

Course I – Numerical Linear Algebra

Bruno Galerne
September 2021



What about this course?

Prerequisite course of mathematics for IoT

Content:

- Basics of numerical linear algebra (matrices etc.) (**starting today**)
- Basics of statistics (**2 courses with Diarra Fall**)
- Basics of mathematical optimization (gradient descent).

Goal:

- Being OK with the mathematical modeling that will appear in many courses.
- Basics of mathematics for **data science expertise class**.

Softwares:

- **R** for statistics
- **Python** for the other topics

Evaluation:

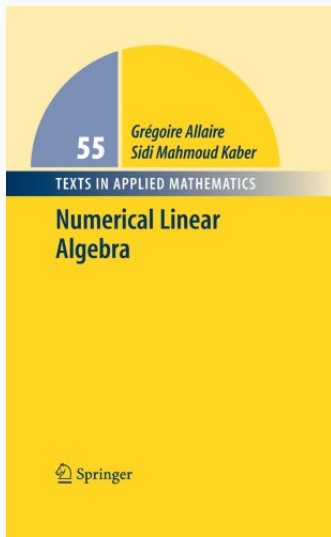
- Individual lab session report (~ 3 reports) with math exercices and numerical experiments.

Who am I?

- A professor of applied mathematics of Université d'Orléans.
 - Lab: Institut Denis Poisson (Université d'Orléans, Université de Tours, CNRS).
 - Research in image processing.
 - Research interests: Mathematical modeling of textures, texture synthesis algorithms, deep neural networks for image processing...
-
- Email: `bruno.galerie@univ-orleans.fr`
 - `https://www.idpoisson.fr/galerie/`

Recalls on linear algebra

- Scientific computing often relies on vector and matrix representation for multi-dimensional arrays (digital signals, images,...).
- The properties of matrices are often central for data science (e.g. correlation matrices for Principal Component Analysis, positive matrices for convex optimization...).
- Handling matrices numerically is different than in theory:
Never solve a linear system $Ax = b$ using $x = A^{-1}b$
- Also useful for ODE, PDE,...

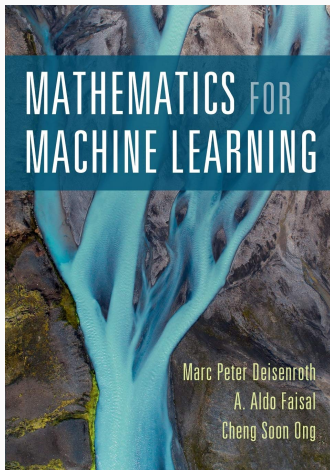


Numerical Linear Algebra,

by Grégoire Allaire and Sidi Mahmoud Kaber,
Texts in Applied Mathematics book series,
Springer

E-book freely available for UO students!

[https://catalogue.univ-orleans.fr/
Record/784674](https://catalogue.univ-orleans.fr/Record/784674)



Mathematics for Machine Learning,
by Marc Peter Deisenroth, A. Aldo Faisal, and
Cheng Soon Ong. Published by Cambridge
University Press, 2020.

E-book freely available!

<https://mml-book.github.io/>

- A *vector space* E is a set with two operations:
 - Addition of elements (called *vectors*):

$$\forall x, y \in E, \quad x + y \in E$$

- Multiplication by a scalar $\lambda \in \mathbb{K}$ with $\mathbb{K} = \mathbb{R}$ or $\mathbb{C}$:

$$\forall x \in E, \lambda \in \mathbb{K}, \quad \lambda x \in E$$

Examples:

- $\mathbb{R}^n$, set of all functions $f : \mathbb{R} \rightarrow \mathbb{C}$.
- Set of polynomial functions: $P(x) = \sum_{k=0}^d a_k x^k, x \in \mathbb{R}$.
- Set of sinusoidal functions:

$$f(t) = \sum_{k \in \mathbb{Z}} c_k e^{ikt}, \quad t \in \mathbb{R} \quad \text{with } c_k = 0 \text{ for all but a finite number of indices.}$$

Vector spaces: Subspaces

Definition

$F \subset E$ is a *vector subspace* of E if F is non empty and F is a vector space for the same operations as E .

Proposition

F is a subspace if and only if

- ❶ $0 \in F$
- ❷ $\forall x, y \in E, \lambda \in \mathbb{K}, \quad x + \lambda y \in F$

Definition (Direct sum)

Two subspaces F_1 and $F_2 \subset E$ are a *direct sum* of E , denoted $E = F_1 \oplus F_2$, if

- ❶ $F_1 \cap F_2 = \{0\}$
- ❷ For all $x \in E, \exists (x_1, x_2) \in F_1 \times F_2, x = x_1 + x_2$.

Remark: One can generalize the definition with more than two subspaces.

Vector spaces: Subspaces

Definition

$F \subset E$ is a *vector subspace* of E if F is non empty and F is a vector space for the same operations as E .

Proposition

F is a subspace if and only if

- ❶ $0 \in F$
- ❷ $\forall x, y \in E, \lambda \in \mathbb{K}, \quad x + \lambda y \in F$

Definition (Direct sum)

Two subspaces F_1 and $F_2 \subset E$ are a *direct sum* of E , denoted $E = F_1 \oplus F_2$, if

- ❶ $F_1 \cap F_2 = \{0\}$
- ❷ For all $x \in E, \exists (x_1, x_2) \in F_1 \times F_2, x = x_1 + x_2$.

Remark: One can generalize the definition with more than two subspaces.

Uniqueness: The above decomposition $x = x_1 + x_2$ with $(x_1, x_2) \in F_1 \times F_2$ is unique.

Vector spaces: Spanned vector spaces

Let $A \subset E$. We denote $\text{Span}(A)$ the vector space spanned by the vectors of A :

$$\text{Span}(A) = \left\{ x \in E, \exists p \in \mathbb{N}, a_1, \dots, a_p \in A, \lambda_1, \dots, \lambda_p \in \mathbb{K}, x = \sum_{k=1}^p \lambda_k a_k \right\}.$$

$\text{Span}(A)$ is the set of all finite linear combinations of vectors of A . This is the smallest subspace containing A .

Definition

Let $A \subset E$ be a family of vectors.

- A is said to be *linearly independent* if for any set $a_1, \dots, a_p$ of distincts vectors of A , $\sum_{k=1}^p \lambda_k a_k = 0$ implies $\lambda_1 = \dots = \lambda_p = 0$.
- A is a *generating family* if $\text{Span}(A) = E$, that is every vector of E can be expressed as a linear combination of vectors of A .
- A is a *basis* of E if A is both generating and linearly independent.

Theorem (Dimension theorem)

If E contains a finite generating family, then it contains a finite basis and all basis have the same number of vectors. This integer is the dimension of E denoted $\dim(E)$.

Vector spaces: Canonical basis

Example: The **canonical basis** of $\mathbb{R}^n$ is given by

$$e_k = \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix} \quad k\text{-th row}$$

Hence one has for any $x \in \mathbb{R}^n$

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_k \\ \vdots \\ x_n \end{pmatrix} = \sum_{k=1}^n x_k e_k = x_1 \begin{pmatrix} 1 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{pmatrix} + \cdots + x_k \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix} + \cdots + x_n \begin{pmatrix} 0 \\ \vdots \\ 0 \\ \vdots \\ 1 \end{pmatrix}.$$

Definition (Linear map)

Let E and F be two vector spaces. A function $U : E \rightarrow F$ is a linear map if $\forall x, y \in E, \lambda \in \mathbb{K}, \quad U(x + \lambda y) = U(x) + \lambda U(y)$.

If $F = E$, then U is called an endomorphism of E .

$$U(0_E) = 0_F$$

Associated subspaces:

- Kernel of U : $\text{Ker}(U) = \{x \in E, U(x) = 0_F\} \subset E$
- Image of U : $\text{Im}(U) = \{u(x), x \in E\} \subset F$

Rank theorem: If E has finite dimension,

$$\dim(\text{Ker}(U)) + \dim(\text{Im}(U)) = \dim(E).$$

From linear maps to matrices

Consider a linear map $U : E \rightarrow F$ and

- $\mathcal{B} = (e_1, \dots, e_p)$ a basis of E
- $\mathcal{C} = (f_1, \dots, f_n)$ a basis of F

then

$$U(x) = U\left(\sum_{j=1}^p x_j e_j\right) = \sum_{j=1}^p x_j U(e_j) = \sum_{j=1}^p x_j \sum_{i=1}^n U_{i,j} f_i$$

where $U_{i,j}$ is the i -th coordinate in the basis $\mathcal{C}$ of the vector $U(e_j)$.

$$(U_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}}$$

is the matrix representation of U for the basis $(\mathcal{B}, \mathcal{C})$.

Recalls on matrices

A matrix is a double entry table:

$$A = \begin{pmatrix} a_{1,1} & \dots & a_{1,p} \\ a_{2,1} & \dots & a_{2,p} \\ \vdots & & \vdots \\ a_{n,1} & \dots & a_{n,p} \end{pmatrix} = (a_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}}$$

This is a matrix with n rows and p columns.

Set of matrices: denoted $\mathcal{M}_{n,p}(\mathbb{K})$.

Square matrices: $\mathcal{M}_{n,n}(\mathbb{K})$ is denoted by $\mathcal{M}_n(\mathbb{K})$

A matrix is a double entry table:

$$A = \begin{pmatrix} a_{1,1} & \dots & a_{1,p} \\ a_{2,1} & \dots & a_{2,p} \\ \vdots & & \vdots \\ a_{n,1} & \dots & a_{n,p} \end{pmatrix} = (a_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}}$$

This is a matrix with n rows and p columns.

Set of matrices: denoted $\mathcal{M}_{n,p}(\mathbb{K})$.

Square matrices: $\mathcal{M}_{n,n}(\mathbb{K})$ is denoted by $\mathcal{M}_n(\mathbb{K})$

What is the dimension of $\mathcal{M}_{n,p}(\mathbb{K})$?

A matrix is a double entry table:

$$A = \begin{pmatrix} a_{1,1} & \dots & a_{1,p} \\ a_{2,1} & \dots & a_{2,p} \\ \vdots & & \vdots \\ a_{n,1} & \dots & a_{n,p} \end{pmatrix} = (a_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}}$$

This is a matrix with n rows and p columns.

Set of matrices: denoted $\mathcal{M}_{n,p}(\mathbb{K})$.

Square matrices: $\mathcal{M}_{n,n}(\mathbb{K})$ is denoted by $\mathcal{M}_n(\mathbb{K})$

What is the dimension of $\mathcal{M}_{n,p}(\mathbb{K})$?

What is the canonical basis of $\mathcal{M}_{n,p}(\mathbb{K})$?

Matrices: Addition and multiplication

Let $A \in \mathcal{M}_{n,p}(\mathbb{K})$ and $B \in \mathcal{M}_{q,r}(\mathbb{K})$ then

- The **addition** $A + B$ is defined if and only if A and B have the same size $((n, p) = (q, r))$ and

$$A + B = (a_{i,j} + b_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}}$$

- The **multiplication** AB is defined if and only if $p = q$ and then

$$AB = (c_{i,k})_{\substack{1 \leq i \leq n \\ 1 \leq k \leq r}} \quad \text{where} \quad c_{i,k} = \sum_{j=1}^p a_{i,j} b_{j,k}.$$

Warning: In Python/numpy you can sum matrices with different sizes thanks to **Broadcasting Rules**

Remark: Matrix multiplication is natural since it corresponds to the composition of the associated linear maps.

For square matrices: In general $AB \neq BA$.

Matrices: Addition and multiplication

Let $A \in \mathcal{M}_{n,p}(\mathbb{K})$ and $B \in \mathcal{M}_{q,r}(\mathbb{K})$ then

- The **addition** $A + B$ is defined if and only if A and B have the same size $((n, p) = (q, r))$ and

$$A + B = (a_{i,j} + b_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}}$$

- The **multiplication** AB is defined if and only if $p = q$ and then

$$AB = (c_{i,k})_{\substack{1 \leq i \leq n \\ 1 \leq k \leq r}} \quad \text{where} \quad c_{i,k} = \sum_{j=1}^p a_{i,j} b_{j,k}.$$

Warning: In Python/numpy you can sum matrices with different sizes thanks to **Broadcasting Rules**

Remark: Matrix multiplication is natural since it corresponds to the composition of the associated linear maps.

For square matrices: In general $AB \neq BA$.

Computational cost of matrix multiplication?

Matrices: Addition and multiplication

Let $A \in \mathcal{M}_{n,p}(\mathbb{K})$ and $B \in \mathcal{M}_{q,r}(\mathbb{K})$ then

- The **addition** $A + B$ is defined if and only if A and B have the same size $((n, p) = (q, r))$ and

$$A + B = (a_{i,j} + b_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}}$$

- The **multiplication** AB is defined if and only if $p = q$ and then

$$AB = (c_{i,k})_{\substack{1 \leq i \leq n \\ 1 \leq k \leq r}} \quad \text{where} \quad c_{i,k} = \sum_{j=1}^p a_{i,j} b_{j,k}.$$

Warning: In Python/numpy you can sum matrices with different sizes thanks to **Broadcasting Rules**

Remark: Matrix multiplication is natural since it corresponds to the composition of the associated linear maps.

For square matrices: In general $AB \neq BA$.

Computational cost of matrix multiplication?

(for large matrices Strassen algorithm is faster)

Identity matrix: $I_n = \text{diag}(1, \dots, 1) = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \ddots & & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$

Properties:

- Associativity:

$$\forall A \in \mathcal{M}_{n,p}(\mathbb{K}), B \in \mathcal{M}_{p,q}(\mathbb{K}), C \in \mathcal{M}_{q,r}(\mathbb{K}), \quad (AB)C = A(BC)$$

- Distributivity: $\forall A, B \in \mathcal{M}_{n,p}(\mathbb{K}), C, D \in \mathcal{M}_{p,q}(\mathbb{K}),$

$$(A + B)C = AC + BC \quad \text{and} \quad A(C + D) = AC + AD$$

- Multiplication with the identity matrix:

$$\forall A \in \mathcal{M}_{n,p}(\mathbb{K}), \quad I_n A = A = A I_p.$$

Proposition

Let $A \in \mathcal{M}_n(\mathbb{K})$. There is equivalence between

- ❶ $\text{Ker}(A) = \{0\}$
- ❷ $\text{Im}(A) = \mathbb{K}^n$
- ❸ There exists $B \in \mathcal{M}_n(\mathbb{K})$ such that $AB = I_n$
- ❹ There exists $B \in \mathcal{M}_n(\mathbb{K})$ such that $BA = I_n$

If any of this holds, then A is said to be invertible and the matrix B is A^{-1} .

The set of invertible matrices is denoted by $GL_n(\mathbb{K})$.

$$AA^{-1} = I_n = A^{-1}A$$

$$(AB)^{-1} = B^{-1}A^{-1}$$

$$(A + B)^{-1} \neq A^{-1} + B^{-1}$$

Matrices: Transpose of a matrix

Definition

Let $A = (a_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}} \in \mathcal{M}_{n,p}(\mathbb{R})$. The transpose of A is the matrix $B = (b_{i,j})_{\substack{1 \leq i \leq p \\ 1 \leq j \leq n}} \in \mathcal{M}_{p,n}(\mathbb{R})$ such that

$$b_{i,j} = a_{j,i}.$$

The transpose is denoted $B = A^T$.

For $A = (a_{i,j})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}} \in \mathcal{M}_{n,p}(\mathbb{R})$, the conjugate transpose is such that $b_{i,j} = \bar{a}_{j,i}$ and is denoted A^* .

Given A , A^T is obtained by exchanging rows and columns:

$$\begin{pmatrix} 1 & 2 & 3 \\ 7 & 8 & 9 \end{pmatrix}^T = \begin{pmatrix} 1 & 7 \\ 2 & 8 \\ 3 & 9 \end{pmatrix}$$

$$(A^T)^T = A$$

$$\text{Linearity: } (A + B)^T = A^T + B^T$$

$$(AB)^T = B^T A^T$$

Same properties are valid for conjugate transpose A^* .

Matrices: Transpose of a matrix

$$(A^T)^T = A$$

$$\text{Linearity: } (A + B)^T = A^T + B^T$$

$$(AB)^T = B^T A^T$$

Same properties are valid for conjugate transpose A^* .

What are the missing quantifiers here?

Multiplication by blocks

For two integers $n, p \in \mathbb{N}^*$, let

- $A, E \in \mathcal{M}_n(\mathbb{K})$
- $B, F \in \mathcal{M}_{n,p}(\mathbb{K})$
- $C, G \in \mathcal{M}_{p,n}(\mathbb{K})$
- $D, H \in \mathcal{M}_p(\mathbb{K})$.

Then the product between block-matrices is obtained by applying the 2×2 matrix product formula:

$$\begin{pmatrix} A & B \\ C & D \end{pmatrix} \begin{pmatrix} E & F \\ G & H \end{pmatrix} = \begin{pmatrix} AE + BG & AF + BH \\ CE + DG & CF + DH \end{pmatrix}$$

Remarks:

- This generalizes to more than 2×2 blocks as long as blocks are compatible.
- Blocks often used to construct matrices by induction with blocks of size 1 and $n - 1$.

- The **trace** of a square matrix is the linear map $\mathcal{M}_n(\mathbb{K}) \rightarrow \mathbb{K}$ defined by

$$\mathrm{Tr}(A) = \sum_{k=1}^n a_{k,k}$$

(sum of the diagonal entries)

- Let $A \in \mathcal{M}_{n,p}(\mathbb{K})$ and $B \in \mathcal{M}_{p,n}(\mathbb{K})$. One has

$$\mathrm{Tr}(AB) = \mathrm{Tr}(BA).$$

Matrices: Diagonal matrices

- A matrix is diagonal if all its off-diagonal coefficients are 0.
- Subspace of dimension n :

$$\begin{aligned} \mathbb{K}^n &\rightarrow \mathcal{M}_n(\mathbb{K}) \\ a = \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix} &\mapsto \text{diag}(a) = \begin{pmatrix} a_1 & 0 & \dots & 0 \\ 0 & a_2 & \dots & 0 \\ \vdots & \ddots & & \vdots \\ 0 & 0 & \dots & a_n \end{pmatrix} \end{aligned}$$

- Matrix multiplication:

$$\text{diag}(a) \text{diag}(b) = \text{diag}(a \odot b) = \begin{pmatrix} a_1 b_1 & 0 & \dots & 0 \\ 0 & a_2 b_2 & \dots & 0 \\ \vdots & \ddots & & \vdots \\ 0 & 0 & \dots & a_n b_n \end{pmatrix}$$

Computing this product is $O(n)$ and not $O(n^3)$!

Matrices: Diagonal matrices

- A matrix is diagonal if all its off-diagonal coefficients are 0.
- Subspace of dimension n :

$$\begin{aligned}\mathbb{K}^n &\rightarrow \mathcal{M}_n(\mathbb{K}) \\ a = \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix} &\mapsto \text{diag}(a) = \begin{pmatrix} a_1 & 0 & \dots & 0 \\ 0 & a_2 & \dots & 0 \\ \vdots & \ddots & & \vdots \\ 0 & 0 & \dots & a_n \end{pmatrix}\end{aligned}$$

- Matrix multiplication:

$$\text{diag}(a) \text{diag}(b) = \text{diag}(a \odot b) = \begin{pmatrix} a_1 b_1 & 0 & \dots & 0 \\ 0 & a_2 b_2 & \dots & 0 \\ \vdots & \ddots & & \vdots \\ 0 & 0 & \dots & a_n b_n \end{pmatrix}$$

Computing this product is $O(n)$ and not $O(n^3)$!

Exercise: What about invertible diagonal matrices?

Matrices: Diagonal matrices

- A matrix is diagonal if all its off-diagonal coefficients are 0.
- Subspace of dimension n :

$$\mathbb{K}^n \rightarrow \mathcal{M}_n(\mathbb{K})$$
$$a = \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix} \mapsto \text{diag}(a) = \begin{pmatrix} a_1 & 0 & \dots & 0 \\ 0 & a_2 & \dots & 0 \\ \vdots & \ddots & & \vdots \\ 0 & 0 & \dots & a_n \end{pmatrix}$$

- Matrix multiplication:

$$\text{diag}(a) \text{diag}(b) = \text{diag}(a \odot b) = \begin{pmatrix} a_1 b_1 & 0 & \dots & 0 \\ 0 & a_2 b_2 & \dots & 0 \\ \vdots & \ddots & & \vdots \\ 0 & 0 & \dots & a_n b_n \end{pmatrix}$$

Computing this product is $O(n)$ and not $O(n^3)$!

Exercise: What about invertible diagonal matrices?

(linked to naive filter inversion in spectral domain)

- Upper triangular matrices: All elements below the diagonal are zero.
- Lower triangular matrices: All elements above the diagonal are zero.
- Solving a linear system with a triangular matrix is easy by variable elimination...

- Upper triangular matrices: All elements below the diagonal are zero.
- Lower triangular matrices: All elements above the diagonal are zero.
- Solving a linear system with a triangular matrix is easy by variable elimination...

Linear systems will be discussed later...

Recalls on determinant:

- $\det : \mathcal{M}_n(\mathbb{R}) \rightarrow \mathbb{R}$

$$\det(A) = \begin{vmatrix} a_{1,1} & \dots & a_{1,n} \\ \vdots & & \vdots \\ a_{n,1} & \dots & a_{n,n} \end{vmatrix}$$

- For 2×2 matrices:

$$A = \begin{vmatrix} a & b \\ c & d \end{vmatrix} = ad - cb.$$

- A is invertible if and only if $\det(A) \neq 0$.
- $\det(AB) = \det(A) \det(B)$
- For $A \in \mathcal{M}_n(\mathbb{R})$, $\det(A)$ is also the signed volume of the n -dimensional parallelepiped spanned by the column or row vectors of the matrix.

$$\det(A) = \det(A^T).$$

- The determinant of a triangular matrix is just the product of its diagonal coefficients (in particular for diagonal matrices).

Matrices: Orthogonal and unitary matrices

Recall that A^T is the transpose of A and A^* is the conjugate transpose of A (different when using complex numbers).

- A matrix $P \in \mathcal{M}_n(\mathbb{R})$ is said to be orthogonal if

$$P^T P = I_n$$

that is, P is invertible and $P^{-1} = P^T$.

- A matrix $U \in \mathcal{M}_n(\mathbb{C})$ is said to be unitary if

$$U^* U = I_n$$

that is, U is invertible and $U^{-1} = U^*$.

Matrices: Orthogonal and unitary matrices

Recall that A^T is the transpose of A and A^* is the conjugate transpose of A (different when using complex numbers).

- A matrix $P \in \mathcal{M}_n(\mathbb{R})$ is said to be orthogonal if

$$P^T P = I_n$$

that is, P is invertible and $P^{-1} = P^T$.

- A matrix $U \in \mathcal{M}_n(\mathbb{C})$ is said to be unitary if

$$U^* U = I_n$$

that is, U is invertible and $U^{-1} = U^*$.

- To go further we need the notion of Euclidean norm and inner product...

Definition

Given a vector space E over $\mathbb{K} = \mathbb{R}$ or $\mathbb{C}$ a norm $\|\cdot\|_N : E \rightarrow [0, +\infty)$ is a non-negative function such that

- ❶ For all $x \in E$, if $\|x\|_N = 0$ then $x = 0$ (*point-separation*)
- ❷ For all $x \in E$ and $\lambda \in \mathbb{K}$, $\|\lambda x\|_N = |\lambda| \|x\|_N$ (*absolutely homogeneous*)
- ❸ For all $x, y \in E$, $\|x + y\|_N \leq \|x\|_N + \|y\|_N$ (*triangle inequality*)

Examples:

- Euclidean norm: For $x \in \mathbb{R}^n$,

$$\|x\| = \sqrt{x_1^2 + \cdots + x_n^2} = \left(\sum_{k=1}^n x_k^2 \right)^{\frac{1}{2}}$$

- For $z \in \mathbb{C}^n$,

$$\|z\| = \sqrt{|z_1|^2 + \cdots + |z_n|^2} = \left(\sum_{k=1}^n |z_k|^2 \right)^{\frac{1}{2}}$$

- The Euclidean norm is also called the ℓ_2 -norm.
- For every $p \geq 1$, the ℓ_p -norm is

$$\|x\|_p = \left(\sum_{k=1}^n |x_k|^p \right)^{\frac{1}{p}}.$$

- For $p = \infty$,

$$\|x\|_\infty = \max_{k=1, \dots, n} |x_k|.$$

- The Euclidean norm is also called the ℓ_2 -norm.
- For every $p \geq 1$, the ℓ_p -norm is

$$\|x\|_p = \left(\sum_{k=1}^n |x_k|^p \right)^{\frac{1}{p}}.$$

- For $p = \infty$,

$$\|x\|_\infty = \max_{k=1, \dots, n} |x_k|.$$

Exercise: Let $x \in \mathbb{R}^n$ be fixed. Show that

$$\lim_{p \rightarrow +\infty} \|x\|_p = \|x\|_\infty.$$

Examples for matrices:

- For $A \in \mathcal{M}_{n,p}(\mathbb{C})$, norm induced by the Euclidean norm:

$$\|A\| = \sup_{\{x, \|x\| \leq 1\}} \|Ax\|.$$

This is a matrix norm: $\|AB\| \leq \|A\|\|B\|$.

- For $A \in \mathcal{M}_{n,p}(\mathbb{C})$ a matrix, the Frobenius norm:

$$\|A\|_F = \sqrt{\text{Tr}(AA^*)} = \sqrt{\text{Tr}(A^*A)} = \left(\sum_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}} |a_{i,j}|^2 \right)^{\frac{1}{2}}.$$

This is not a matrix norm.

Theorem

In finite dimension all norms are equivalent.

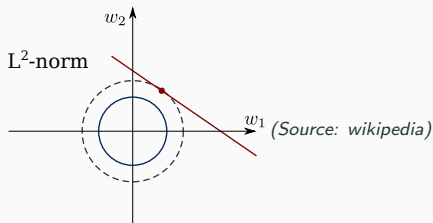
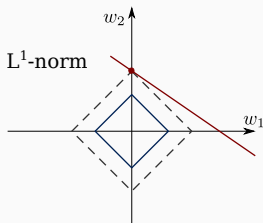
True but misleading for data science!

The choice of norms play a key role for data processing algorithms:

For example, for linear regression of noisy data $y \in \mathbb{R}^n$ with covariate matrix $X \in \mathcal{M}_{n,p}(\mathbb{R})$, the LASSO algorithm solves the optimization problem

$$\min_{\beta \in \mathbb{R}^p} \|y - X\beta\|_2^2 + \lambda \|\beta\|_1.$$

- Minimize squared ℓ_2 to fit data corrupted by Gaussian noise
- Minimize ℓ_1 to select solution that are sparse (as many zero coefficient as possible: simpler explanation model).



Inner product

Definition

Given a vector space E over $\mathbb{K} = \mathbb{R}$ or $\mathbb{C}$ an inner product $\langle \cdot, \cdot \rangle_P : E \times E \rightarrow \mathbb{K}$ is a function such that

- ❶ For all $x, y \in E$, $\langle x, y \rangle_P = \overline{\langle y, x \rangle_P}$. (*Hermitian symmetry*)
- ❷ For all $x \in E$, $y \mapsto \langle x, y \rangle_P$ is linear. (*right linearity*)
- ❸ For all $x \in E$, $\langle x, x \rangle_P \geq 0$ and if $\langle x, x \rangle_P = 0$ then $x = 0$. (*positive definite*)

Associated norm: $x \mapsto \sqrt{\langle x, x \rangle_P}$ is the norm associated to the inner product.

Euclidean inner product in $\mathbb{R}^n$: For $x, y \in \mathbb{R}^n$,

$$\langle x, y \rangle = \sum_{k=1}^n x_k y_k = x^T y.$$

Hermitian inner product in $\mathbb{C}^n$: For $x, y \in \mathbb{C}^n$,

$$\langle x, y \rangle = \sum_{k=1}^n \bar{x}_k y_k = x^* y.$$

- Two non zero vectors x, y are orthogonal if $\langle x, y \rangle = 0$.
- A family of non zeros vectors is said to be orthogonal if all pair of vectors are orthogonal.
- A family of orthogonal vectors is linearly independent.
- A basis $\mathcal{B} = \{x_1, \dots, x_n\}$ of $\mathbb{K}^n$ is orthonormal if

$$\forall 1 \leq i \neq j \leq n, \quad \|x_i\| = 1 \quad \text{and} \quad \langle x_i, x_j \rangle = 0.$$

- Two subspaces F_1 and F_2 are in direct orthogonal sum if $E = F_1 \oplus F_2$ and for all $x_1 \in F_1$, $x_2 \in F_2$, $\langle x_1, x_2 \rangle = 0$. One denotes $E = F_1 \overset{\perp}{\oplus} F_2$
- **Pythagoras's theorem:** For all $x \in E$ with decomposition $x = x_1 + x_2$ one has

$$\|x\|^2 = \|x_1\|^2 + \|x_2\|^2.$$

- $x \mapsto x_1$ is the orthogonal projection onto F_1 , similarly for F_2 .
- If $u_1, \dots, u_p$ are $p \leq n$ unit orthogonal vectors, then the orthogonal projection of x onto $F = \text{Span}(u_1, \dots, u_p)$ is

$$p_F(x) = \sum_{j=1}^p \langle x, u_j \rangle u_j.$$

Real case: Let $A \in \mathcal{M}_{n,p}(\mathbb{R})$, then

$$\forall x \in \mathbb{R}^p, y \in \mathbb{R}^n, \quad \langle Ax, y \rangle = (Ax)^T y = x^T A^T y = x^T (A^T y) = \langle x, A^T y \rangle.$$

Complex case: Let $A \in \mathcal{M}_{n,p}(\mathbb{C})$, then

$$\forall x \in \mathbb{C}^p, y \in \mathbb{C}^n, \quad \langle Ax, y \rangle = \langle x, A^* y \rangle.$$

Matrices: Orthogonal matrices

- A matrix $P \in \mathcal{M}_n(\mathbb{R})$ is said to be orthogonal if

$$P^T P = I_n$$

that is, P is invertible and $P^{-1} = P^T$.

- P is orthogonal if and only if its columns form an orthonormal basis of $\mathbb{R}^n$.
- Same for the rows.
- If P is orthogonal, the map $x \mapsto Px$ conserve the scalar product (and thus the norms, and thus the distances):

$$\langle Px, Py \rangle = \langle x, P^T Py \rangle = \langle x, y \rangle$$

- $x \mapsto Px$ is a rotation or symmetry + rotation map.

Numerically: Best matrices to invert!

- A matrix $A \in \mathcal{M}_n(\mathbb{R})$ is said to be *symmetric* if $A^T = A$, that is its coefficients above and below the diagonal are the same.

Definition (Positive matrices)

A symmetric matrix $A \in \mathcal{M}_n(\mathbb{R})$ is said to be *positive* if

$$\forall x \in \mathbb{R}^n \setminus \{0\}, \quad \langle x, Ax \rangle = x^T Ax = \sum_{i,j=1}^n a_{i,j} x_i x_j > 0.$$

A symmetric matrix $A \in \mathcal{M}_n(\mathbb{R})$ is said to be *non negative* if $\forall x \in \mathbb{R}^n$, $x^T Ax \geq 0$.

- A is positive if and only if $x \mapsto \langle x, Ax \rangle$ is a norm. It is the norm associated to A . Then, $(x, y) \mapsto \langle x, Ay \rangle$ is a scalar product.
- In probability, covariance matrices are non negative.

- A matrix $A \in \mathcal{M}_n(\mathbb{C})$ is said to be *Hermitian* if $A^* = A$, that is its coefficients above are the conjugates of the ones below.

Definition (Positive matrices)

An Hermitian matrix $A \in \mathcal{M}_n(\mathbb{C})$ is said to be *positive* if

$$\forall x \in \mathbb{C}^n \setminus \{0\}, \quad \langle x, Ax \rangle = x^* Ax = \sum_{i,j=1}^n a_{i,j} \bar{x}_i x_j > 0.$$

An Hermitian matrix $A \in \mathcal{M}_n(\mathbb{C})$ is said to be *non negative* if $\forall x \in \mathbb{C}^n$, $x^T Ax \geq 0$.

- A is positive if and only if $x \mapsto \langle x, Ax \rangle$ is a norm. It is the norm associated to A . Then, $(x, y) \mapsto \langle x, Ay \rangle$ is a scalar product.

Linear systems: Theory and algorithms

Square linear systems

Consider a square linear system

$$Ax = b$$

with

- $A \in \mathcal{M}_n(\mathbb{K})$ is a square matrix (known)
- $b \in \mathbb{K}^n$ is the “right-hand side” (known)
- $x \in \mathbb{K}^n$ is the **unknown** of the problem.

The goal is to “revert” the linear mixing induced by $x \mapsto Ax$.

Example:

$$\begin{cases} 2x_1 + x_3 &= 1 \\ x_1 + 2x_2 &= 3 \\ x_2 + 5x_3 &= 7 \end{cases} \Leftrightarrow \begin{pmatrix} 2 & 0 & 3 \\ 1 & 2 & 0 \\ 0 & 1 & 5 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 1 \\ 3 \\ 7 \end{pmatrix}$$

Consider a square linear system

$$Ax = b.$$

Theorem (Solutions of a square linear system)

There two cases:

- ① *If the matrix A is invertible, there exists a unique solution of the linear system $Ax = b$ which is $x = A^{-1}b$.*
- ② *If the matrix A is not invertible then one of the following alternatives holds:*
 - *either the right-hand side b belongs to the range of A and there exists an infinity of solutions that differ one from the other by addition of an element of the kernel of A ,*
 - *or the right-hand side b does not belong to the range of A and there are no solutions.*

Consider a rectangular linear system

$$Ax = b$$

with

- $A \in \mathcal{M}_{n,p}(\mathbb{K})$ is a rectangular matrix (known)
- $b \in \mathbb{K}^n$ is the “right-hand side” (known)
- $x \in \mathbb{K}^p$ is the **unknown** of the problem.

Two cases:

- When $n < p$, we say that the system is *underdetermined*: there are more unknowns than equations
- When $n > p$, we say that the system is *overdetermined*: there are fewer unknowns than equations.

Theorem (Solution of rectangular systems)

- **Existence:** *There exists at least one solution of linear system $Ax = b$ if and only if the right-hand side b belongs to $\text{Im}(A)$.*
- **Uniqueness:** *The solution is unique if and only if the kernel of A is reduced to the zero vector. Two solutions differ by an element of the kernel of A .*

Geometric point of view: The set of solution $S = \{x \in \mathbb{K}^n, Ax = b\}$ is either empty or the affine subspace $x_0 + \text{Ker}(A)$ where x_0 is one of the solution.

Remarks:

- If $n < p$, then $\dim(\text{Ker}(A)) \geq p - n \geq 1$, and if there exists a solution of the linear system, there exists an infinity of them.
- If there is no solution, one turns to best approximate solution using least-squares approximation:

$$\min_{x \in \mathbb{R}^n} \|Ax - b\|^2$$

- We now discuss algorithms to solve **square** linear systems with **invertible matrices**.
- **Goal:** Compute $x = A^{-1}b$ but **without computing the inverse matrix** A^{-1} .

Algorithm 1: Triangular system

- Solving a linear system with a triangular matrix is easy by variable elimination.
- **Details on whiteboard...**
- It requires only $O(n^2)$ operations (same cost as multiplying by A^{-1} even if it was known...).

Algorithm 2: LU factorization

Theorem (LU factorization)

Let $A = (a_{i,j})_{1 \leq i,j \leq n}$ be a matrix of order n such that all its diagonal submatrices of order $k = 1, \dots, n$

$$\Delta_k = (a_{i,j})_{1 \leq i,j \leq k}$$

are invertible. There exists a unique pair of matrices (L, U) , with U upper triangular and L lower triangular with a unit diagonal (i.e., $l_{i,i} = 1$), such that $A = LU$.

$$A = LU = \begin{pmatrix} 1 & 0 & \dots & 0 \\ l_{2,1} & 1 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ l_{n,1} & \dots & l_{n,n-1} & 1 \end{pmatrix} \begin{pmatrix} u_{1,1} & u_{1,2} & \dots & u_{1,n} \\ 0 & u_{2,2} & \dots & u_{2,n} \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & u_{n,n} \end{pmatrix}$$

Algorithm 2: LU factorization

- Computing the LU factorization costs $\sim \frac{n^3}{3}$ operations.

Solve linear system using LU factorization:

$$Ax = b \quad \Leftrightarrow \quad LUx = b$$

So why is it useful?

Algorithm 2: LU factorization

- Computing the LU factorization costs $\sim \frac{n^3}{3}$ operations.

Solve linear system using LU factorization:

$$Ax = b \quad \Leftrightarrow \quad LUx = b$$

So why is it useful?

- ① Solve $Ly = b$ (triangular system)
- ② Solve $Ux = y$ (triangular system)

In the end one has

$$x = U^{-1}y = U^{-1}L^{-1}b = (LU)^{-1}b = A^{-1}b.$$

- If you have several systems to solve with A it is worthwhile to store the LU factorization once and for all.

Algorithm 2: LU factorization

- Computing the LU factorization costs $\sim \frac{n^3}{3}$ operations.

Solve linear system using LU factorization:

$$Ax = b \quad \Leftrightarrow \quad LUx = b$$

So why is it useful?

- ① Solve $Ly = b$ (triangular system)
- ② Solve $Ux = y$ (triangular system)

In the end one has

$$x = U^{-1}y = U^{-1}L^{-1}b = (LU)^{-1}b = A^{-1}b.$$

- If you have several systems to solve with A it is worthwhile to store the LU factorization once and for all.

Determinant computation using LU factorization:

Exercise: Express $\det(A)$ in function of its LU factorization.

Algorithm 3: Cholesky factorization

Hypothesis on A : A is a real symmetric positive matrix.

Theorem (Cholesky factorization)

Let A be a real symmetric positive matrix. There exists a unique real lower triangular matrix L , having positive diagonal entries, such that

$$A = LL^T.$$

Linear systems: Same as for LU factorization.

Other applications: Useful for norm associated to A and A^{-1} :

$$\langle x, Ax \rangle = \langle x, LL^T x \rangle = \langle L^T x, L^T x \rangle = \|L^T x\|^2.$$

$$\langle x, A^{-1}x \rangle = \langle x, (LL^T)^{-1}x \rangle = \langle x, (L^T)^{-1}L^{-1}x \rangle = \|L^{-1}x\|^2.$$

Algorithm 4: QR factorization

Theorem (QR factorization)

Let A be a real nonsingular matrix. There exists a unique pair (Q, R) , where Q is an orthogonal matrix and R is an upper triangular matrix with positive diagonal entries, satisfying

$$A = QR.$$

- Proof is based on the Gram–Schmidt orthonormalization process.

Solve linear system using LU factorization:

$$Ax = b \quad \Leftrightarrow \quad QRx = b \quad \Leftrightarrow \quad Rx = Q^T b$$

- ① Update the right-hand side: compute $Q^T b$
- ② Solve the triangular system $Rx = Q^T b$

Spectral Theory of Matrices

Change of basis for vectors

Let E be a vector space of dimension n (e.g. $\mathbb{K}^n$). Let us consider **two** different basis

$$\mathcal{B} = (e_1, \dots, e_n) \quad \text{and} \quad \mathcal{B}' = (e'_1, \dots, e'_n).$$

Definition

The **change of basis** from $\mathcal{B}$ to $\mathcal{B}'$ is the matrix

$$P_{\mathcal{B} \rightarrow \mathcal{B}'} = \text{Mat}(\text{id}_E, \mathcal{B}', \mathcal{B}),$$

that is the matrix whose j -th column are the coordinates of e'_j in the basis $\mathcal{B}$.

- If $x \in E$ has coordinates $X \in \mathbb{K}^n$ in the basis $\mathcal{B}$ and $X' \in \mathbb{K}^n$ in the basis $\mathcal{B}'$, then

$$X = P_{\mathcal{B} \rightarrow \mathcal{B}'} X'.$$

- By composition,

$$P_{\mathcal{B} \rightarrow \mathcal{B}'} P_{\mathcal{B}' \rightarrow \mathcal{B}} = I_n \quad \Leftrightarrow \quad P_{\mathcal{B} \rightarrow \mathcal{B}'} = P_{\mathcal{B}' \rightarrow \mathcal{B}}^{-1}.$$

Change of basis for endomorphisms

Given a linear map $U : E \rightarrow E$ and the same two basis, we have two matrices:

$$A = \text{Mat}(U, \mathcal{B}) \quad \text{and} \quad A' = \text{Mat}(U, \mathcal{B}').$$

Relation between the matrices:

$$\text{Mat}(U, \mathcal{B}') = P_{\mathcal{B}' \rightarrow \mathcal{B}} \text{Mat}(U, \mathcal{B}) P_{\mathcal{B} \rightarrow \mathcal{B}'}$$

that is

$$A' = P_{\mathcal{B} \rightarrow \mathcal{B}'}^{-1} A P_{\mathcal{B} \rightarrow \mathcal{B}'}$$

Proof: Denote $x \in E$, $y = U(x)$ and X, X' and Y, Y' the respective coordinates. We have

$$Y' = A' X'$$

but also

$$Y' = P_{\mathcal{B} \rightarrow \mathcal{B}'}^{-1} Y = P_{\mathcal{B} \rightarrow \mathcal{B}'}^{-1} A X = P_{\mathcal{B} \rightarrow \mathcal{B}'}^{-1} A P_{\mathcal{B} \rightarrow \mathcal{B}'} X'$$

Hence for all X' ,

$$A' X' = P_{\mathcal{B} \rightarrow \mathcal{B}'}^{-1} A P_{\mathcal{B} \rightarrow \mathcal{B}'} X'.$$

- The columns of an invertible matrices form a basis of $\mathbb{R}^n$.
- Given a matrix A (representing an endomorphism U in the canonical basis), all the matrices that represent U are

$$\{P^{-1}AP, P \text{ invertible}\}$$

where the basis vectors $(e'_1, \dots, e'_n)$ are the column of P .

- The question of matrix reduction is: Is there a basis $\mathcal{B}' =$ invertible matrix P such that $P^{-1}AP$ is “simple”.
- “Simple” generally means diagonal, but may be upper triangular...

Eigenvalues and eigenvectors

Let $A \in \mathcal{M}_n(\mathbb{K})$ be a square matrix.

If $\lambda \in \mathbb{C}$ and $v \in \mathbb{K}^n$, $v \neq 0$, are such that

$$Av = \lambda v$$

then λ is an *eigenvalue* and v is an *eigenvector*.

Definition

- $\lambda \in \mathbb{K}$ is an *eigenvalue* of A if $\text{Ker}(A - \lambda I_n) \neq \{0\}$. Then,
- $\text{Ker}(A - \lambda I_n)$ is the *eigensubspace* associated with λ .
- A **non zero** vector $v \in \text{Ker}(A - \lambda I_n)$, that is such that $Av = \lambda v$ is an *eigenvector* of A associated with λ .
- The *multiplicity* of λ is the $\dim(\text{Ker}(A - \lambda I_n))$. λ is said *simple* if the multiplicity is one.

- Eigenvalues can be complex even for real matrices.

- Eigenvalues are the roots of the characteristic polynomial

$$\chi_A(X) = \det(A - XI_n)$$

- If A is a square matrix with eigenvalues λ_i , $i = 1, \dots, n$

$$\operatorname{Tr}(A) = \sum_{i=1}^n \lambda_i$$

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

If there exists a basis $\mathcal{B}'$ formed of eigenvectors $(e'_j)_{1 \leq j \leq n}$ of A associated with λ_j then

$$A = PDP^{-1} \quad \text{with} \quad D = \text{diag}(\lambda_1, \dots, \lambda_n).$$

We say that A is diagonalizable.

In general, not every matrix is diagonalizable...

In general, not every matrix is diagonalizable...

- **Definition:** A matrix $A \in \mathcal{M}_n(\mathbb{C})$ is normal if it commutes with its conjugate matrix, that is,

$$AA^* = A^*A.$$

Theorem (Diagonalization)

A matrix $A \in \mathcal{M}_n(\mathbb{C})$ is normal (i.e. $AA^ = A^*A$) if and only if there exists a unitary matrix U such that*

$$A = U \operatorname{diag}(\lambda_1, \dots, \lambda_n) U^*$$

where $(\lambda_1, \dots, \lambda_n)$ are the eigenvalues of A .

- For real matrices, still need to use complex unitary basis to diagonalize (e.g. convolution in Fourier basis).

Theorem (Diagonalization of symmetric matrices)

A matrix $A \in \mathcal{M}_n(\mathbb{R})$ is real symmetric ($A^T = A$) if and only if there exists an **orthogonal** matrix Q ($Q^T = Q^{-1}$) and **real eigenvalues** $(\lambda_1, \dots, \lambda_n)$ such that

$$A = Q \operatorname{diag}(\lambda_1, \dots, \lambda_n) Q^T$$

Diagonalization: Real symmetric case

Theorem (Diagonalization of symmetric matrices)

A matrix $A \in \mathcal{M}_n(\mathbb{R})$ is real symmetric ($A^T = A$) if and only if there exists an **orthogonal** matrix Q ($Q^T = Q^{-1}$) and **real eigenvalues** $(\lambda_1, \dots, \lambda_n)$ such that

$$A = Q \operatorname{diag}(\lambda_1, \dots, \lambda_n) Q^T$$

Diagonalization: Real symmetric case

Theorem (Diagonalization of symmetric matrices)

A matrix $A \in \mathcal{M}_n(\mathbb{R})$ is real symmetric ($A^T = A$) if and only if there exists an **orthogonal** matrix Q ($Q^T = Q^{-1}$) and **real eigenvalues** $(\lambda_1, \dots, \lambda_n)$ such that

$$A = Q \operatorname{diag}(\lambda_1, \dots, \lambda_n) Q^T$$

Exercise: Show that the j -th column of Q is an eigenvector for the eigenvalue λ_j .

Diagonalization: Real symmetric case

Theorem (Diagonalization of symmetric matrices)

A matrix $A \in \mathcal{M}_n(\mathbb{R})$ is real symmetric ($A^T = A$) if and only if there exists an **orthogonal** matrix Q ($Q^T = Q^{-1}$) and **real eigenvalues** $(\lambda_1, \dots, \lambda_n)$ such that

$$A = Q \operatorname{diag}(\lambda_1, \dots, \lambda_n) Q^T$$

Exercise: Show that the j -th column of Q is an eigenvector for the eigenvalue λ_j .

- We say that A is diagonalizable in an orthonormal basis with real eigenvalues.
- If $(u_1, \dots, u_n)$ are the corresponding basis of eigenvectors, then for all $x \in \mathbb{R}^n$,

$$x = \sum_{j=1}^n \langle x, u_j \rangle u_j \quad \text{and} \quad Ax = \sum_{j=1}^n \lambda_j \langle x, u_j \rangle u_j$$

Exercise: Show that A a real symmetric matrix is positive if and only if all its eigenvalues are positive

Theorem (Diagonalization of symmetric matrices)

A matrix $A \in \mathcal{M}_n(\mathbb{R})$ is Hermitian ($A = A^*$) if and only if there exists a unitary matrix U ($U^* = U^{-1}$) and real eigenvalues $(\lambda_1, \dots, \lambda_n)$ such that

$$A = U \operatorname{diag}(\lambda_1, \dots, \lambda_n) U^*$$

- Eigenvalues and eigenvectors are only defined for square matrices.
- Given a rectangular matrix $A \in \mathcal{M}_{m,n}(\mathbb{C})$ we can use the eigenvalues of associated square matrices to decompose the matrix...
- A^*A is Hermitian and has real, non negative eigenvalues (i.e. A^*A is non negative).
- We call singular values of a matrix $A \in \mathcal{M}_{m,n}(\mathbb{C})$ are the non negative square roots of the n eigenvalues of A^*A .
- The non zero eigenvalues of A^*A are the same as the ones of AA^* .
- Hence A has necessarily $r \leq \min(m, n)$ positive singular values.

Singular value decomposition (SVD) of a matrix

Theorem (SVD factorization)

Let $A \in \mathcal{M}_{m,n}(\mathbb{C})$ be a matrix having r positive singular values

$$\mu_1 \geq \mu_2 \geq \cdots \geq \mu_r > 0.$$

There exists two unitary matrices $U \in \mathcal{M}_n(\mathbb{C})$, $V \in \mathcal{M}_m(\mathbb{C})$, such that

$$A = V \begin{pmatrix} \Sigma & 0 \\ 0 & 0 \end{pmatrix} U^*$$

where $\Sigma = \text{diag}(\mu_1, \mu_2, \dots, \mu_r)$.

Denoting (u_j) and (v_j) the columns of U and V , A is written as the sum of r rank one matrices:

$$A = \sum_{j=1}^r \mu_j v_j u_j^*.$$

Proposition (SVD and low-rank approximation)

Let

$$A = V \begin{pmatrix} \Sigma & 0 \\ 0 & 0 \end{pmatrix} U^* = \sum_{j=1}^r \mu_j v_j u_j^*$$

be the SVD factorization of A (with the ordered singular values). Then for $1 \leq k < r$, the matrix

$$A_k = \sum_{j=1}^k \mu_j v_j u_j^*$$

is the best approximation of A by matrices of rank k , in the following sense: for all matrices $X \in \mathcal{M}_{m,n}(\mathbb{C})$ of rank less than k we have

$$\|A - A_k\| \leq \|A - X\| \quad (\text{induced norm}),$$

and this minimal error is $\|A - A_k\| = \mu_{k+1}$.