

Numerical Computing

Lecture 4: Linear Algebra Foundations

Francisco Richter Mendoza

Università della Svizzera Italiana

Faculty of Informatics

Lugano, Switzerland

Today's Topics

- ▶ Vector and matrix norms
- ▶ Condition numbers
- ▶ Sensitivity analysis
- ▶ Special matrix classes
- ▶ Singular Value Decomposition
- ▶ Applications and examples

Vector Norms

Definition

Function $\|\cdot\| : \mathbb{R}^n \rightarrow \mathbb{R}$ is a **vector norm** if:

1. **Positivity:** $\|x\| \geq 0$, with $\|x\| = 0 \iff x = 0$
2. **Homogeneity:** $\|\alpha x\| = |\alpha| \|x\|$
3. **Triangle inequality:** $\|x + y\| \leq \|x\| + \|y\|$

Common p -norms

$$\|x\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{1/p}$$

- ▶ $\|x\|_1 = \sum_{i=1}^n |x_i|$ (**Manhattan**)
- ▶ $\|x\|_2 = \sqrt{\sum_{i=1}^n x_i^2}$ (**Euclidean**)
- ▶ $\|x\|_\infty = \max_{1 \leq i \leq n} |x_i|$ (**Maximum**)

Matrix Norms

Induced Matrix Norm

$$\|A\| = \max_{x \neq 0} \frac{\|Ax\|}{\|x\|} = \max_{\|x\|=1} \|Ax\|$$

Common Matrix Norms

- ▶ $\|A\|_1 = \max_j \sum_{i=1}^m |a_{ij}|$ (max column sum)
- ▶ $\|A\|_\infty = \max_i \sum_{j=1}^n |a_{ij}|$ (max row sum)
- ▶ $\|A\|_2 = \sigma_{\max}(A)$ (largest singular value)
- ▶ $\|A\|_F = \sqrt{\sum_{i,j} |a_{ij}|^2}$ (Frobenius)

Key property: $\|Ax\| \leq \|A\| \|x\|$ for induced norms

Linear Algebra Foundations

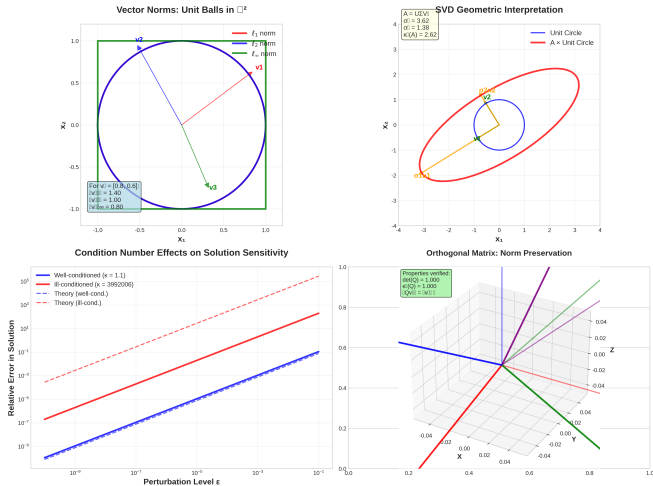


Figure: Vector norms unit balls, SVD geometric interpretation, condition number effects, and orthogonal matrix properties demonstrating fundamental linear algebra concepts.

Condition Numbers

Definition

For invertible matrix A :

$$\kappa(A) = \|A\| \|A^{-1}\|$$

For 2-norm: $\kappa_2(A) = \frac{\sigma_{\max}(A)}{\sigma_{\min}(A)}$

Properties

- ▶ $\kappa(A) \geq 1$ always
- ▶ $\kappa(I) = 1$ (identity perfectly conditioned)
- ▶ $\kappa(\alpha A) = \kappa(A)$ for $\alpha \neq 0$
- ▶ $\kappa_2(Q) = 1$ for orthogonal Q

Interpretation: Measures how "close" matrix is to being singular

Sensitivity Analysis

Linear System Perturbation

For $Ax = b$ with perturbed RHS $b + \Delta b$:

$$\frac{\|\Delta x\|}{\|x\|} \leq \kappa(A) \frac{\|\Delta b\|}{\|b\|}$$

Matrix Perturbation

For $(A + \Delta A)(x + \Delta x) = b$:

$$\frac{\|\Delta x\|}{\|x\|} \approx \kappa(A) \frac{\|\Delta A\|}{\|A\|}$$

Key insight: Condition number bounds relative error amplification

Example: $\kappa = 10^{12}$ means 10^{-16} perturbation $\Rightarrow 10^{-4}$ error

Matrix Conditioning Analysis

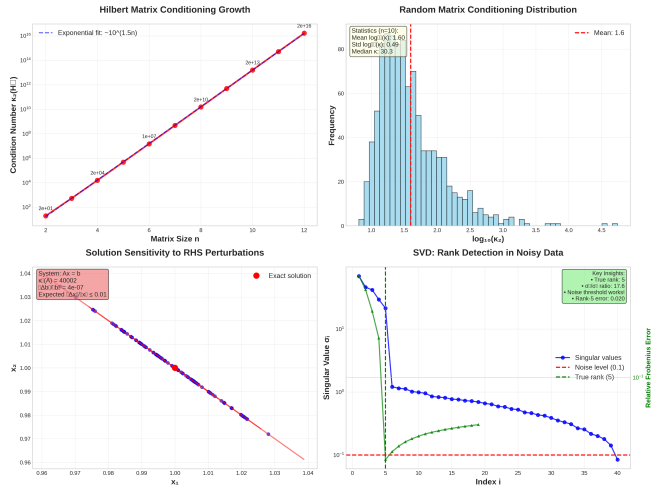


Figure: Hilbert matrix conditioning growth, random matrix statistics, perturbation analysis visualization, and SVD rank detection in noisy data.

Orthogonal Matrices

Definition

Matrix Q is **orthogonal** if $Q^T Q = I$, i.e., $Q^T = Q^{-1}$

Properties

- ▶ $\|Qx\|_2 = \|x\|_2$ (**preserves Euclidean norm**)
- ▶ $\det(Q) = \pm 1$
- ▶ $\kappa_2(Q) = 1$ (**perfectly conditioned**)
- ▶ $(Qx) \cdot (Qy) = x \cdot y$ (**preserves inner products**)

Geometric interpretation: Orthogonal transformations are rotations and reflections

Applications: QR decomposition, least squares, eigenvalue algorithms

Symmetric Matrices

Definition

Matrix A is **symmetric** if $A = A^T$

Spectral Theorem

Every symmetric matrix A has:

- ▶ Real eigenvalues: $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$
- ▶ Orthogonal eigenvectors: $A = Q\Lambda Q^T$
- ▶ $\kappa_2(A) = \frac{|\lambda_{\max}|}{|\lambda_{\min}|}$

Positive Definite

A is **positive definite** if $x^T A x > 0$ for all $x \neq 0$

- ▶ Equivalent: all eigenvalues > 0
- ▶ Cholesky decomposition: $A = LL^T$

Special Matrix Classes

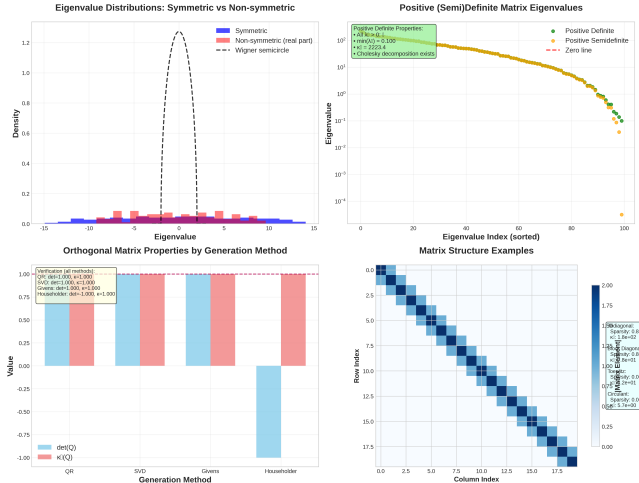


Figure: Symmetric matrix eigenvalue distributions, positive definite properties, orthogonal matrix generation methods, and structured matrix sparsity patterns.

Singular Value Decomposition (SVD)

Theorem

Every matrix $A \in \mathbb{R}^{m \times n}$ has SVD:

$$A = U \Sigma V^T$$

where:

- ▶ $U \in \mathbb{R}^{m \times m}$ orthogonal (left singular vectors)
- ▶ $V \in \mathbb{R}^{n \times n}$ orthogonal (right singular vectors)
- ▶ $\Sigma \in \mathbb{R}^{m \times n}$ diagonal with $\sigma_1 \geq \sigma_2 \geq \dots \geq 0$

Key Properties

- ▶ $\|A\|_2 = \sigma_1$, $\|A\|_F = \sqrt{\sum_i \sigma_i^2}$
- ▶ $\text{rank}(A)$ = number of nonzero singular values
- ▶ $\kappa_2(A) = \sigma_1 / \sigma_r$ where $r = \text{rank}(A)$

SVD Applications

Low-Rank Approximation

Best rank- k approximation: $A_k = \sum_{i=1}^k \sigma_i u_i v_i^T$

Error: $\|A - A_k\|_2 = \sigma_{k+1}$

Principal Component Analysis

Data matrix X : SVD reveals principal directions of variation

Pseudoinverse

For $A = U\Sigma V^T$: $A^+ = V\Sigma^+ U^T$

where Σ^+ inverts nonzero singular values

Applications: Image compression, data analysis, least squares

SVD Applications

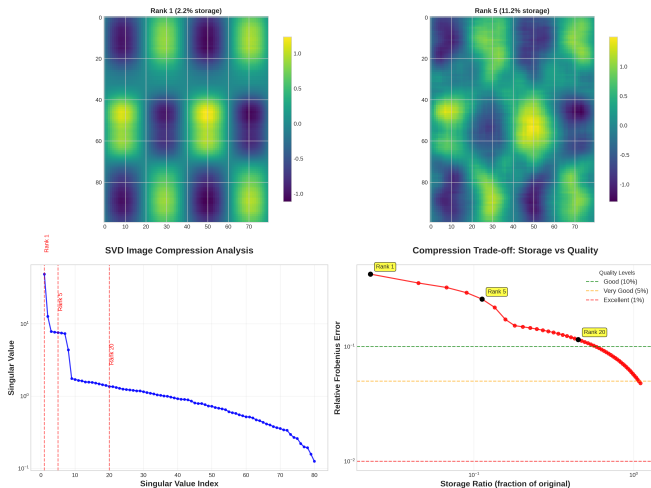


Figure: SVD image compression demonstration showing original vs compressed images, singular value decay, and storage-quality trade-offs for different compression ranks.

Computational Considerations

Matrix-Vector Operations

- ▶ Matrix-vector product: $O(mn)$ operations
- ▶ Matrix-matrix product: $O(mnp)$ operations
- ▶ Matrix inversion: $O(n^3)$ operations

Conditioning and Accuracy

- ▶ Well-conditioned: $\kappa \approx 1$ to 10^3
- ▶ Moderately conditioned: $\kappa \approx 10^3$ to 10^{12}
- ▶ Ill-conditioned: $\kappa > 10^{12}$ (problematic in double precision)

Rule of thumb: Lose $\log_{10}(\kappa)$ digits of accuracy

Practical Examples

Hilbert Matrix

$$H_{ij} = \frac{1}{i+j-1}: \kappa_2(H_{10}) \approx 1.6 \times 10^{13}$$

Lesson: Some problems are inherently ill-conditioned

Vandermonde Matrix

$V_{ij} = x_i^{j-1}$: Condition number grows exponentially with n

Application: Polynomial interpolation

Random Matrices

Typical condition number: $\kappa \approx 10^1$ to 10^3

Insight: Most matrices are reasonably well-conditioned

Key Takeaways

- ▶ **Norms:** Measure size of vectors and matrices
- ▶ **Condition numbers:** Quantify sensitivity to perturbations
- ▶ **Orthogonal matrices:** Preserve geometry, perfectly conditioned
- ▶ **Symmetric matrices:** Real eigenvalues, orthogonal eigenvectors
- ▶ **SVD:** Universal tool for matrix analysis and approximation
- ▶ **Computational cost:** $O(n^3)$ for most dense linear algebra

Next Lecture: Direct Methods for Linear Systems