
Linear Transformations and Reconstruction: Choosing a Decoder Basis

ECE 57000 — September 11, 2026

David I. Inouye

| Wednesday made the SVD a way to read a linear transformation

- Orthogonal factors preserve lengths; singular values set directional scalings.
- The largest singular value bounds how much any input can stretch.
- Rank counts the nonzero singular values: how many independent directions survive.

Next question: When can we recover an input from its transformed output?

| Rank counts the **independent directions that survive**

$$\text{rank}(A) = \#\{j : \sigma_j > 0\} \leq \min(m, d)$$

For example:

$$\Sigma = \begin{bmatrix} 3 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{bmatrix} \implies \text{rank}(A) = 2$$

- Two independent input directions produce nonzero output directions.
- The third input direction disappears: $A\mathbf{v}_3 = \mathbf{0}$.
- Rotations and reflections change which directions these are, but not the rank.

| A null direction makes **different inputs indistinguishable**

The null space contains all directions that map to zero:

$$\text{null}(A) = \{\boldsymbol{h} : A\boldsymbol{h} = \mathbf{0}\}.$$

For any nonzero $\boldsymbol{h} \in \text{null}(A)$:

$$A(\boldsymbol{x} + \boldsymbol{h}) = A\boldsymbol{x} + A\boldsymbol{h}$$

(Distributivity)

$$= A\boldsymbol{x}$$

(Null direction)

| A square operator is invertible when **no direction collapses**

In 1D, $y = ax$ can be undone by $x = y/a$ exactly when $a \neq 0$.

For square $A = U\Sigma V^T \in \mathbb{R}^{d \times d}$:

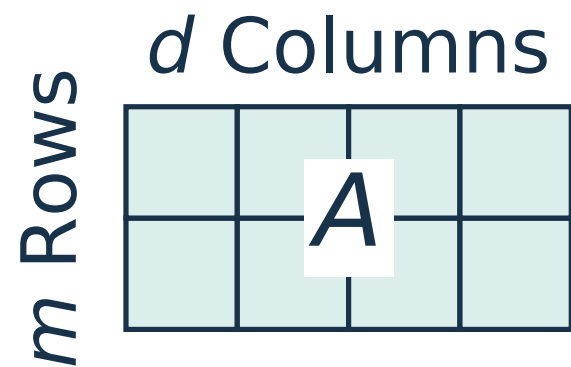
$$\text{all } \sigma_j > 0 \quad \implies \quad A^{-1} = V\Sigma^{-1}U^T, \quad \Sigma^{-1} = \text{diag}(1/\sigma_1, \dots, 1/\sigma_d).$$

$A^{-1}A = (V\Sigma^{-1}U^T)(U\Sigma V^T)$	(Substitute)
$= V\Sigma^{-1}\Sigma V^T$	$(U^T U = I_d)$
$= VV^T$	$(\Sigma^{-1}\Sigma = I_d)$
$= I_d$	$(VV^T = I_d)$

| Recovery depends on **preserving every input direction**

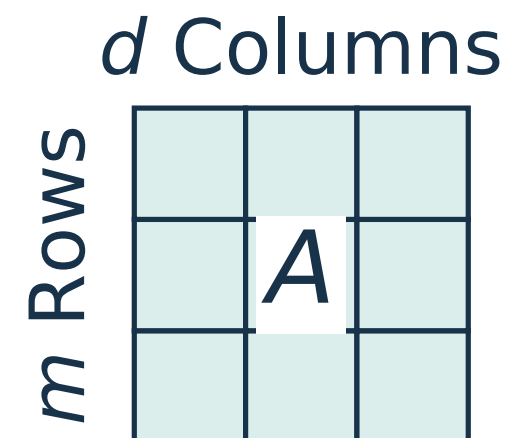
$y = Ax$ maps d input coordinates to m output coordinates.

$$m < d$$



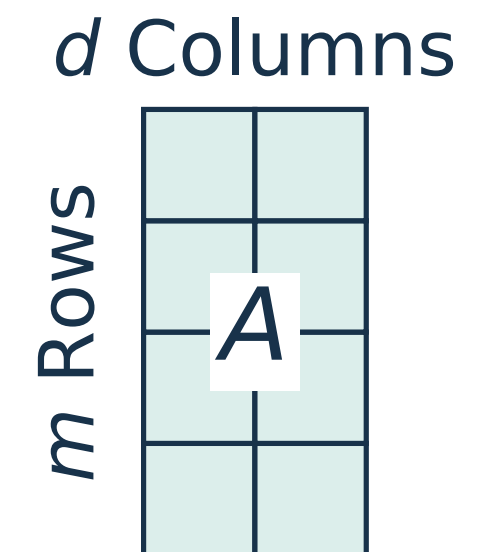
Recovery Impossible
Some Input Direction Is Lost

$$m = d$$



Recovery If $\text{rank}(A) = d$
Every Input Direction Must Survive

$$m > d$$



Recovery If $\text{rank}(A) = d$
Extra Outputs Can Be Redundant

Recovering **every** $x \in \mathbb{R}^d$ requires $\text{rank}(A) = d$.

| The pseudoinverse **undoes every nonzero scaling**

For **any** $A = U\Sigma V^T \in \mathbb{R}^{m \times d}$, define:

$$\underbrace{A^\dagger}_{d \times m} = \underbrace{V}_{d \times d} \underbrace{\Sigma^\dagger}_{d \times m} \underbrace{U^T}_{m \times m}.$$

$\Sigma^\dagger$ has the **transposed shape** of Σ , with zero off-diagonal entries:

$$(\Sigma^\dagger)_{jj} = \begin{cases} 1/\sigma_j, & \sigma_j > 0, \\ 0, & \sigma_j = 0. \end{cases}$$

- Reverse the output reorientation, undo nonzero scalings, then reorient back.
- A zero scaling destroyed information; leaving it zero cannot restore it.
- The **Moore–Penrose pseudoinverse** exists even when A has no inverse.

| The pseudoinverse recovers inputs **exactly** when no input direction is lost

For $y = Ax$, recovery of **every** input means:

$$\hat{x} = A^\dagger y = x \text{ for all } x \in \mathbb{R}^d \iff A^\dagger A = I_d.$$

Let $r = \text{rank}(A)$ and $V_r = [\mathbf{v}_1 \cdots \mathbf{v}_r]$ contain its surviving input directions.

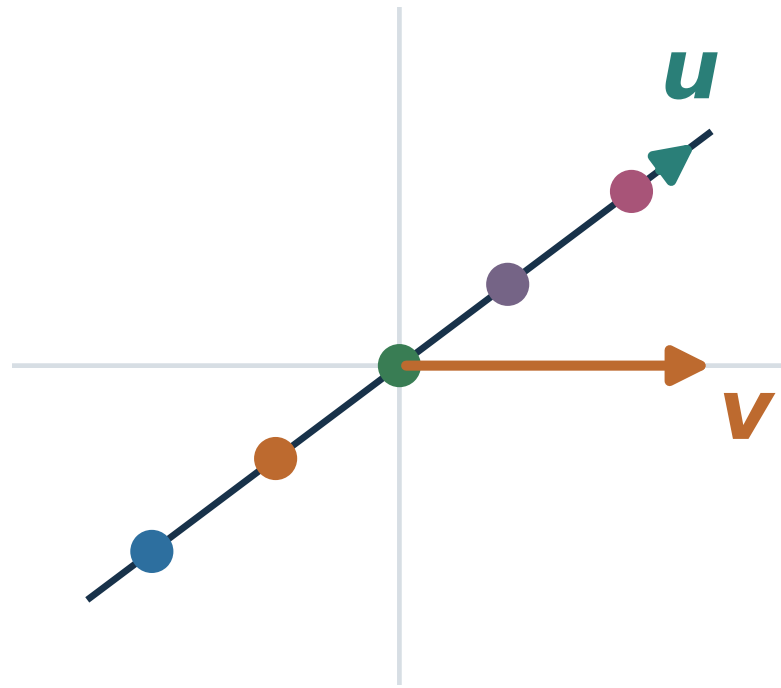
$A^\dagger A = V \Sigma^\dagger U^T U \Sigma V^T$	(Substitute)
$= V(\Sigma^\dagger \Sigma) V^T$	($U^T U = I_m$)
$= V_r V_r^T$	(Keep nonzero directions)

$\Sigma^\dagger \Sigma$ has r diagonal ones and $d - r$ zeros. Thus $A^\dagger A = I_d$ **exactly when** $r = d$: full column rank, requiring $m \geq d$.

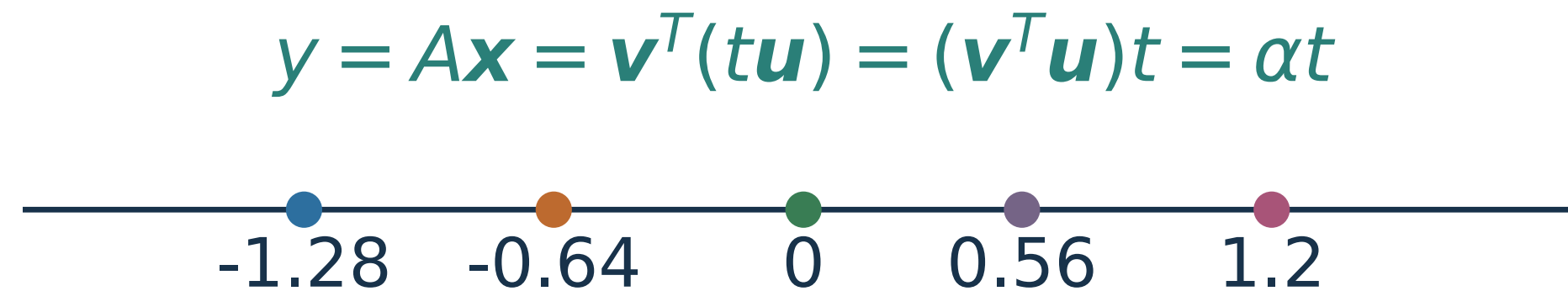
| A restricted input subspace can be **recovered exactly**

Inputs: $x = tu \in \mathbb{R}^2$, $\|u\|_2 = 1$. Encoder: $A = v^T$, with $\alpha = v^T u \neq 0$.

Inputs on a Line in 2D



One Coordinate Identifies Each Input



Decode with $B = u/\alpha$: knowing the line lets us undo its nonzero scaling.

$$\hat{x} = By = \frac{u}{\alpha}(\alpha t) = tu = x.$$

If $v = u$, then $\alpha = 1$, $y = t$, and $B = A^\dagger = u$. **In general, $B \neq A^\dagger$.**

| Read what an operator does from its singular values

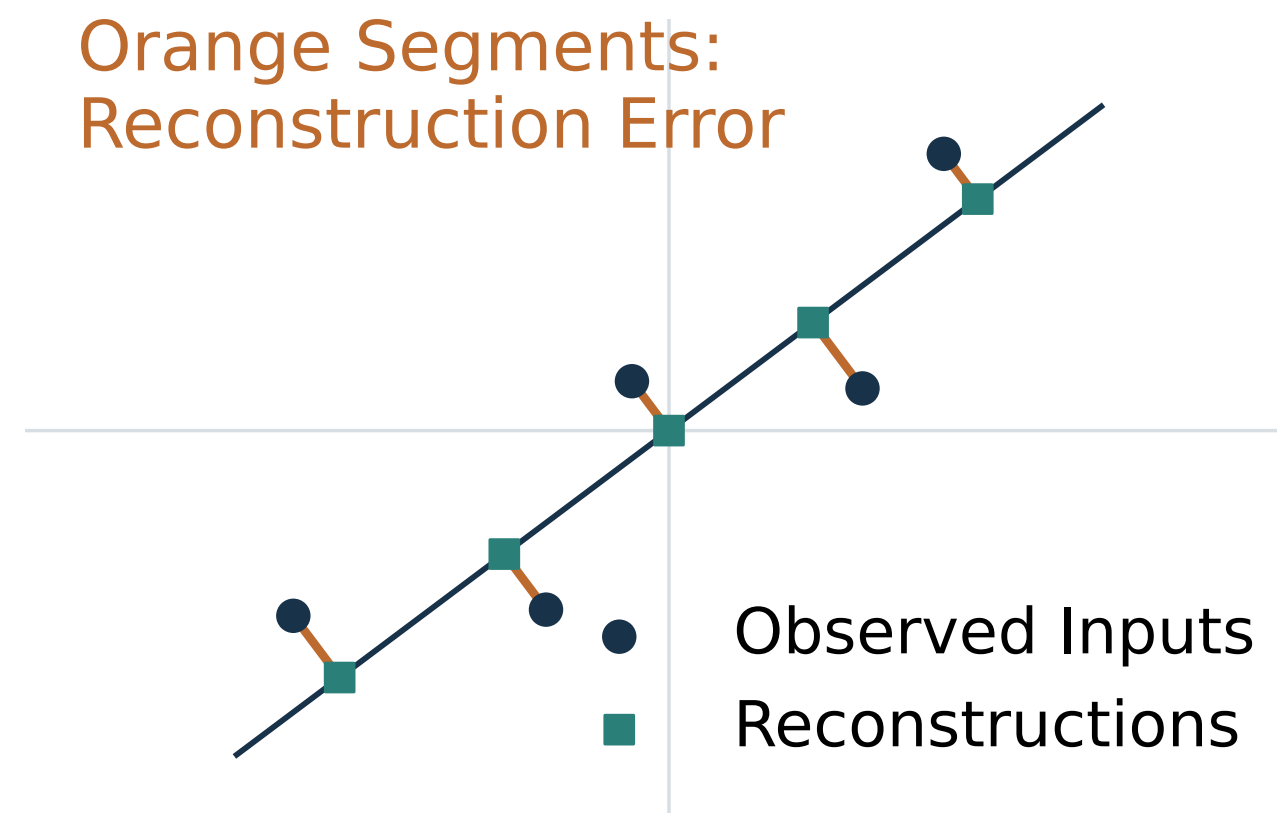
A sensor-processing system uses $A = U \operatorname{diag}(4, 1, 0) V^T$.

Can you recover every sensor input? What is the largest possible stretching factor? Could this operator preserve every pairwise distance? Explain which input direction is invisible to the system.

Rank 2; recovery is impossible in general; maximum scaling is 4. The direction v_3 disappears. All distances cannot be preserved.

| Data near a subspace can have **small reconstruction error**

What if the inputs are **close to** a line, but do not lie exactly on it?



$$y = u^T x, \quad \hat{x} = uy, \quad \text{error} = \|x - \hat{x}\|_2^2.$$

Which direction minimizes total reconstruction error?

Reconstruction and Principal Component Analysis

David I. Inouye

Dynamic Content · Last Updated: September 11, 2026

| A useful reduction should reconstruct **the observed data well**

Return to the cells described by many measured features.

- Fewer coordinates cannot preserve every possible input exactly.
- The observed data may nevertheless lie near a lower-dimensional subspace.
- We will judge a representation by how well it reconstructs those observations.

What reconstruction error are we willing to accept in exchange for fewer features?

| Four steps lead to the **PCA** formulation

We will turn “reconstruct the data well” into a problem involving one matrix.

1. Define what a good reconstruction means — Next

Choose a representation size and measure reconstruction error.

2. Choose a convenient decoder basis.

An orthonormal basis can represent the same reconstruction subspace.

3. Find the best encoder for that decoder.

The closest point in the subspace determines the coordinates.

4. Formulate the subspace choice.

Combine these results into the PCA optimization problem.

| Centering separates the **mean from variation**

For notation simplicity, we now call the **raw data** $\widetilde{X}$ and the **centered data** X .

$$\mu = \frac{1}{n} \sum_{i=1}^n \widetilde{x}_i, \quad X = \widetilde{X} - \mathbf{1}_n \mu^T.$$

- Each row is $x_i^T = (\widetilde{x}_i - \mu)^T$.
- Every feature column of X has mean zero.
- **From here on, X and x_i always denote centered data.**

Add the mean back to reconstruct raw measurements: $\widehat{\widetilde{x}}_i = \mu + \widehat{x}_i$.

| A linear encoder and decoder make **reconstruction explicit**

For one centered observation, use the same operator convention as $y = Ax$:

$$z_i = \underbrace{E}_{k \times d} x_i, \quad \hat{x}_i = \underbrace{D}_{d \times k} z_i.$$

Stacking observations as rows gives:

$$\underbrace{Z}_{n \times k} = \underbrace{X}_{n \times d} \underbrace{E^T}_{d \times k}, \quad \underbrace{\hat{X}}_{n \times d} = \underbrace{Z}_{n \times k} \underbrace{D^T}_{k \times d}.$$

The batch reconstruction is $\hat{X} = XE^T D^T$, with rank at most k .

| Reconstruction error formalizes **one meaning of preservation**

Choose a linear encoder and decoder to minimize squared reconstruction error:

$$\min_{E,D} \|X - XE^T D^T\|_F^2.$$

The **Frobenius norm** measures the error across every entry:

$$\|M\|_F^2 = \sum_{i,j} M_{ij}^2, \quad \|X - \widehat{X}\|_F^2 = \sum_{i=1}^n \|\mathbf{x}_i - \widehat{\mathbf{x}}_i\|_2^2.$$

| Four steps lead to the **PCA** formulation

We have an error measure. Now consider which reconstructions the decoder can produce.

1. Define what a good reconstruction means — Established

Choose a representation size and measure reconstruction error.

2. Choose a convenient decoder basis — Next

An orthonormal basis can represent the same reconstruction subspace.

3. Find the best encoder for that decoder.

The closest point in the subspace determines the coordinates.

4. Formulate the subspace choice.

Combine these results into the PCA optimization problem.

| Matrix–vector multiplication **combines the matrix's columns**

For a 3×2 decoder, the familiar row dot products can be regrouped by column:

$$\begin{bmatrix} d_{11} & d_{12} \\ d_{21} & d_{22} \\ d_{31} & d_{32} \end{bmatrix} \begin{bmatrix} z_1 \\ z_2 \end{bmatrix} = \underbrace{\begin{bmatrix} d_{11}z_1 + d_{12}z_2 \\ d_{21}z_1 + d_{22}z_2 \\ d_{31}z_1 + d_{32}z_2 \end{bmatrix}}_{\text{one dot product per row}} = z_1 \underbrace{\begin{bmatrix} d_{11} \\ d_{21} \\ d_{31} \end{bmatrix}}_{\mathbf{d}_1} + z_2 \underbrace{\begin{bmatrix} d_{12} \\ d_{22} \\ d_{32} \end{bmatrix}}_{\mathbf{d}_2}.$$

For any number of columns: $\hat{\mathbf{x}} = D\mathbf{z} = \sum_{j=1}^k z_j \mathbf{d}_j$.

- **Let $\mathbf{z}$ be arbitrary for now.** Every reconstruction lies in $\mathcal{S} = \text{span}\{\mathbf{d}_1, \dots, \mathbf{d}_k\}$.
- First choose how D describes that subspace; then choose the best coordinates.

| Scaling a decoder direction **only changes its coordinates**

A line can use a unit vector u or the scaled vector $3u$.

$$\underbrace{u}_{\text{unit basis}} \underbrace{t}_{\text{coordinate}} = \underbrace{(3u)}_{\text{scaled basis}} \underbrace{(t/3)}_{\text{adjusted coordinate}} .$$

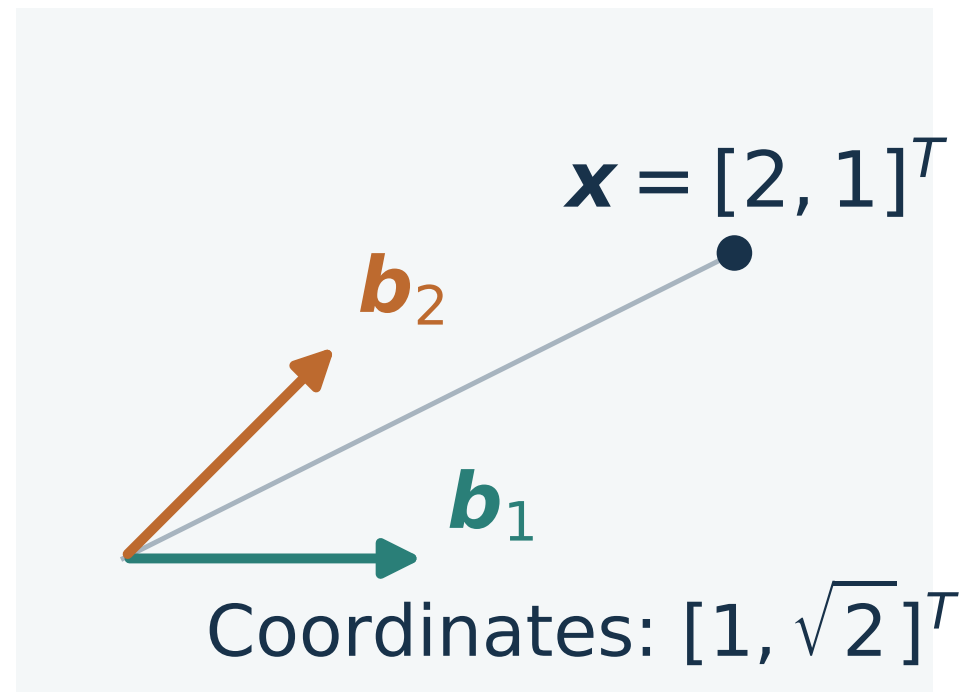
Both describe **the same point on the same line**.

- A nonzero column can be normalized if we rescale its coordinate to compensate.
- Unit columns remove an arbitrary scale choice; no reconstruction is lost.
- Next: can we choose perpendicular columns while keeping the same subspace?

| Perpendicular unit directions can describe **the same subspace**

Both unit-vector pairs span the plane: either represents **every** $x = [x_1, x_2]^T$.

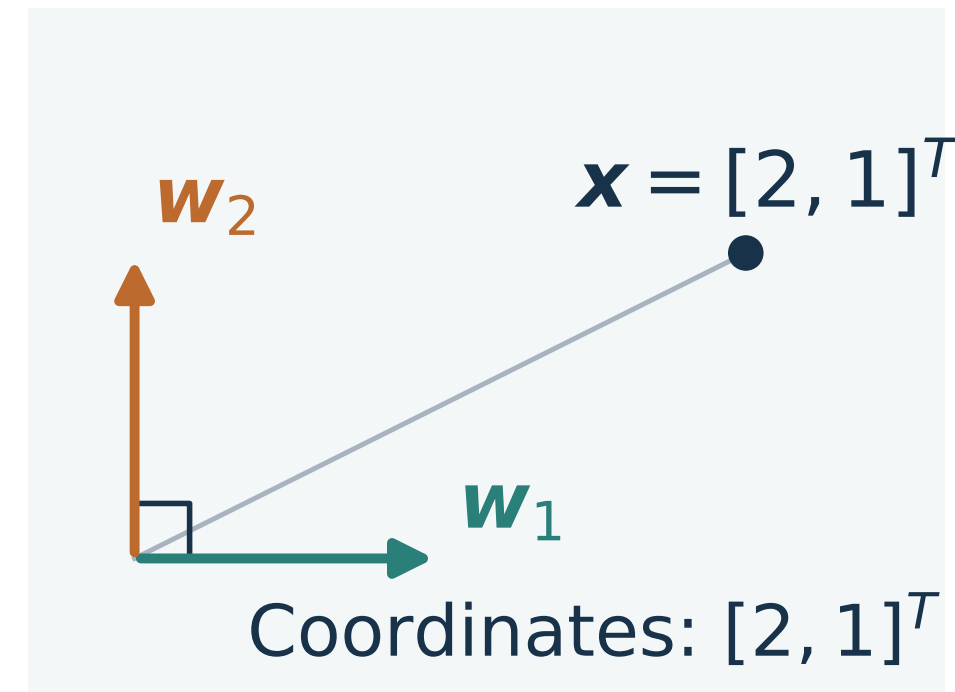
Unit Skew Directions



For $b_1 = [1, 0]^T$, $b_2 = [1, 1]^T / \sqrt{2}$:

$$x = (x_1 - x_2)b_1 + \sqrt{2}x_2b_2.$$

Orthonormal Directions



For $w_1 = [1, 0]^T$, $w_2 = [0, 1]^T$:

$$x = (w_1^T x)w_1 + (w_2^T x)w_2.$$

Same plane; simpler coordinates: an orthonormal basis gives each coefficient by a dot product.

| QR lets us choose an **orthonormal decoder without changing reconstruction**

Start with any encoder $E_0 \in \mathbb{R}^{k \times d}$ and decoder $D_0 \in \mathbb{R}^{d \times k}$, where $k \leq d$. Suppose D_0 does not have orthonormal columns.

The **thin QR decomposition** writes $D_0 = WR$: $W^T W = I_k$ and R is upper triangular. Here $W \in \mathbb{R}^{d \times k}$ and $R \in \mathbb{R}^{k \times k}$.

Can we choose $D = W$ and a new E to keep every reconstruction the same?

$$\begin{aligned} D_0 E_0 x &= (WR) E_0 x && \text{(QR of the original decoder)} \\ &= W(R E_0) x && \text{(Regroup the factors)} \\ &= D E x && \text{(Choose } D = W, E = R E_0) \end{aligned}$$

Same reconstruction for every input: an orthonormal decoder is without loss of generality. Next, fix $D = W$ and find its best encoder.