

Rank

Stephen Boyd and Sanjay Lall
Revision by Amirhossein Afsharrad

EE263
Stanford University

Rank of a matrix

we define the **rank** of $A \in \mathbb{R}^{m \times n}$ as

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

(nontrivial) facts:

- ▶ $\text{rank}(A) = \text{rank}(A^T)$
- ▶ $\text{rank}(A)$ is maximum number of independent columns (or rows) of A
hence $\text{rank}(A) \leq \min(m, n)$

$$\begin{array}{l} \downarrow \text{range}(A) \subseteq \mathbb{R}^m \\ \neq \\ \uparrow \text{range}(A^T) \subseteq \mathbb{R}^n \end{array}$$

$$A \in \mathbb{R}^{m \times n}$$

$$A = [a_1 \ \dots \ a_n]$$

$$\text{rank}(A) = r < \begin{matrix} m, n \\ \uparrow \quad \swarrow \\ 100 \quad 50 \\ r = 30 \end{matrix}$$

* remove ~~any~~ column
 a

such that the rank
will not change

keep doing that until you have a basis

Rank of a matrix

- ▶ **rank**(A) is maximum number of independent columns of A
 - ▶ to see this, notice that if the columns of A are independent, then the number of columns r is the rank, since the columns are a basis for the range
 - ▶ and if not, then there must be one column in the span of the others, so remove it, and repeat if necessary
 - ▶ all other independent sets of columns must have no more than r elements.
- ▶ proof of **rank**(A) = **rank**(A^T) uses QR, to come

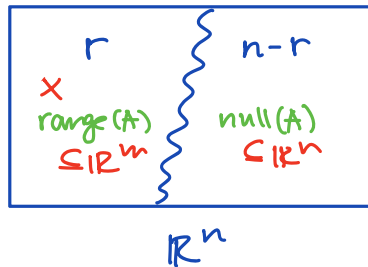
Conservation of dimension

$$A \in \mathbb{R}^{m \times n}$$

$$\underbrace{\dim \text{range}(A) + \dim \text{null}(A)}_{\text{rank}(A)} = n$$

input dimension

• WRONG intuition



- ▶ **rank**(A) is dimension of set 'hit' by the mapping $y = Ax$
- ▶ **dim null**(A) is dimension of set of x 'crushed' to zero by $y = Ax$
- ▶ 'conservation of dimension': each dimension of input is either crushed to zero or ends up in output
- ▶ roughly speaking:
 - ▶ n is number of degrees of freedom in input x
 - ▶ **dim null**(A) is number of degrees of freedom lost in the mapping from x to $y = Ax$
 - ▶ **rank**(A) is number of degrees of freedom in output y
- ▶ proof using QR

Coding interpretation of rank

$$A \in \mathbb{R}^{20 \times 10}$$

$$B \in \mathbb{R}^{10 \times 20}$$

$$BA = I$$

$$\text{rank}(BC) \leq \min\{\text{rank}(B), \text{rank}(C)\}$$

AB is never the identity

$$AB = I_{20 \times 20} \rightarrow \text{rank}(AB) = 20 < \text{rank}(A) = 10$$

contradiction

► hence if $A = BC$ with $B \in \mathbb{R}^{m \times r}$, $C \in \mathbb{R}^{r \times n}$, then $\text{rank}(A) \leq r$

► **converse**: if $\text{rank}(A) = r$ then $A \in \mathbb{R}^{m \times n}$ factors as $A = BC$ with $B \in \mathbb{R}^{m \times r}$, $C \in \mathbb{R}^{r \times n}$

► $\text{rank}(A) = r$ is minimum size vector needed to faithfully reconstruct y from x

$$A \in \mathbb{R}^{1000 \times 2000}$$

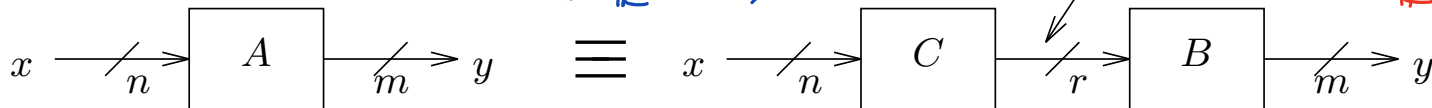
$$\text{rank}(A) = 10$$

$\rightarrow B, C$ exist

$$\text{s.t. } A = BC$$

$$B \in \mathbb{R}^{1000 \times 10}, C \in \mathbb{R}^{10 \times 2000}$$

$\text{rank}(A)$ lines



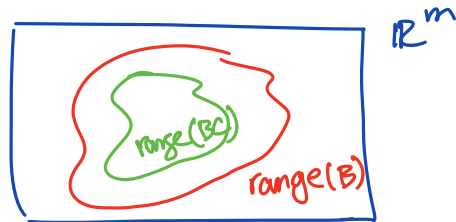
$$y = B(Cx)$$

$\mathbb{R}^{10}$
 $\mathbb{R}^{2000}$
 $\mathbb{R}^{1000}$

Coding interpretation of rank

$$y \in \text{range}(BC) \rightarrow y = B \underbrace{Cx}_z \rightarrow y = Bz \rightarrow y \in \text{range}(B)$$

- ▶ $\text{rank}(BC) \leq \text{rank}(B)$ because $\text{range}(BC) \subset \text{range}(B)$
- ▶ transpose implies $\text{rank}(BC) \leq \text{rank}(C)$
- ▶ factorization converse comes from QR



Low-rank factorization: latent factors

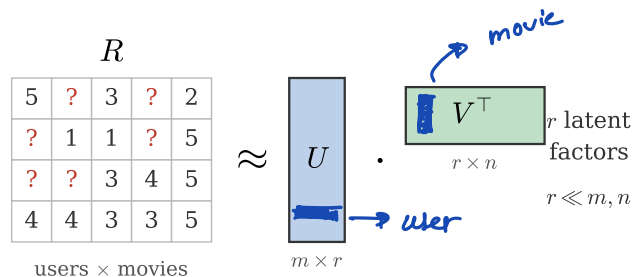


if $A = BC$ with inner dimension r , then r numbers (the *latent factors*) mediate every input-output pair

recommender systems: a ratings matrix factors as $R \approx UV^T$

- ▶ a row of U is a user's taste, a row of V a movie's genre mix
- ▶ r hidden factors explain all ratings, and fill in the missing ones

the same idea represents users, items, and words as short *embedding* vectors in $\mathbb{R}^r$; store $(m+n)r$ numbers instead of mn



Application: fast matrix-vector multiplication

$$A \in \mathbb{R}^{1000 \times 2000}, \text{rank}(A)=10 \rightarrow B, C \text{ exist}$$
$$\text{s.t. } A=BC$$
$$B \in \mathbb{R}^{1000 \times 10}, C \in \mathbb{R}^{10 \times 2000}$$

$$y = B(Cx)$$

- ▶ need to compute matrix-vector product $y = Ax$, $A \in \mathbb{R}^{m \times n}$
- ▶ A has known factorization $A = BC$, $B \in \mathbb{R}^{m \times r}$
- ▶ computing $y = Ax$ directly: mn operations $\rightarrow 2 \times 10^6$ 3000×10
- ▶ computing $y = Ax$ as $y = B(Cx)$ (compute $z = Cx$ first, then $y = Bz$): $rn + mr = (m + n)r$ operations
- ▶ savings can be considerable if $r \ll \min\{m, n\}$

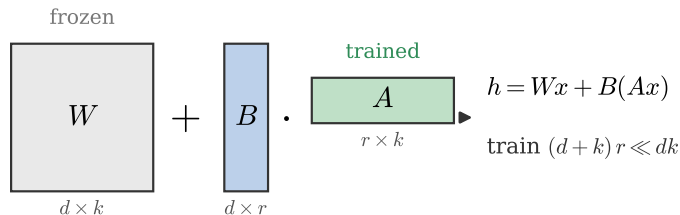
$$\frac{3 \times 10^4}{2 \times 10^6} = 1.5 \times 10^{-2}$$

Application: low-rank adaptation (LoRA)

a pretrained weight matrix $W \in \mathbb{R}^{d \times k}$ is huge, and fine-tuning all of it is expensive

LoRA freezes W and learns a *low-rank* update $\Delta W = BA$, with $B \in \mathbb{R}^{d \times r}$, $A \in \mathbb{R}^{r \times k}$, $r \ll d, k$

- ▶ forward pass $h = Wx + B(Ax)$ (the same fast-product trick)
- ▶ trains $(d + k)r$ parameters instead of dk , often well under 1%
- ▶ many small adapters can share one frozen base model



Low-rank approximation: compression

real matrices (images, data tables) are often close to a *low-rank* matrix: a rank- r approximation needs only $r(m + n)$ numbers instead of mn

rank 5



2% of the data

rank 20



7% of the data

rank 50



18% of the data

full (rank 512)



100%

rank measures *redundancy*: a few independent patterns already capture most of the data (the best rank- r approximation comes from the SVD, later)

Full rank matrices

for $A \in \mathbb{R}^{m \times n}$ we always have $\text{rank}(A) \leq \min(m, n)$

we say A is *full rank* if $\text{rank}(A) = \min(m, n)$

- ▶ for **square** matrices, full rank means nonsingular
- ▶ for **skinny** matrices ($m \geq n$), full rank means columns are independent
- ▶ for **fat** matrices ($m \leq n$), full rank means rows are independent