

CSC338 WINTER 2022

WEEK 4 - SYSTEMS OF LINEAR EQUATIONS

Ilir Dema

University of Toronto

Feb 4, 2022

WHAT IS VECTOR NORM?

DEFINITION

- Let V be a vector space. **Norm** on V is a map $\|\cdot\| : V \rightarrow \mathbb{R}^{\geq 0}$ satisfying the following properties:
 - (I) $\forall v \in V, \|v\| \geq 0$ and $\|v\| = 0$ if and only if $v = 0$
 - (II) $\forall v \in V, \forall a \in \mathbb{R}, \|av\| = |a|\|v\|$
 - (III) $\forall v, w \in V, \|v + w\| \leq \|v\| + \|w\|$

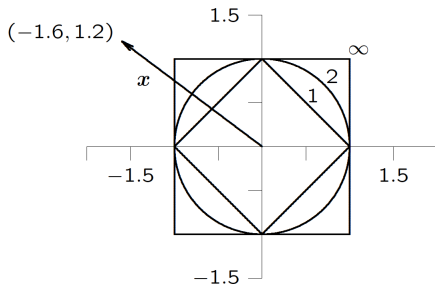
EXAMPLES

Let $V = \mathbb{R}^n, v = (x_1, \dots, x_n)$.

- $\|v\|_1 = \sum_{i=1}^n |x_i|$ (1-norm)
- $\|v\|_2 = (\sum_{i=1}^n |x_i|^2)^{1/2}$ (2-norm)
- $\|v\|_\infty = \max_{1 \leq i \leq n} |x_i|$ (∞ -norm)
- Exercise: Prove that the above three are indeed norms.
- Exercise: Prove that (in general): $|||u|| - ||v||| \leq \|u - v\|$

GEOMETRICAL INTERPREATION

Drawing shows unit sphere in two dimensions for each of these norms:



Norms have following values for vector shown:

$$\|x\|_1 = 2.8, \quad \|x\|_2 = 2.0, \quad \|x\|_\infty = 1.6$$

WHAT IS MATRIX NORM?

DEFINITION

- Let V be a vector space of dimension n and $A \in M_{n \times n}$ a matrix, representing some linear transformation of V .
- Definition

$$\|A\| = \max_{x \neq 0, x \in V} \frac{\|Ax\|}{\|x\|}$$

- Exercise: Prove that the defined matrix norm satisfies the properties of norms.
- Exercise: Prove that for any two $n \times n$ matrices, $\|AB\| \leq \|A\| \|B\|$.
- Exercise: Prove that for any matrix A and any vector x , $\|Ax\| \leq \|A\| \|x\|$.

EQUIVALENCY OF NORMS

DEFINITION

- Two norms $\|\cdot\|_1, \|\cdot\|_2$ are called equivalent if there exist positive numbers a, b such that $\|x\|_1 \leq a\|x\|_2$ and $\|x\|_2 \leq b\|x\|_1$ for any vector x .

EXERCISES

- Prove equivalency of $\|\cdot\|_1, \|\cdot\|_2, \|\cdot\|_\infty$ norms in $\mathbb{R}^n$.
- Let A be matrix and $\|A\|_1 = \max_j \sum_{i=1}^n |a_{ij}|, \|A\|_\infty = \max_i \sum_{j=1}^n |a_{ij}|$. Are they equivalent? Also, are they equivalent to the previously defined matrix norm?
- Exercise: Prove that $\|A\| = \max_{\|x\|=1} \|Ax\|$.

ILL-CONDITIONED SYSTEMS

EXAMPLE

- When there is a small change in one or more coefficients in a system, if the system is well conditioned, the change in the solution will also be small
- In the case of ill conditioned systems, a small change in some of the coefficients, will result in large changes in the solution.

$$\begin{bmatrix} 1 & 2 \\ 0.48 & 0.99 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 3 \\ 1.47 \end{bmatrix} \implies \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \end{bmatrix}$$

$$\begin{bmatrix} 1 & 2 \\ 0.49 & 0.99 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 3 \\ 1.47 \end{bmatrix} \implies \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 3 \\ 0 \end{bmatrix}$$

- Q: What is an adequate measure of size for linear systems?

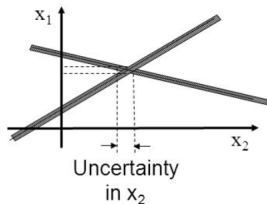
GEOMETRICAL INTERPRETATION

Consider the graphical interpretation for a 2-equation system:

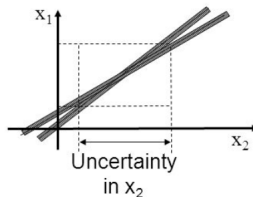
$$\begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \begin{Bmatrix} x_1 \\ x_2 \end{Bmatrix} = \begin{Bmatrix} b_1 \\ b_2 \end{Bmatrix} \quad (1)$$

$$(2)$$

We can plot the two linear equations on a graph of x_1 vs. x_2 .



Well-conditioned



Ill-conditioned

CONDITIONING NUMBER OF A MATRIX

DEFINITION

- For nonsingular A , $\text{cond}(A) = \|A\| \|A^{-1}\|$.
- For singular A , $\text{cond}(A) = \infty$.

PROPERTIES OF $\text{cond}(A)$

- For any matrix A , $\text{cond}(A) \geq 1$.
- $\text{cond}(I) = 1$
- For any nonzero number a , $\text{cond}(aA) = \text{cond}(A)$
- For a diagonal matrix A , $\text{cond}(A) = (\max_i |a_{ii}|) / (\min_i |a_{ii}|)$
- $\|A\| \|A^{-1}\| = (\max_{\|x\|=1} \|Ax\|) (\min_{\|x\|=1} \|Ax\|)^{-1}$, therefore the conditioning number of a matrix describes the ratio to max stretch that A inflicts on unit vectors versus min stretch that A inflicts on unit vectors.

ESTIMATING ERRORS

DEFINITION

- Let x be the true value of a vector, and $\hat{x}$ an approximate value of it. Recall that absolute error of x is the difference $\Delta x = \hat{x} - x$.
- Note that $\|\Delta x\| = \|x - \hat{x}\|$.
- Relative error $\delta x = \frac{\Delta x}{\|x\|}$ will be defined only for nonzero x .

THEOREM

Let $Ax = b$ be a linear system. Then the following hold:

$$\|\delta x\| \leq \text{cond}(A) \|\delta b\|$$

$$\|\delta x\| \leq \text{cond}(A) \|\delta A\|$$

FULL BACKWARD ERROR

ANALYSIS

Let $\mathbf{Ax} = \mathbf{b}$, $\hat{\mathbf{A}}\hat{\mathbf{x}} = \hat{\mathbf{b}}$, $\hat{\mathbf{A}} = \mathbf{A} + \Delta\mathbf{A}$, $\hat{\mathbf{x}} = \mathbf{x} + \Delta\mathbf{x}$, $\hat{\mathbf{b}} = \mathbf{b} + \Delta\mathbf{b}$.

It's not hard to see that

$$\Delta\mathbf{Ax} + \mathbf{A}\Delta\mathbf{x} + \Delta\mathbf{A}\Delta\mathbf{x} = \Delta\mathbf{b}$$

hence ingoring second degree differences, we get the linear relationship

$$\mathbf{A}\Delta\mathbf{x} = \Delta\mathbf{b} - \Delta\mathbf{Ax}.$$

Since $\mathbf{A}$ is invertible,

$$\Delta\mathbf{x} = \mathbf{A}^{-1}(\Delta\mathbf{b} - \Delta\mathbf{Ax}).$$

MATRIX CONDITIONING NUMBER

FACTS

- Conditioning number of a matrix is not easy to compute because it involves inverting the matrix (a $O(n^3)$ complexity operation) and computing the norm of both.
- An alternative way (same complexity though!) is to obtain singular value decomposition $\mathbf{A} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}$ where $\mathbf{U}, \mathbf{V}$ are orthogonal (that is $\det(\mathbf{U}), \det(\mathbf{V}) = \pm 1$) and $\mathbf{\Sigma} = \text{diag}(\sigma_i)$. Then $\text{cond}(\mathbf{A}) = \sigma_1 / \sigma_n$.
- An evaluation of $\text{cond}(\mathbf{A})$ can be obtained by observing that if $\mathbf{Ax} = \mathbf{y}$ then $\frac{\|\mathbf{x}\|}{\|\mathbf{y}\|} \leq \|\mathbf{A}^{-1}\|$. The equality can be achieved for a careful choice of $\mathbf{x}, \mathbf{y}$. Usually $\mathbf{y}$ is a vector with components ± 1 with signs chosen such that $\|\mathbf{x}\|$ is maximized.
- Software packages have functions to compute it, e.g. `numpy.linalg.cond`.

THE RESIDUAL

WHAT IS IT?

- Suppose a numerical algorithm for $\mathbf{Ax} = \mathbf{b}$ yields $\hat{\mathbf{x}}$.
- One way to verify our solution, is to plug it in the original system, obtaining $\hat{\mathbf{b}} = \mathbf{A}\hat{\mathbf{x}}$ and evaluate $\mathbf{r} = \mathbf{b} - \mathbf{A}\hat{\mathbf{x}}$.
- Some computation (please do it on your own in the form of an exercise!) shows that

$$\frac{\|\Delta \mathbf{x}\|}{\|\hat{\mathbf{x}}\|} \leq \text{cond}(\mathbf{A}) \frac{\|\mathbf{r}\|}{\|\mathbf{A}\| \|\hat{\mathbf{x}}\|} \leq \text{cond}(\mathbf{A}) \frac{\|\Delta \mathbf{A}\|}{\|\mathbf{A}\|}$$

where $\Delta \mathbf{A}$ is the backward error on $\mathbf{A}$.

- This shows residual is good if $\mathbf{A}$ is well conditioned.

RESIDUAL IS NO GOOD: EXAMPLE

THE RESIDUAL NOT GOOD IF MATRIX IS NOT WELL CONDITIONED

$$\mathbf{Ax} = \begin{bmatrix} 0.913 & 0.659 \\ 0.457 & 0.330 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 0.254 \\ 0.127 \end{bmatrix} = \mathbf{b}$$

Consider two approximate solutions

$$\hat{\mathbf{x}}_1 = \begin{bmatrix} 0.6391 \\ -0.5 \end{bmatrix}, \hat{\mathbf{x}}_2 = \begin{bmatrix} 0.999 \\ -1.001 \end{bmatrix}$$

The norms of their residuals are

$$\|\mathbf{r}_1\| = 7.0 \times 10^{-5}, \|\mathbf{r}_2\| = 2.4 \times 10^{-2}$$

The real solution is $\begin{bmatrix} 1 \\ -1 \end{bmatrix}$. What can we conclude?

ITERATIVE REFINEMENT

HERE IS HOW

- Recall the residual $\mathbf{r} = \mathbf{b} - \mathbf{A}\hat{\mathbf{x}}$. The process can be repeated solving $\mathbf{r} = \mathbf{A}\hat{\mathbf{z}}$ and forming the approximate solution $\mathbf{x} + \mathbf{z}$.
- Denoting $\mathbf{x}_n$ the solution obtained in the iteration n , it's not hard to see

$$\|\mathbf{x}_{n+1} - \mathbf{x}\| \leq \|\mathbf{A}^{-1}\| \|\Delta\mathbf{A}\| \|\mathbf{x}_n - \mathbf{x}\|.$$

where $\Delta\mathbf{A}$ is the error in the n -th iteration of Gauss Elimination.

- Clearly, if $\|\mathbf{A}^{-1}\| \|\Delta\mathbf{A}\| < 1$ the procedure converges.
- It can be shown $\|\Delta\mathbf{A}\| \leq \rho n \epsilon_{mach} \|\mathbf{A}\|$ where ρ is a constant factor computed during the Gauss elimination process.

SHERMAN-MORRISON FORMULA

Sherman-Morrison formula gives inverse of matrix resulting from rank-one change to matrix whose inverse is already known:

$$(\mathbf{A} - \mathbf{u}\mathbf{v}^T)^{-1} = \mathbf{A}^{-1} + \mathbf{A}^{-1}\mathbf{u}(1 - \mathbf{v}^T\mathbf{A}^{-1}\mathbf{u})^{-1}\mathbf{v}^T\mathbf{A}^{-1},$$

where $\mathbf{u}$ and $\mathbf{v}$ are n -vectors

Evaluation of formula requires $\mathcal{O}(n^2)$ work (for matrix-vector multiplications) rather than $\mathcal{O}(n^3)$ work required for inversion

SHERMAN-MORRISON EXAMPLE

Consider rank-one modification

$$\begin{bmatrix} 2 & 4 & -2 \\ 4 & 9 & -3 \\ -2 & -1 & 7 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 2 \\ 8 \\ 10 \end{bmatrix}$$

(with 3, 2 entry changed) of system whose LU factorization was computed in earlier example

One way to choose update vectors:

$$\mathbf{u} = \begin{bmatrix} 0 \\ 0 \\ -2 \end{bmatrix} \quad \text{and} \quad \mathbf{v} = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix},$$

so matrix of modified system is $\mathbf{A} - \mathbf{u}\mathbf{v}^T$

EXAMPLE (CONTINUED)

Using LU factorization of A to solve $Az = u$
and $Ay = b$,

$$z = \begin{bmatrix} -3/2 \\ 1/2 \\ -1/2 \end{bmatrix} \quad \text{and} \quad y = \begin{bmatrix} -1 \\ 2 \\ 2 \end{bmatrix}$$

Final step computes updated solution

$$x = y + \frac{v^T y}{1 - v^T z} z =$$
$$\begin{bmatrix} -1 \\ 2 \\ 2 \end{bmatrix} + \frac{2}{1 - 1/2} \begin{bmatrix} -3/2 \\ 1/2 \\ -1/2 \end{bmatrix} = \begin{bmatrix} -7 \\ 4 \\ 0 \end{bmatrix}$$

PROOF

To solve linear system $(A - uv^T)x = b$ with new matrix, use formula to obtain

$$\begin{aligned}x &= (A - uv^T)^{-1}b \\&= A^{-1}b + A^{-1}u(1 - v^T A^{-1}u)^{-1}v^T A^{-1}b,\end{aligned}$$

which can be implemented by steps

1. Solve $Az = u$ for z , so $z = A^{-1}u$
2. Solve $Ay = b$ for y , so $y = A^{-1}b$
3. Compute $x = y + ((v^T y)/(1 - v^T z))z$

CHOLESKY FACTORIZATION

Symmetric Positive Definite Matrices

If A is symmetric and positive definite, then LU factorization can be arranged so that

$$U = L^T, \quad \text{that is,} \quad A = LL^T,$$

where L is lower triangular with positive diagonal entries

Algorithm for computing *Cholesky factorization* derived by equating corresponding entries of A and LL^T and generating them in correct order

CHOLESKY FACTORIZATION (CONTINUED)

Features of Cholesky algorithm symmetric positive definite matrices:

- All n square roots are of positive numbers, so algorithm well defined
- No pivoting required for numerical stability
- Only lower triangle of $\mathbf{A}$ accessed, and hence upper triangular portion need not be stored
- Only $n^3/6$ multiplications and similar number of additions required

SYMMETRIC MATRICES

For symmetric indefinite $\mathbf{A}$, Cholesky factorization not applicable, and some form of pivoting generally required for numerical stability

Factorization of form

$$\mathbf{PAP}^T = \mathbf{LDL}^T,$$

with $\mathbf{L}$ unit lower triangular and $\mathbf{D}$ either tridiagonal or block diagonal with 1×1 and 2×2 diagonal blocks, can be computed stably using symmetric pivoting strategy

REFERENCES

Michael T. Heath
Scientific Computing (Revised Second Edition)
SIAM