf{A})_i)^q \right)^{1/q}$, where $\sigma(\mathbf{A})_i$ is the i^{th} singular value of $\mathbf{A}$.

Example (with $r = \min\{n, p\}$ and $\sigma_i = \sigma(\mathbf{A})_i$)

$$\|\mathbf{A}\|_1 = \|\mathbf{A}\|_* := \sum_{i=1}^r \sigma_i \equiv \text{trace} \left(\sqrt{\mathbf{A}^T \mathbf{A}} \right) \quad (\text{Nuclear/trace})$$

$$\|\mathbf{A}\|_2 = \|\mathbf{A}\|_F := \sqrt{\sum_{i=1}^r (\sigma_i)^2} \equiv \sqrt{\sum_{i=1}^n \sum_{j=1}^p |a_{ij}|^2} \quad (\text{Frobenius})$$

$$\|\mathbf{A}\|_\infty = \|\mathbf{A}\| := \max_{i=1, \dots, r} \{\sigma_i\} \equiv \max_{\mathbf{x} \neq \mathbf{0}} \frac{\|\mathbf{A}\mathbf{x}\|}{\|\mathbf{x}\|} \quad (\text{Spectral/matrix})$$

Matrix norms contd.

Problem (Rank- r approximation)

Find $\arg \min_{\mathbf{X}} \|\mathbf{X} - \mathbf{Y}\|_F$ subject to: $\text{rank}(\mathbf{X}) \leq r$.

Matrix norms contd.

Problem (Rank- r approximation)

Find $\arg \min_{\mathbf{X}} \|\mathbf{X} - \mathbf{Y}\|_F$ subject to: $\text{rank}(\mathbf{X}) \leq r$.

Solution (Eckart–Young–Mirsky Theorem)

$$\begin{aligned} \arg \min_{\mathbf{X}: \text{rank}(\mathbf{X}) \leq r} \|\mathbf{X} - \mathbf{Y}\|_F &= \arg \min_{\mathbf{X}: \text{rank}(\mathbf{X}) \leq r} \|\mathbf{X} - \mathbf{U}\Sigma_{\mathbf{Y}}\mathbf{V}^T\|_F, \quad (\text{SVD}) \\ &= \arg \min_{\mathbf{X}: \text{rank}(\mathbf{X}) \leq r} \|\mathbf{U}^T\mathbf{X}\mathbf{V} - \Sigma_{\mathbf{Y}}\|_F, \quad (\text{unit. invar. of } \|\cdot\|_F) \\ &= \mathbf{U} \left(\arg \min_{\mathbf{Z}: \text{rank}(\mathbf{Z}) \leq r} \|\mathbf{Z} - \Sigma_{\mathbf{Y}}\|_F \right) \mathbf{V}^T, \quad (\text{Let } \mathbf{Z} = \mathbf{U}\mathbf{X}\mathbf{V}^T) \\ &= \mathbf{U}H_r(\Sigma_{\mathbf{Y}})\mathbf{V}^T, \quad (r\text{-sparse approx. of the diagonal entries}) \end{aligned}$$

Singular value hard thresholding operator H_r performs the **best rank- r approximation** of a matrix via sparse approximation: We keep the r **largest singular values** of the matrix and set the rest to zero.

Matrix norms contd.

Definition (Operator norm)

The **operator norm** between ℓ_q and ℓ_r ($1 \leq q, r \leq \infty$) of a matrix $\mathbf{A}$ is defined as

$$\|\mathbf{A}\|_{q \rightarrow r} = \sup_{\|\mathbf{x}\|_q \leq 1} \|\mathbf{A}\mathbf{x}\|_r$$

Problem

Show that $\|\mathbf{A}\|_{2 \rightarrow 2} = \|\mathbf{A}\|$ i.e., ℓ_2 to ℓ_2 operator norm is the *spectral* norm.

Solution

$$\begin{aligned} \|\mathbf{A}\|_{2 \rightarrow 2} &= \sup_{\|\mathbf{x}\|_2 \leq 1} \|\mathbf{A}\mathbf{x}\|_2 = \sup_{\|\mathbf{x}\|_2 \leq 1} \|\mathbf{U}\Sigma\mathbf{V}^T\mathbf{x}\|_2 \quad (\text{using SVD of } \mathbf{A}) \\ &= \sup_{\|\mathbf{x}\|_2 \leq 1} \|\Sigma\mathbf{V}^T\mathbf{x}\|_2 \quad (\text{rotational invariance of } \|\cdot\|_2) \\ &= \sup_{\|\mathbf{z}\|_2 \leq 1} \|\Sigma\mathbf{z}\|_2 \quad (\text{letting } \mathbf{V}^T\mathbf{x} = \mathbf{z}) \\ &= \sup_{\|\mathbf{z}\|_2 \leq 1} \sqrt{\sum_{i=1}^{\min(n,p)} \sigma_i^2 z_i^2} = \sigma_{\max} = \|\mathbf{A}\| \quad \square \end{aligned}$$

Matrix norms contd.

Other examples

- ▶ The $\|\mathbf{A}\|_{\infty \rightarrow \infty}$ (norm induced by ℓ_∞ -norm) also denoted $\|\mathbf{A}\|_\infty$, is the **max-row-sum norm**:

$$\|\mathbf{A}\|_{\infty \rightarrow \infty} := \sup\{\|\mathbf{Ax}\|_\infty \mid \|\mathbf{x}\|_\infty \leq 1\} = \max_{i=1,\dots,n} \sum_{j=1}^p |a_{ij}|.$$

- ▶ The $\|\mathbf{A}\|_{1 \rightarrow 1}$ (norm induced by ℓ_1 -norm) also denoted $\|\mathbf{A}\|_1$, is the **max-column-sum norm**:

$$\|\mathbf{A}\|_{1 \rightarrow 1} := \sup\{\|\mathbf{Ax}\|_1 \mid \|\mathbf{x}\|_1 \leq 1\} = \max_{i=1,\dots,p} \sum_{j=1}^n |a_{ij}|.$$

Matrix norms contd.

Useful relation for operator norms

The following **identity** holds

$$\|\mathbf{A}\|_{q \rightarrow r} := \max_{\|\mathbf{z}\|_r \leq 1, \|\mathbf{x}\|_q = 1} \langle \mathbf{z}, \mathbf{A}\mathbf{x} \rangle = \max_{\|\mathbf{x}\|_{q'} \leq 1, \|\mathbf{z}\|_{r'} = 1} \langle \mathbf{A}^T \mathbf{z}, \mathbf{x} \rangle =: \|\mathbf{A}^T\|_{q' \rightarrow r'}$$

whenever $1/q + 1/q' = 1 = 1/r + 1/r'$.

Example

1. $\|\mathbf{A}\|_{\infty \rightarrow 1} = \|\mathbf{A}^T\|_{1 \rightarrow \infty}$.
2. $\|\mathbf{A}\|_{2 \rightarrow 1} = \|\mathbf{A}^T\|_{2 \rightarrow \infty}$.
3. $\|\mathbf{A}\|_{\infty \rightarrow 2} = \|\mathbf{A}^T\|_{1 \rightarrow 2}$.

*Matrix norms contd.

Computation of operator norms

► The computation of some **operator norms** is NP-hard* [4]; these include:

1. $\|\mathbf{A}\|_{\infty \rightarrow 1}$
2. $\|\mathbf{A}\|_{2 \rightarrow 1}$
3. $\|\mathbf{A}\|_{\infty \rightarrow 2}$

► But some of them are **approximable** [11]; these include

1. $\|\mathbf{A}\|_{\infty \rightarrow 1}$ (via Gronthendieck factorization)
2. $\|\mathbf{A}\|_{\infty \rightarrow 2}$ (via Pietz factorization)

*: See Recitation 1.

Matrix norms contd.

Matrix & vector norm analogy

Vectors	$\ \mathbf{x}\ _1$	$\ \mathbf{x}\ _2$	$\ \mathbf{x}\ _\infty$
Matrices	$\ \mathbf{X}\ _*$	$\ \mathbf{X}\ _F$	$\ \mathbf{X}\ $

Definition (Dual of a matrix)

The **dual norm** of $\mathbf{A} \in \mathbb{R}^{n \times p}$ is defined as

$$\|\mathbf{A}\|^* = \sup \left\{ \text{trace}(\mathbf{A}^T \mathbf{X}) \mid \|\mathbf{X}\| \leq 1 \right\}.$$

Matrix & vector dual norm analogy

Vector primal norm	$\ \mathbf{x}\ _1$	$\ \mathbf{x}\ _2$	$\ \mathbf{x}\ _\infty$
Vector dual norm	$\ \mathbf{x}\ _\infty$	$\ \mathbf{x}\ _2$	$\ \mathbf{x}\ _1$
Matrix primal norm	$\ \mathbf{X}\ _*$	$\ \mathbf{X}\ _F$	$\ \mathbf{X}\ $
Matrix dual norm	$\ \mathbf{X}\ $	$\ \mathbf{X}\ _F$	$\ \mathbf{X}\ _*$

Matrix norms contd.

Definition (Nuclear norm computation)

$$\begin{aligned}\|\mathbf{A}\|_* &:= \|\boldsymbol{\sigma}(\mathbf{A})\|_1 \quad \text{where } \boldsymbol{\sigma}(\mathbf{A}) \text{ is a vector of singular values of } \mathbf{A} \\ &= \min_{\mathbf{U}, \mathbf{V}: \mathbf{A} = \mathbf{U}\mathbf{V}^H} \|\mathbf{U}\|_F \|\mathbf{V}\|_F = \min_{\mathbf{U}, \mathbf{V}: \mathbf{A} = \mathbf{U}\mathbf{V}^H} \frac{1}{2} (\|\mathbf{U}\|_F^2 + \|\mathbf{V}\|_F^2)\end{aligned}$$

Additional useful properties are below:

- ▶ Nuclear vs. Frobenius: $\|\mathbf{A}\|_F \leq \|\mathbf{A}\|_* \leq \sqrt{\text{rank}(\mathbf{A})} \cdot \|\mathbf{A}\|_F$
- ▶ Hölder for matrices: $|\langle \mathbf{A}, \mathbf{B} \rangle| \leq \|\mathbf{A}\|_p \|\mathbf{B}\|_q$, when $\frac{1}{p} + \frac{1}{q} = 1$
- ▶ We have
 1. $\|\mathbf{A}\|_{2 \rightarrow 2} \leq \|\mathbf{A}\|_F$
 2. $\|\mathbf{A}\|_{2 \rightarrow 2} \leq \|\mathbf{A}\|_{1 \rightarrow 1} \|\mathbf{A}\|_{\infty \rightarrow \infty}$
 3. $\|\mathbf{A}\|_{2 \rightarrow 2} \leq \|\mathbf{A}\|_{1 \rightarrow 1}$ when $\mathbf{A}$ is self-adjoint.

Linear systems

Problem (Solving a linear system)

Which is the best method for solving the linear system

$$\mathbf{Ax} = \mathbf{b} ?$$

Solving a linear system via optimization

To find a solution to the linear system, we can also solve the optimization problem

$$\min_{\mathbf{x}} f_{\mathbf{A},\mathbf{b}}(\mathbf{x}) := \frac{1}{2} \langle \mathbf{Ax}, \mathbf{x} \rangle - \langle \mathbf{b}, \mathbf{x} \rangle$$

which is seen to have a solution satisfying $\mathbf{Ax} = \mathbf{b}$ by solving $\nabla_{\mathbf{x}} f_{\mathbf{A},\mathbf{b}}(\mathbf{x}) = 0$.

- ▶ $f_{\mathbf{A},\mathbf{b}}$ is a quadratic function with **Lipschitz-gradient** ($L = \|\mathbf{A}\|$).
- ▶ If $\mathbf{A}$ is a $p \times p$ symmetric positive definite matrix, (i.e., $\mathbf{A} = \mathbf{A}^T \succ 0$), $f_{\mathbf{A}}$ is also **strongly convex** ($\mu = \lambda_1(\mathbf{A})$, the smallest eigenvalue of $\mathbf{A}$).
- ▶ if $\mathbf{A}$ is not symmetric, but full column rank, we can consider the system

$$\mathbf{A}^T \mathbf{Ax} = \mathbf{A}^T \mathbf{b}$$

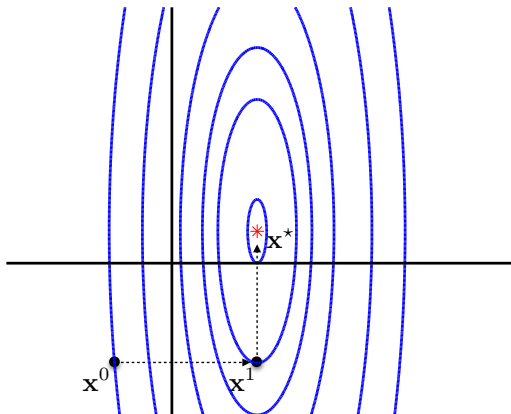
which can be seen as: $\Phi \mathbf{x} = \mathbf{y}$ where Φ is symmetric and positive definite.

Linear systems

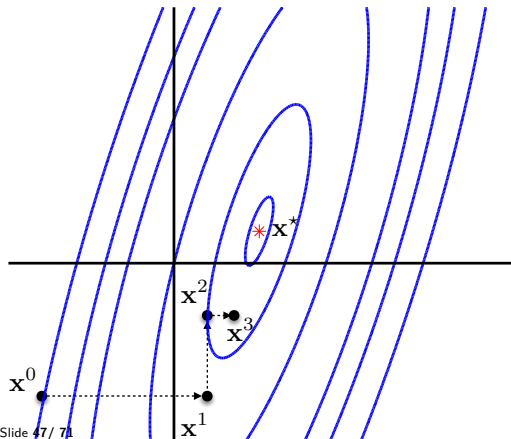
Remark

If Φ is diagonal and positive definite, given a starting point $\mathbf{x}^0 \in \text{dom}(f)$, successive minimization of $f_{\Phi, \mathbf{y}}(\mathbf{x})$ along the coordinate axes yield $\mathbf{x}^*$ is at most p steps.

Diagonal Φ



Non-diagonal Φ



How can we adapt to the geometry of Φ ?

Conjugate gradients method - Φ symmetric and positive definite

Generate a set of *conjugate* directions $\{\mathbf{p}^0, \mathbf{p}^1, \dots, \mathbf{p}^{p-1}\}$ such that

$$\langle \mathbf{p}^i, \Phi \mathbf{p}^j \rangle = 0 \quad \text{for all } i \neq j \quad (\text{which also implies linear independence}).$$

Successively minimize $f_{\Phi, \mathbf{y}}$ along the individual conjugate directions. Let

$$\mathbf{r}^k = \Phi \mathbf{x}^k - \mathbf{y} \quad \text{and} \quad \mathbf{x}^{k+1} = \mathbf{x}^k + \alpha_k \mathbf{p}^k,$$

where α_k is the minimizer of $f_{\Phi, \mathbf{y}}(\mathbf{x})$ along $\mathbf{x}^k + \alpha \mathbf{p}^k$, i.e.,

$$\alpha_k = - \frac{\langle \mathbf{r}^k, \mathbf{p}^k \rangle}{\langle \mathbf{p}^k, \Phi \mathbf{p}^k \rangle}$$

Theorem

For any $\mathbf{x}^0 \in \mathbb{R}^p$ the sequence $\{\mathbf{x}^k\}$ generated by the conjugate directions algorithm converges to the solution $\mathbf{x}^*$ of the linear system in **at most** p steps.

Intuition

The conjugate directions adapt to the geometry of the problem, taking the role of the canonical directions when Φ is a generic symmetric positive definite matrix.

Conjugate gradients method

Intuition

The conjugate directions adapt to the geometry of the problem, taking the role of the canonical directions when Φ is a generic symmetric positive definite matrix.

Back to diagonal

For a generic symmetric positive definite Φ , let us consider the variable $\bar{\mathbf{x}} := \mathbf{S}^{-1}\mathbf{x}$, with

$$\mathbf{S} = [\mathbf{p}^0, \dots, \mathbf{p}^{p-1}]$$

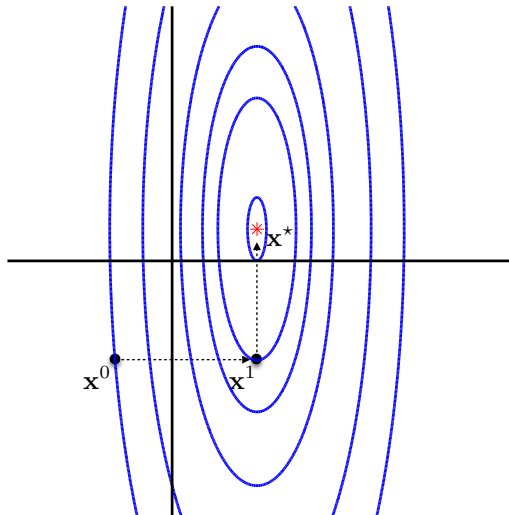
where $\{\mathbf{p}^k\}$ are the conjugate directions with respect to Φ . $f_{\Phi, \mathbf{y}}(\mathbf{x})$ now becomes

$$\bar{f}_{\Phi, \mathbf{y}}(\bar{\mathbf{x}}) := f_{\Phi, \mathbf{y}}(\mathbf{S}\bar{\mathbf{x}}) = \frac{1}{2} \langle \bar{\mathbf{x}}, (\mathbf{S}^T \Phi \mathbf{S}) \bar{\mathbf{x}} \rangle - \langle \mathbf{S}^T \mathbf{y}, \bar{\mathbf{x}} \rangle.$$

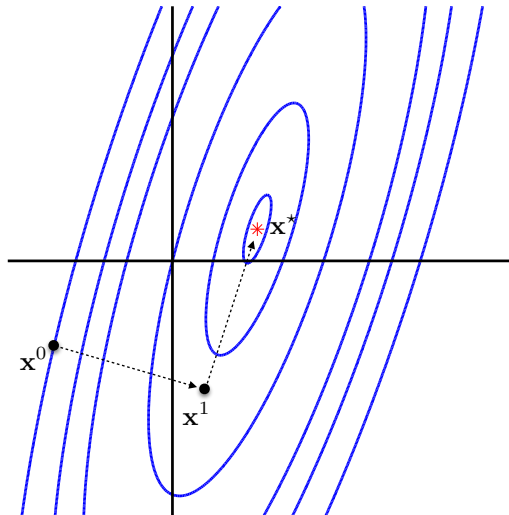
By the conjugacy property, $\langle \mathbf{p}^i, \Phi \mathbf{p}^j \rangle = 0$, $\forall i \neq j$, the matrix $\mathbf{S}^T \Phi \mathbf{S}$ is diagonal. Therefore, we can find the minimum of $\bar{f}(\bar{\mathbf{x}})$ in at most p steps along the canonical directions in $\bar{\mathbf{x}}$ space, which are the $\{\mathbf{p}^k\}$ directions in $\mathbf{x}$ space.

Conjugate directions naturally adapt to the linear operator

Diagonal Φ



Non-diagonal Φ



Conjugate gradients method

Theorem

For any $\mathbf{x}^0 \in \mathbb{R}^p$ the sequence $\{\mathbf{x}^k\}$ generated by the conjugate directions algorithm converges to the solution $\mathbf{x}^*$ of the linear system in **at most** p steps.

Proof.

Since $\{\mathbf{p}^k\}$ are linearly independent, they span $\mathbb{R}^p$. Therefore, we can write

$$\mathbf{x}^* - \mathbf{x}^0 = a_0 \mathbf{p}^0 + a_1 \mathbf{p}^1 + \cdots + a_{p-1} \mathbf{p}^{p-1}$$

for some values of the coefficients a_k . By multiplying on the left by $(\mathbf{p}^k)^T \Phi$ and using the conjugacy property, we obtain

$$a_k = \frac{\langle \mathbf{p}^k, \Phi(\mathbf{x}^* - \mathbf{x}^0) \rangle}{\langle \mathbf{p}^k, \Phi \mathbf{p}^k \rangle}.$$

Since $\mathbf{x}^k = \mathbf{x}^{k-1} + \alpha_{k-1} \mathbf{p}^{k-1}$, we have $\mathbf{x}^k = \mathbf{x}^0 + \alpha_0 \mathbf{p}^0 + \alpha_1 \mathbf{p}^1 + \cdots + \alpha_{k-1} \mathbf{p}^{k-1}$. By premultiplying by $(\mathbf{p}^k)^T \Phi$ and using the conjugacy property, we obtain $\langle \mathbf{p}^k, \Phi(\mathbf{x}^k - \mathbf{x}^0) \rangle = 0$ which implies

$$\langle \mathbf{p}^k, \Phi(\mathbf{x}^* - \mathbf{x}^0) \rangle = \langle \mathbf{p}^k, \Phi(\mathbf{x}^* - \mathbf{x}^k) \rangle = \langle \mathbf{p}^k, \mathbf{y} - \Phi \mathbf{x}^k \rangle = -\langle \mathbf{p}^k, \mathbf{r}^k \rangle$$

$$\text{so that } a_k = -\frac{\langle \mathbf{p}^k, \mathbf{r}^k \rangle}{\langle \mathbf{p}^k, \Phi \mathbf{p}^k \rangle} = \alpha_k.$$

□

Conjugate gradients method

How can we efficiently generate a set of conjugate directions?

Iteratively generate the new descent direction $\mathbf{p}^k$ from the previous one:

$$\mathbf{p}^k = -\mathbf{r}^k + \beta_k \mathbf{p}^{k-1}$$

For ensuring conjugacy $\langle \mathbf{p}^k, \Phi \mathbf{p}^{k-1} \rangle = 0$, we need to choose β_k as

$$\beta_k = \frac{\langle \mathbf{r}^k, \Phi \mathbf{p}^{k-1} \rangle}{\langle \mathbf{p}^{k-1}, \Phi \mathbf{p}^{k-1} \rangle} .$$

Lemma

The directions $\{\mathbf{p}^0, \mathbf{p}^1, \dots, \mathbf{p}^p\}$ form a conjugate directions set.

Conjugate gradients method

Conjugate gradients (CG) method

1 Initialization:

1.a Choose $\mathbf{x}^0 \in \text{dom}(f)$ arbitrarily as a starting point.

1.b Set $\mathbf{r}^0 = \Phi \mathbf{x}^0 - \mathbf{y}$, $\mathbf{p}^0 = -\mathbf{r}^0$, $k = 0$.

2. While $\mathbf{r}^k \neq \mathbf{0}$, generate a sequence $\{\mathbf{x}^k\}_{k \geq 0}$ as:

$$\begin{aligned}\alpha_k &= -\frac{\langle \mathbf{r}^k, \mathbf{p}^k \rangle}{\langle \mathbf{p}^k, \Phi \mathbf{p}^k \rangle} \\ \mathbf{x}^{k+1} &= \mathbf{x}^k + \alpha_k \mathbf{p}^k \\ \mathbf{r}^{k+1} &= \Phi \mathbf{x}^{k+1} - \mathbf{y} \\ \beta_{k+1} &= \frac{\langle \mathbf{r}^{k+1}, \Phi \mathbf{p}^k \rangle}{\langle \mathbf{p}^k, \Phi \mathbf{p}^k \rangle} \\ \mathbf{p}^{k+1} &= -\mathbf{r}^{k+1} + \beta_{k+1} \mathbf{p}^k \\ k &= k + 1\end{aligned}$$

Theorem

Since the directions $\{\mathbf{p}^0, \mathbf{p}^1, \dots, \mathbf{p}^k\}$ are conjugate, CG converges in at most p steps.

Other properties of the conjugate gradient method

Theorem

If Φ has only r distinct eigenvalues, then the CG iterations will terminate at the solution in at most r iterations.

Theorem

If Φ has eigenvalues $\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_p$, we have that

$$\|\mathbf{x}^{k+1} - \mathbf{x}^*\|_{\Phi} \leq \left(\frac{\lambda_{p-k} - \lambda_1}{\lambda_{p-k} + \lambda_1} \right) \|\mathbf{x}^0 - \mathbf{x}^*\|_{\Phi},$$

where the local norm is given by $\|\mathbf{x}\|_{\Phi} = \sqrt{\mathbf{x}^T \Phi \mathbf{x}}$.

Theorem

Conjugate gradients algorithm satisfy the following iteration invariant for the solution of $\Phi \mathbf{x} = \mathbf{y}$

$$\|\mathbf{x}^{k+1} - \mathbf{x}^*\|_{\Phi} \leq 2 \left(\frac{\sqrt{\kappa(\Phi)} - 1}{\sqrt{\kappa(\Phi)} + 1} \right)^k \|\mathbf{x}^0 - \mathbf{x}^*\|_{\Phi},$$

where the condition number of Φ is defined as $\kappa(\Phi) := \|\Phi\| \|\Phi^{-1}\| = \frac{\lambda_p}{\lambda_1}$.

Matrix perturbation inequalities

- ▶ In the theorems below $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{p \times p}$ are **symmetric** positive semi-definite matrices with **spectra** $\{\lambda_i(\mathbf{A})\}_{i=1}^p$ and $\{\lambda_i(\mathbf{B})\}_{i=1}^p$ where $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p$.

Theorem (Lidskii inequality)

$\lambda_{i_1}(\mathbf{A} + \mathbf{B}) + \dots + \lambda_{i_n}(\mathbf{A} + \mathbf{B}) \leq \lambda_{i_1}(\mathbf{A}) + \dots + \lambda_{i_n}(\mathbf{A}) + \lambda_{i_1}(\mathbf{B}) + \dots + \lambda_{i_n}(\mathbf{B})$,
for any $1 \leq i_1 \leq \dots \leq i_n \leq p$.

Theorem (Weyl inequality)

$$\lambda_{i+j-1}(\mathbf{A} + \mathbf{B}) \leq \lambda_i(\mathbf{A}) + \lambda_j(\mathbf{B}), \quad \text{for any } i, j \geq 1 \text{ and } i + j - 1 \leq p.$$

Theorem (Interlacing property)

Let $\mathbf{A}_n = \mathbf{A}(1:n, 1:n)$, then

$$\lambda_{n+1}(\mathbf{A}_{n+1}) \leq \lambda_n(\mathbf{A}_n) + \lambda_n(\mathbf{A}_{n+1}) \quad \text{for } n = 1, \dots, p.$$

- ▶ These inequalities **hold** in the **more general setting** when λ_i are replaced by σ_i .
- ▶ The list goes on to include **Wedin** bounds, **Wielandt-Hoffman** bounds and so on.
- ▶ More on such inequalities can be found in **Terry Tao's blog (254A, Notes 3a)**.

*Tensors

1. Basic tensor definitions
2. Notation and preliminaries
3. Tensors decompositions
4. Tensor rank
5. Banach's result on supersymmetric tensors

*: PhD material

Basic definitions

- **Tensors** provide natural and concise mathematical representations of **data**.

Definition (Tensor)

An **order** m tensor in p -**dimensional** space is a mathematical object that has p *indices* and p^m *components* and obeys certain transformation rules.

- ▶ In the literature, **rank** is used interchangeably with **order**, i.e., an **order- k** tensor is also referred to as **k th-rank** tensor.
- ▶ We often use **order** instead of **rank** so that it is not confused with the **rank of a tensor**.
- ▶ Furthermore, **mode** or **way** is also used to refer to the **order** of a tensor.
- ▶ **Tensors** are **multidimensional arrays** and are a generalization of:
 1. **scalars** - **tensors** with *no indices*; i.e., order **zero** tensor.
 2. **vectors** - **tensors** with exactly *one index*; i.e., order **one** tensor.
 3. **matrices** - **tensors** with exactly *two indices*; i.e., order **two** tensor.
- ▶ A **third-order** tensor has exactly *three indices*.
- ▶ A **higher-order** tensor has *greater than two indices*; i.e., a tensor of order ≥ 2 .

Notation & preliminaries

Notation & preliminaries

- ▶ The notation conforms to [7] which is the main reference for this material.
- ▶ Higher-order tensors are denoted by **boldface Euler script letters**, e.g. $\mathcal{A}$.
- ▶ Element $(i, j, k, \dots)$ of a **tensor** $\mathcal{A}$ are denoted by $a_{ijk\dots}$.
- ▶ The m th element in a sequence is denoted by a **superscript in parentheses**, e.g. $\mathbf{A}^{(m)}$ denotes the m th matrix in a sequence.
- ▶ **Subarrays** of a tensor are formed when a **subset of the indices** of the elements of a tensor are fixed.
- ▶ **Fibers** are the higher-order analogue of matrix rows and columns, defined by *fixing every index but one*.
- ▶ **Slices** are 2-dimensional sections of a tensor, defined by **fixing all but 2 indices**.
For instance, the horizontal, lateral, and frontal **slices** of a third-order tensor $\mathcal{A}$ are denoted by $\mathbf{A}_{i::}$, $\mathbf{A}_{:j:}$, & $\mathbf{A}_{::k}$ (or more compactly $\mathbf{A}_i$, $\mathbf{A}_j$, & $\mathbf{A}_k$) respectively.

Notation & preliminaries

Notation & preliminaries

- ▶ The notation conforms to [7] which is the main reference for this material.
- ▶ Higher-order tensors are denoted by **boldface Euler script letters**, e.g. $\mathcal{A}$.
- ▶ Element $(i, j, k, \dots)$ of a **tensor** $\mathcal{A}$ are denoted by $a_{ijk\dots}$.
- ▶ The m th element in a sequence is denoted by a **superscript in parentheses**, e.g. $\mathbf{A}^{(m)}$ denotes the m th matrix in a sequence.
- ▶ **Subarrays** of a tensor are formed when a **subset of the indices** of the elements of a tensor are fixed.
- ▶ **Fibers** are the higher-order analogue of matrix rows and columns, defined by *fixing every index but one*.
- ▶ **Slices** are 2-dimensional sections of a tensor, defined by **fixing all but 2 indices**.
For instance, the horizontal, lateral, and frontal **slices** of a third-order tensor $\mathcal{A}$ are denoted by $\mathbf{A}_{i::}$, $\mathbf{A}_{:j:}$, & $\mathbf{A}_{::k}$ (or more compactly $\mathbf{A}_i$, $\mathbf{A}_j$, & $\mathbf{A}_k$) respectively.

Curse of dimensionality

Storage of an **order- m tensor** with mode sizes p requires p^m elements.

Notation & preliminaries contd.

- Tensors are **linear vector** spaces.

Definition (Norm)

The **norm** of a tensor $\mathcal{A} \in \mathbb{R}^{p_1 \times p_2 \times \cdots \times p_k}$ is given by

$$\|\mathcal{A}\| = \sqrt{\sum_{i_1=1}^{p_1} \sum_{i_2=1}^{p_2} \cdots \sum_{i_k=1}^{p_k} a_{i_1 i_2 \dots i_k}^2}$$

- This is the analogue to the matrix **Frobenius norm**.

Definition (Inner product)

The **inner product** of two **same-sized** tensors $\mathcal{X}, \mathcal{Y} \in \mathbb{R}^{p_1 \times p_2 \times \cdots \times p_k}$ is given by

$$\langle \mathcal{X}, \mathcal{Y} \rangle = \sum_{i_1=1}^{p_1} \sum_{i_2=1}^{p_2} \cdots \sum_{i_k=1}^{p_k} x_{i_1 i_2 \dots i_k} y_{i_1 i_2 \dots i_k}$$

- It follows immediately that $\langle \mathcal{A}, \mathcal{A} \rangle = \|\mathcal{A}\|^2$.

Notation & preliminaries contd.

Rank-one tensors

A k -way tensor $\mathcal{A} \in \mathbb{R}^{p_1 \times p_2 \times \dots \times p_k}$ is **rank-one** if it can be written as the *outer product* of k vectors, i.e.

$$\mathcal{A} = \mathbf{v}^{(1)} \circ \mathbf{v}^{(2)} \circ \dots \circ \mathbf{v}^{(k)}$$

where “ $\circ$ ” represents the *vector outer product*.

- ▶ Each element of the tensor is the product of the corresponding vector elements:

$$x_{i_1 i_2 \dots i_k} = v_{i_1}^{(1)} v_{i_2}^{(2)} \dots v_{i_k}^{(k)} \quad \forall 1 \leq i_n \leq p_n.$$

Notation & preliminaries contd.

Rank-one tensors

A k -way tensor $\mathcal{A} \in \mathbb{R}^{p_1 \times p_2 \times \dots \times p_k}$ is **rank-one** if it can be written as the *outer product* of k vectors, i.e.

$$\mathcal{A} = \mathbf{v}^{(1)} \circ \mathbf{v}^{(2)} \circ \dots \circ \mathbf{v}^{(k)}$$

where “ $\circ$ ” represents the *vector outer product*.

- Each element of the tensor is the product of the corresponding vector elements:

$$x_{i_1 i_2 \dots i_k} = v_{i_1}^{(1)} v_{i_2}^{(2)} \dots v_{i_k}^{(k)} \quad \forall 1 \leq i_n \leq p_n.$$

Definition (Cubical tensors)

A tensor $\mathcal{A} \in \mathbb{R}^{p_1 \times \dots \times p_k}$ is **cubical** if every mode is **same size**, i.e. $p_1 = \dots = p_k = p$; as a shorthand an order- k cubical tensors is denoted as $\mathbf{A} \in \otimes^k \mathbb{R}^p$.

Definition (Symmetric tensors)

A cubical tensor $\mathbf{A} \in \otimes^k \mathbb{R}^p$ is **symmetric** (also referred to as **super-symmetric**) if its k -way representations are **invariant** to permutations of the array indices: i.e. for all indices $i_1, i_2, \dots, i_k \in [p]$ and any permutation π on k :

$$a_{i_1 i_2 \dots i_k} = a_{i_{\pi(1)} i_{\pi(2)} \dots i_{\pi(k)}}.$$

Notation & preliminaries contd.

Why tensors are important?

Multivariate functions are related to **multidimensional arrays** or *tensors*:

Take a function $f(\mathbf{x}_1, \dots, \mathbf{x}_p)$; take a tensor-product grid and get a **tensor**, i.e.

$$a_{i_1 i_2 \dots i_p} = f(\mathbf{x}_1(i_1), \dots, \mathbf{x}_p(i_p))$$

Notation & preliminaries contd.

Why tensors are important?

Multivariate functions are related to **multidimensional arrays** or *tensors*:

Take a function $f(\mathbf{x}_1, \dots, \mathbf{x}_p)$; take a tensor-product grid and get a **tensor**, i.e.

$$a_{i_1 i_2 \dots i_p} = f(\mathbf{x}_1(i_1), \dots, \mathbf{x}_p(i_p))$$

Where does tensors come from?

- ▶ n -th derivative of a multivariate function $f(x_1, \dots, x_p)$, i.e. $\nabla^n f(x_1, \dots, x_p)$
- ▶ p -dimensional PDE: $\Delta u = f$, $u = u(\mathbf{x}_1, \dots, \mathbf{x}_p)$
- ▶ Data (images, video, hyperspectral images, etc)
- ▶ Latent variable models, joint probability distributions
- ▶ Many others

Tensor decomposition

Definition (Tensor decomposition [7])

Tensor decomposition refers to the factorization of a tensor into a finite sum of component rank-one tensors.

- ▶ This is the analogue of the [SVD for matrices](#) and is also known as **parallel factors** and **canonical decompositions**.

Tensor decomposition

Definition (Tensor decomposition [7])

Tensor decomposition refers to the factorization of a tensor into a finite sum of component rank-one tensors.

- This is the analogue of the [SVD for matrices](#) and is also known as **parallel factors** and **canonical decompositions**.

Example

Given a order-3 tensor $\mathcal{A} \in \mathbb{R}^{p_1 \times p_2 \times p_3}$, it's decomposition attempts to express it as

$$\mathcal{A} \approx \sum_{r=1}^R \mathbf{x}_r \circ \mathbf{y}_r \circ \mathbf{z}_r,$$

where $R > 0$ is integer and for $r = 1, \dots, R$, $\mathbf{x}_r \in \mathbb{R}^{p_1}$, $\mathbf{y}_r \in \mathbb{R}^{p_2}$, and $\mathbf{z}_r \in \mathbb{R}^{p_3}$. Elementwise, this decomposition can be written as

$$a_{ijk} \approx \sum_{r=1}^R x_{ir} y_{jr} z_{kr} \quad \text{for } i = 1, \dots, p_1, j = 1, \dots, p_2, k = 1, \dots, p_3.$$

Tensor decomposition contd.

Definition (Factor matrices)

Given a decomposition $\mathcal{A} \approx \sum_{r=1}^R \mathbf{x}_r \circ \mathbf{y}_r \circ \mathbf{z}_r$, the **factor matrices** refers to the combination of the vectors from the rank-one components, i.e. $\mathbf{X} = [\mathbf{x}_1 \ \mathbf{x}_2 \ \cdots \ \mathbf{x}_R]$ and similarly for $\mathbf{Y}$ and $\mathbf{Z}$.

- ▶ Thus tensor decomposition can be concisely written as

$$\mathcal{A} \approx [[\mathbf{X}, \mathbf{Y}, \mathbf{Z}]] \equiv \sum_{r=1}^R \mathbf{x}_r \circ \mathbf{y}_r \circ \mathbf{z}_r.$$

- ▶ If we assume that the columns of $\mathbf{X}$, $\mathbf{Y}$, and $\mathbf{Z}$ are **normalized** with the weights absorbed in a vector $\boldsymbol{\lambda}$, then the tensor decomposition can further be expressed as

$$\mathcal{A} = [[\boldsymbol{\lambda}; \mathbf{X}, \mathbf{Y}, \mathbf{Z}]] \equiv \sum_{r=1}^R \lambda_r \mathbf{x}_r \circ \mathbf{y}_r \circ \mathbf{z}_r.$$

Tensor rank

Definition (Tensor rank)

The **rank** of a tensor $\mathcal{A}$ denoted $\text{rank}(\mathcal{A})$ is the smallest number of rank-one tensors that generate $\mathcal{A}$ as their sum.

- ▶ This is the smallest number of components in an exact tensor decomposition where “exact” means the decomposition holds with *equality*:

$$\mathcal{A} = [[\mathbf{X}, \mathbf{Y}, \mathbf{Z}]] \equiv \sum_{r=1}^R \mathbf{x}_r \circ \mathbf{y}_r \circ \mathbf{z}_r.$$

- ▶ An **exact** tensor decomposition with $R = \text{rank}(\mathcal{A})$ is called **rank decomposition**.
- ▶ This is the exact analogue of the definition of a matrix rank but the properties of a matrix and a tensor ranks are quite different.

Tensors rank contd.

Tensor rank approximation: caveat!

Not much is known about the **generalizability** of **matrix notions to tensors** particularly *rank approximation*.

- ▶ The equivalence of the **Eckart-Young-Mirsky theorem** for rank- k approximation of matrices does **not** exist for tensors.
 1. For instance, summing k of the factors of a third-order tensor of rank R does not necessarily yield a best rank- k approximation.
 2. Kolda [6] gave an example where the best rank- k approximation of a tensor is **not** a factor in the best rank-2 approximation.
- ▶ The notion of tensor (symmetric) rank is considerably more delicate than matrix (symmetric) rank. For instance:
 1. **Not** clear *a priori* that the symmetric rank should even be finite [3].
 2. Removal of the best rank-1 approximation of a general tensor may increase the tensor rank of the residual [9].
- ▶ It is **NP-hard** to compute the rank of a tensor in general; only **approximations** of **(super) symmetric** tensors possible [1].

★ Tensors as multilinear maps

- ▶ Just as a matrix can be **pre- & post-multiplied** by a pair of matrices, an order- k tensor can be *multiplied on k -sides* by k -matrices.

Definition (Multilinear maps with tensors)

For a set of matrices $\{\mathbf{X}_i \in \mathbb{R}^{p \times m_i} \mid i \in [k]\}$, the $(i_1, i_2, \dots, i_k)$ -th entry of a k -way array representation of $\mathcal{A}(\mathbf{X}_1, \dots, \mathbf{X}_k) \in \mathbb{R}^{m_1 \times \dots \times m_k}$ is

$$[\mathcal{A}(\mathbf{X}_1, \dots, \mathbf{X}_k)]_{i_1 \dots i_k} := \sum_{j_1, \dots, j_k \in [p]} a_{j_1 j_2 \dots j_k} [X_1]_{j_1 i_1} [X_2]_{j_2 i_2} \dots [X_k]_{j_k i_k},$$

where $[\mathbf{X}_i]_{jk}$ is the (j, k) entry of a matrix $\mathbf{X}_i$.

Example

1. If $\mathbf{A}$ is a **matrix** ($k = 2$), then

$$\mathbf{A}(\mathbf{X}_1, \mathbf{X}_2) = \mathbf{X}_1^T \mathbf{A} \mathbf{X}_2$$

2. For a **matrix** $\mathbf{A}$ and a **vector** $\mathbf{x} \in \mathbb{R}^p$, we can express $\mathbf{A}\mathbf{x}$ as

$$\mathbf{A}(\mathbf{I}, \mathbf{x}) = \mathbf{A}\mathbf{x}$$

3. With the **canonical basis** $\{\mathbf{e}_{i_1}, \mathbf{e}_{i_2}, \dots, \mathbf{e}_{i_k}\}$ we have

$$\mathbf{A}(\mathbf{e}_{i_1}, \mathbf{e}_{i_2}, \dots, \mathbf{e}_{i_k}) = A_{i_1, i_2, \dots, i_k}$$

★ Tensor compression and Tucker decomposition

- ▶ The **Tucker decomposition** is a form of **higher-order PCA**.
- ▶ It also goes by many other names, see [7].

Definition (Tucker decomposition [7])

The **Tucker decomposition** decomposes a tensor into a **core tensor** multiplied (or transformed) by a matrix along each mode.

Example

- ▶ In the case of a third-order tensor $\mathcal{A} \in \mathbb{R}^{p_1 \times p_2 \times p_3}$, we have

$$\mathcal{A} = \sum_{r_1=1}^{R_1} \sum_{r_2=1}^{R_2} \sum_{r_3=1}^{R_3} g_{r_1 r_2 r_3} \mathbf{x}_{r_1} \circ \mathbf{y}_{r_2} \circ \mathbf{z}_{r_3} = [[\mathcal{G}; \mathbf{X}, \mathbf{Y}, \mathbf{Z}]].$$

- ▶ The matrices $\mathbf{X} \in \mathbb{R}^{p_1 \times R_1}$, $\mathbf{Y} \in \mathbb{R}^{p_2 \times R_2}$, and $\mathbf{Z} \in \mathbb{R}^{p_3 \times R_3}$ are the factor matrices and are the **principal components** in each mode.
- ▶ The tensor $\mathcal{G} \in \mathbb{R}^{R_1 \times R_2 \times R_3}$ is the **core tensor** and its entries show the level of interaction between different components.

★ Banach's results for tensors

- ▶ Banach proved that the maximal overlap between a symmetric tensor and a rank-1 tensor is attained at a symmetric rank-1 tensor.
- ▶ Unfortunately, this—seemingly trivial result—is not obvious. That is, if $\mathbf{U} \in \text{Sym}^k(\mathbb{C}^p)$ is a k -index totally symmetric vector with d dimensions per index, then

$$\max_{\arg \mathbf{X} = \mathbf{x}_1 \circ \dots \circ \mathbf{x}_k, \|\mathbf{x}_i\|_2 = 1} |\langle \mathbf{X}, \mathbf{U} \rangle|^2$$

fulfills $\mathbf{x}_1 = \dots = \mathbf{x}_n$.

References I

- [1] Animashree Anandkumar, Rong Ge, Daniel Hsu, Sham M Kakade, and Matus Telgarsky.
Tensor decompositions for learning latent variable models.
Journal of machine learning research, 15:2773–2832, 2014.
(Cited on page 72.)
- [2] Waheed U Bajwa, Jarvis D Haupt, Gil M Raz, Stephen J Wright, and Robert D Nowak.
Toeplitz-structured compressed sensing matrices.
In *IEEE/SP Workshop Stat. Sig. Proc.*, pages 294–298. IEEE, 2007.
(Cited on page 25.)
- [3] Pierre Comon, Gene Golub, Lek-Heng Lim, and Bernard Mourrain.
Symmetric tensors and symmetric tensor rank.
SIAM Journal on Matrix Analysis and Applications, 30(3):1254–1279, 2008.
(Cited on page 72.)
- [4] Simon Foucart and Holger Rauhut.
A mathematical introduction to compressive sensing, volume 1.
Birkhäuser Basel, 2013.
(Cited on pages 24 and 46.)
- [5] Gene H Golub and Charles F Van Loan.
Matrix computations, volume 3.
JHU Press, 2012.
(Cited on page 33.)

References II

- [6] Tamara G Kolda.
Orthogonal tensor decompositions.
SIAM Journal on Matrix Anal. and Apps., 23(1):243–255, 2001.
(Cited on page 72.)
- [7] Tamara G Kolda and Brett W Bader.
Tensor decompositions and applications.
SIAM Rev., 51(3):455–500, 2009.
(Cited on pages 61, 62, 68, 69, and 74.)
- [8] Holger Rauhut, Justin Romberg, and Joel A Tropp.
Restricted isometries for partial random circulant matrices.
App. Comp. Harmon. Anal., 32(2):242–254, 2012.
(Cited on page 25.)
- [9] Alwin Stegeman and Pierre Comon.
Subtracting a best rank-1 approximation may increase tensor rank.
Linear Algebra and its Applications, 433(7):1276–1300, 2010.
(Cited on page 72.)
- [10] Gilbert Strang.
Applied linear algebra, 1976.
(Cited on page 34.)

References III

- [11] Joel A. Tropp.
Column subset selection, matrix factorization, and eigenvalue optimization.
In *Proc. ACM-SIAM Symp. Discrete Algorithms*, 2009.
(Cited on page 46.)
- [12] Alp Yurtsever, Joel A Tropp, Olivier Fercoq, Madeleine Udell, and Volkan Cevher.
Scalable semidefinite programming.
SIAM Journal on Mathematics of Data Science, 3(1):171–200, 2021.
(Cited on page 37.)