

COMS 4771 Lecture 19

1. Principal component analysis
2. Singular value decomposition

REPRESENTATION LEARNING

USEFUL REPRESENTATIONS OF DATA

Representation learning:

- ▶ **Given:** raw feature vectors $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)} \in \mathbb{R}^d$.
- ▶ **Goal:** learn a “useful” feature transformation $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^k$.
(Often $k \ll d$ —i.e., *dimensionality reduction*—but not always.)

USEFUL REPRESENTATIONS OF DATA

Representation learning:

- ▶ **Given:** raw feature vectors $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)} \in \mathbb{R}^d$.
- ▶ **Goal:** learn a “useful” feature transformation $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^k$.
(Often $k \ll d$ —i.e., *dimensionality reduction*—but not always.)

Can then use ϕ as a feature map for supervised learning.

USEFUL REPRESENTATIONS OF DATA

Representation learning:

- ▶ **Given:** raw feature vectors $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)} \in \mathbb{R}^d$.
- ▶ **Goal:** learn a “useful” feature transformation $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^k$.
(Often $k \ll d$ —i.e., *dimensionality reduction*—but not always.)

Can then use ϕ as a feature map for supervised learning.

Some previously encountered examples:

- ▶ Feature maps corresponding to pos. def. kernels (+approximations).
(Usually *data-oblivious*—feature map doesn’t depend on the data.)
- ▶ Centering $\mathbf{x} \mapsto \mathbf{x} - \boldsymbol{\mu}$
(Effect: all features have mean 0.)
- ▶ Standardization $\mathbf{x} \mapsto \text{Diag}(\sigma_1, \sigma_2, \dots, \sigma_d)^{-1}(\mathbf{x} - \boldsymbol{\mu})$.
(Effect: all features have mean 0 and unit variance.)

USEFUL REPRESENTATIONS OF DATA

Representation learning:

- ▶ **Given:** raw feature vectors $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)} \in \mathbb{R}^d$.
- ▶ **Goal:** learn a “useful” feature transformation $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^k$.
(Often $k \ll d$ —i.e., *dimensionality reduction*—but not always.)

Can then use ϕ as a feature map for supervised learning.

Some previously encountered examples:

- ▶ Feature maps corresponding to pos. def. kernels (+approximations).
(Usually *data-oblivious*—feature map doesn’t depend on the data.)
- ▶ Centering $\mathbf{x} \mapsto \mathbf{x} - \boldsymbol{\mu}$
(Effect: all features have mean 0.)
- ▶ Standardization $\mathbf{x} \mapsto \text{Diag}(\sigma_1, \sigma_2, \dots, \sigma_d)^{-1}(\mathbf{x} - \boldsymbol{\mu})$.
(Effect: all features have mean 0 and unit variance.)

What other properties of a feature representation may be desirable?

PRINCIPAL COMPONENT ANALYSIS

DIMENSIONALITY REDUCTION VIA PROJECTIONS

Projections

- ▶ **Input:** $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)} \in \mathbb{R}^d$, target dimensionality $k \in \mathbb{N}$.
- ▶ **Output:** a k -dimensional subspace, represented by an orthonormal basis $\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_k \in \mathbb{R}^d$.
- ▶ **Projection:** Formally, projection of $\mathbf{x} \in \mathbb{R}^d$ to $\text{span}(\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_k)$ is

$$\underbrace{\left(\sum_{i=1}^k \mathbf{q}_i \mathbf{q}_i^\top \right)}_{\Pi} \mathbf{x} = \sum_{i=1}^k \langle \mathbf{q}_i, \mathbf{x} \rangle \mathbf{q}_i \in \mathbb{R}^d.$$

But, we can simply represent the projection in terms of its coefficients w.r.t. the orthonormal basis $\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_k$:

$$\phi(\mathbf{x}) := \begin{bmatrix} \langle \mathbf{q}_1, \mathbf{x} \rangle \\ \langle \mathbf{q}_2, \mathbf{x} \rangle \\ \vdots \\ \langle \mathbf{q}_k, \mathbf{x} \rangle \end{bmatrix} \in \mathbb{R}^k.$$

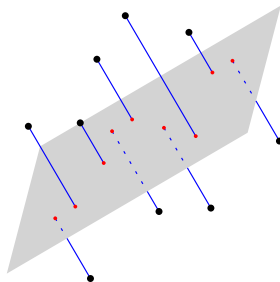
PROJECTION OF MINIMUM RESIDUAL SQUARED ERROR

Minimize residual squared error

Objective: find k -dimensional projector $\Pi: \mathbb{R}^d \rightarrow \mathbb{R}^d$ such that the average residual squared error

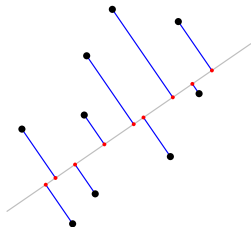
$$\frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - \Pi \mathbf{x}^{(i)} \right\|_2^2$$

is as small as possible.



PROJECTION OF MINIMUM RESIDUAL SQUARED ERROR

$k = 1$ case ($\Pi = \mathbf{q}\mathbf{q}^\top$)

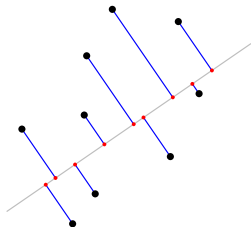


Objective: find unit vector $\mathbf{q} \in \mathbb{R}^d$ to minimize

$$\frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - \mathbf{q}\mathbf{q}^\top \mathbf{x}^{(i)} \right\|_2^2$$

PROJECTION OF MINIMUM RESIDUAL SQUARED ERROR

$k = 1$ case ($\Pi = \mathbf{q}\mathbf{q}^\top$)

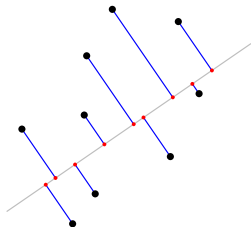


Objective: find unit vector $\mathbf{q} \in \mathbb{R}^d$ to minimize

$$\begin{aligned} & \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - \mathbf{q}\mathbf{q}^\top \mathbf{x}^{(i)} \right\|_2^2 \\ &= \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} \right\|_2^2 - \mathbf{q}^\top \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}^{(i)} \mathbf{x}^{(i)\top} \right) \mathbf{q} \end{aligned}$$

PROJECTION OF MINIMUM RESIDUAL SQUARED ERROR

$k = 1$ case ($\Pi = \mathbf{q}\mathbf{q}^\top$)



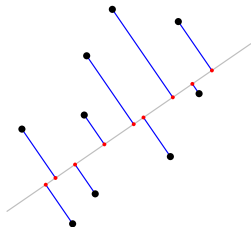
Objective: find unit vector $\mathbf{q} \in \mathbb{R}^d$ to minimize

$$\begin{aligned} & \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - \mathbf{q}\mathbf{q}^\top \mathbf{x}^{(i)} \right\|_2^2 \\ &= \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} \right\|_2^2 - \mathbf{q}^\top \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}^{(i)} \mathbf{x}^{(i)\top} \right) \mathbf{q} \\ &= \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} \right\|_2^2 - \mathbf{q}^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) \mathbf{q} \end{aligned}$$

(where $\mathbf{x}^{(i)\top}$ is i -th row of $\mathbf{X} \in \mathbb{R}^{n \times d}$).

PROJECTION OF MINIMUM RESIDUAL SQUARED ERROR

$k = 1$ case ($\Pi = \mathbf{q}\mathbf{q}^\top$)



Objective: find unit vector $\mathbf{q} \in \mathbb{R}^d$ to minimize

$$\begin{aligned} & \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - \mathbf{q}\mathbf{q}^\top \mathbf{x}^{(i)} \right\|_2^2 \\ &= \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} \right\|_2^2 - \mathbf{q}^\top \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}^{(i)} \mathbf{x}^{(i)\top} \right) \mathbf{q} \\ &= \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} \right\|_2^2 - \mathbf{q}^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) \mathbf{q} \end{aligned}$$

(where $\mathbf{x}^{(i)\top}$ is i -th row of $\mathbf{X} \in \mathbb{R}^{n \times d}$).

$$\arg \min_{\mathbf{q} \in \mathbb{R}^d: \|\mathbf{q}\|_2=1} \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - \mathbf{q}\mathbf{q}^\top \mathbf{x}^{(i)} \right\|_2^2 \equiv \arg \max_{\mathbf{q} \in \mathbb{R}^d: \|\mathbf{q}\|_2=1} \mathbf{q}^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) \mathbf{q}.$$

EIGENDECOMPOSITIONS

Every symmetric matrix $\mathbf{A} \in \mathbb{R}^{d \times d}$ guaranteed to have eigendecomposition with real eigenvalues:

$$\begin{array}{ccccccc} \boxed{\phantom{\mathbf{A}}} & = & \boxed{\phantom{\mathbf{V}}} & \boxed{\phantom{\mathbf{\Lambda}}} & \boxed{\phantom{\mathbf{V}^\top}} & = & \sum_{i=1}^d \lambda_i \mathbf{v}_i \mathbf{v}_i^\top \\ \mathbf{A} & & \mathbf{V} & \mathbf{\Lambda} & \mathbf{V}^\top & & \\ (d \times d) & & (d \times d) & (d \times d) & (d \times d) & & \end{array}$$

real **eigenvalues**: $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d$ ($\mathbf{\Lambda} = \text{Diag}(\lambda_1, \lambda_2, \dots, \lambda_d)$);

corresponding orthonormal **eigenvectors**: $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_d$ ($\mathbf{V} = [\mathbf{v}_1 | \mathbf{v}_2 | \dots | \mathbf{v}_d]$).

EIGENDECOMPOSITIONS

Every symmetric matrix $A \in \mathbb{R}^{d \times d}$ guaranteed to have eigendecomposition with real eigenvalues:

$$\begin{array}{c} \boxed{} \\ \mathbf{A} \\ (d \times d) \end{array} = \begin{array}{c} \boxed{} \\ \mathbf{V} \\ (d \times d) \end{array} \begin{array}{c} \boxed{} \\ \mathbf{\Lambda} \\ (d \times d) \end{array} \begin{array}{c} \boxed{} \\ \mathbf{V}^\top \\ (d \times d) \end{array} = \sum_{i=1}^d \lambda_i \mathbf{v}_i \mathbf{v}_i^\top$$

real **eigenvalues**: $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d$ ($\mathbf{\Lambda} = \text{Diag}(\lambda_1, \lambda_2, \dots, \lambda_d)$);
corresponding orthonormal **eigenvectors**: $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_d$ ($\mathbf{V} = [\mathbf{v}_1 | \mathbf{v}_2 | \dots | \mathbf{v}_d]$).

Fixed-point characterization of eigenvectors:

$$A\mathbf{v}_i = \lambda_i \mathbf{v}_i.$$

EIGENDECOMPOSITIONS

Variational characterization of eigenvectors:

$$\begin{aligned} \max_{\mathbf{q} \in \mathbb{R}^d} \mathbf{q}^\top \mathbf{A} \mathbf{q} \\ \text{s.t. } \|\mathbf{q}\|_2 = 1 \end{aligned}$$

- ▶ Maximum value: λ_1 (top eigenvalue)
- ▶ Maximizer: $\mathbf{v}_1$ (top eigenvector)

EIGENDECOMPOSITIONS

Variational characterization of eigenvectors:

$$\begin{aligned} \max_{\mathbf{q} \in \mathbb{R}^d} \mathbf{q}^\top \mathbf{A} \mathbf{q} \\ \text{s.t. } \|\mathbf{q}\|_2 = 1 \end{aligned}$$

- ▶ Maximum value: λ_1 (top eigenvalue)
- ▶ Maximizer: $\mathbf{v}_1$ (top eigenvector)

For $i > 1$,

$$\begin{aligned} \max_{\mathbf{q} \in \mathbb{R}^d} \mathbf{q}^\top \mathbf{A} \mathbf{q} \\ \text{s.t. } \|\mathbf{q}\|_2 = 1 \\ \langle \mathbf{q}, \mathbf{v}_j \rangle = 0 \quad \forall j < i \end{aligned}$$

- ▶ Maximum value: λ_i (i -th largest eigenvalue)
- ▶ Maximizer: $\mathbf{v}_i$ (i -th eigenvector)

PRINCIPAL COMPONENT ANALYSIS ($k = 1$)

$k = 1$ case ($\mathbf{\Pi} = \mathbf{q}\mathbf{q}^\top$)

$$\arg \min_{\mathbf{q} \in \mathbb{R}^d: \|\mathbf{q}\|_2=1} \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - \mathbf{q}\mathbf{q}^\top \mathbf{x}^{(i)} \right\|_2^2 \equiv \arg \max_{\mathbf{q} \in \mathbb{R}^d: \|\mathbf{q}\|_2=1} \mathbf{q}^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) \mathbf{q}.$$

Solution: eigenvector of $\mathbf{X}^\top \mathbf{X}$ corresponding to largest eigenvalue (i.e., the top eigenvector $\mathbf{v}_1$).

PRINCIPAL COMPONENT ANALYSIS ($k = 1$)

$k = 1$ case ($\mathbf{\Pi} = \mathbf{q}\mathbf{q}^\top$)

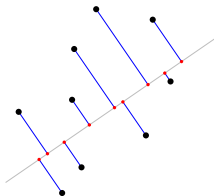
$$\arg \min_{\mathbf{q} \in \mathbb{R}^d: \|\mathbf{q}\|_2=1} \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - \mathbf{q}\mathbf{q}^\top \mathbf{x}^{(i)} \right\|_2^2 \equiv \arg \max_{\mathbf{q} \in \mathbb{R}^d: \|\mathbf{q}\|_2=1} \mathbf{q}^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) \mathbf{q}.$$

Solution: eigenvector of $\mathbf{X}^\top \mathbf{X}$ corresponding to largest eigenvalue (i.e., the top eigenvector $\mathbf{v}_1$).

$$\mathbf{q}^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) \mathbf{q} = \frac{1}{n} \sum_{i=1}^n (\mathbf{q}^\top \mathbf{x}^{(i)})^2 = \text{empirical variance of } \mathbf{q}^\top \mathbf{x}$$

(assuming $\frac{1}{n} \sum_{i=1}^n \mathbf{x}^{(i)} = \mathbf{0}$ —i.e., “centering” already applied).

top eigenvector $\equiv$ direction of maximum variance



PRINCIPAL COMPONENT ANALYSIS (GENERAL k)

General k case ($\Pi = QQ^\top$)

$$\arg \min_{\substack{Q \in \mathbb{R}^{d \times k}: \\ Q^\top Q = I}} \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - QQ^\top \mathbf{x}^{(i)} \right\|_2^2 \quad \equiv \quad \arg \max_{\substack{Q \in \mathbb{R}^{d \times k}: \\ Q^\top Q = I}} \text{tr} \left(Q^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) Q \right).$$

Solution: k eigenvectors of $\mathbf{X}^\top \mathbf{X}$ corresponding to k largest eigenvalue

PRINCIPAL COMPONENT ANALYSIS (GENERAL k)

General k case ($\Pi = QQ^\top$)

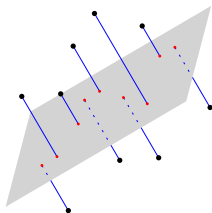
$$\arg \min_{\substack{Q \in \mathbb{R}^{d \times k}: \\ Q^\top Q = I}} \frac{1}{n} \sum_{i=1}^n \left\| \mathbf{x}^{(i)} - QQ^\top \mathbf{x}^{(i)} \right\|_2^2 \equiv \arg \max_{\substack{Q \in \mathbb{R}^{d \times k}: \\ Q^\top Q = I}} \text{tr} \left(Q^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) Q \right).$$

Solution: k eigenvectors of $\mathbf{X}^\top \mathbf{X}$ corresponding to k largest eigenvalue

$$\text{tr} \left(Q^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) Q \right) = \sum_{i=1}^k \text{empirical variance of } \mathbf{q}_i^\top \mathbf{x}$$

(assuming $\frac{1}{n} \sum_{i=1}^n \mathbf{x}^{(i)} = \mathbf{0}$ —i.e., “centering” already applied).

top k eigenvectors $\equiv k$ -dim. subspace of maximum variance



PRINCIPAL COMPONENT ANALYSIS (PCA)

Data matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$

Rank k PCA (k dimensional linear subspace)

- ▶ Get top k eigenvectors $\mathbf{V} = [\mathbf{v}_1 | \mathbf{v}_2 | \dots | \mathbf{v}_k]$ of

$$\frac{1}{n} \mathbf{X}^\top \mathbf{X} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}^{(i)} \mathbf{x}^{(i)\top}.$$

- ▶ *Feature map*: $\phi(\mathbf{x}) := (\langle \mathbf{v}_1, \mathbf{x} \rangle, \langle \mathbf{v}_2, \mathbf{x} \rangle, \dots, \langle \mathbf{v}_k, \mathbf{x} \rangle) \in \mathbb{R}^k$
- ▶ *Decorrelating property*:

$$\frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}^{(i)}) \phi(\mathbf{x}^{(i)})^\top = \text{Diag}(\lambda_1, \lambda_2, \dots, \lambda_k)$$

- ▶ *Reconstruction*: $\mathbf{x} \mapsto \mathbf{V} \phi(\mathbf{x})$

PRINCIPAL COMPONENT ANALYSIS (PCA)

Data matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$

Rank k PCA with centering (k dimensional affine subspace)

- Get top k eigenvectors $\mathbf{V} = [\mathbf{v}_1 | \mathbf{v}_2 | \dots | \mathbf{v}_k]$ of

$$\frac{1}{n} \sum_{i=1}^n (\mathbf{x}^{(i)} - \boldsymbol{\mu})(\mathbf{x}^{(i)} - \boldsymbol{\mu})^\top$$

where $\boldsymbol{\mu} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}^{(i)}$.

- *Feature map*: $\phi(\mathbf{x}) := (\langle \mathbf{v}_1, \mathbf{x} - \boldsymbol{\mu} \rangle, \langle \mathbf{v}_2, \mathbf{x} - \boldsymbol{\mu} \rangle, \dots, \langle \mathbf{v}_k, \mathbf{x} - \boldsymbol{\mu} \rangle) \in \mathbb{R}^k$
- *Decorrelating property*:

$$\frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}^{(i)}) = \mathbf{0}$$

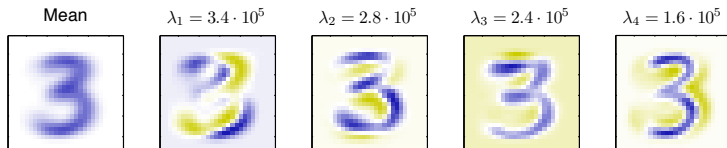
$$\frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}^{(i)}) \phi(\mathbf{x}^{(i)})^\top = \text{Diag}(\lambda_1, \lambda_2, \dots, \lambda_k)$$

- *Reconstruction*: $\mathbf{x} \mapsto \boldsymbol{\mu} + \mathbf{V} \phi(\mathbf{x})$

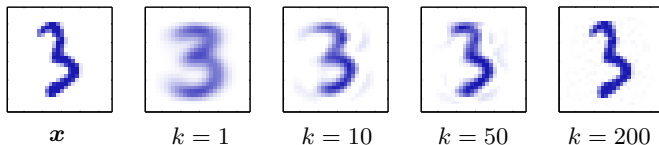
EXAMPLE: COMPRESSING DIGITS IMAGES

16×16 pixel images of handwritten 3s (as vectors in $\mathbb{R}^{256}$)

Mean μ and eigenvectors v_1, v_2, v_3, v_4



Reconstructions:



Only have to store k numbers per image,
along with the mean μ and k eigenvectors ($256(k + 1)$ numbers).

EXAMPLE: EIGENFACES

92×112 pixel images of faces (as vectors in $\mathbb{R}^{10304}$)



100 example images



top $k = 48$ eigenvectors

OTHER EXAMPLES

- ▶ $\boldsymbol{x} \in \mathbb{R}^d$: movement of stock prices for d different stocks in one day.

OTHER EXAMPLES

- ▶ $x \in \mathbb{R}^d$: movement of stock prices for d different stocks in one day.

Principal component: combination of stocks that account for the most variation in stock price movement.

OTHER EXAMPLES

- ▶ $x \in \mathbb{R}^d$: movement of stock prices for d different stocks in one day.

Principal component: combination of stocks that account for the most variation in stock price movement.

- ▶ $x \in \{1, 2, \dots, 5\}^d$: levels at which various terms describe an individual (e.g., “jolly”, “impulsive”, “outgoing”, “conceited”, “meddlesome”)

OTHER EXAMPLES

- ▶ $x \in \mathbb{R}^d$: movement of stock prices for d different stocks in one day.

Principal component: combination of stocks that account for the most variation in stock price movement.

- ▶ $x \in \{1, 2, \dots, 5\}^d$: levels at which various terms describe an individual (e.g., “jolly”, “impulsive”, “outgoing”, “conceited”, “meddlesome”)

Principal components: major personality axes in a population (e.g., “extroversion”, “agreeableness”, “conscientiousness”)

OTHER EXAMPLES

- ▶ $x \in \mathbb{R}^d$: movement of stock prices for d different stocks in one day.

Principal component: combination of stocks that account for the most variation in stock price movement.

- ▶ $x \in \{1, 2, \dots, 5\}^d$: levels at which various terms describe an individual (e.g., “jolly”, “impulsive”, “outgoing”, “conceited”, “meddlesome”)

Principal components: major personality axes in a population (e.g., “extroversion”, “agreeableness”, “conscientiousness”)

- ▶ ...

COMPUTATION

POWER METHOD

Problem: Given matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$, compute the top eigenvector of $\mathbf{X}^\top \mathbf{X}$.

- ▶ **Initialize** with random $\hat{\mathbf{v}} \in \mathbb{R}^d$.
- ▶ **Repeat:**
 1. $\hat{\mathbf{v}} := \mathbf{X}^\top \mathbf{X} \hat{\mathbf{v}}$.
 2. $\hat{\mathbf{v}} := \hat{\mathbf{v}} / \|\hat{\mathbf{v}}\|_2$.

POWER METHOD

Problem: Given matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$, compute the top eigenvector of $\mathbf{X}^\top \mathbf{X}$.

► **Initialize** with random $\hat{\mathbf{v}} \in \mathbb{R}^d$.

► **Repeat:**

1. $\hat{\mathbf{v}} := \mathbf{X}^\top \mathbf{X} \hat{\mathbf{v}}.$

2. $\hat{\mathbf{v}} := \hat{\mathbf{v}} / \|\hat{\mathbf{v}}\|_2.$

Theorem: For any $\varepsilon \in (0, 1)$, with high probability (over choice of initial $\hat{\mathbf{v}}$),

$$\hat{\mathbf{v}}^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) \hat{\mathbf{v}} \geq (1 - \varepsilon) \cdot \text{top eigenvalue of } \frac{1}{n} \mathbf{X}^\top \mathbf{X}$$

after $O\left(\frac{1}{\varepsilon} \log \frac{d}{\varepsilon}\right)$ iterations.

POWER METHOD

Problem: Given matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$, compute the top eigenvector of $\mathbf{X}^\top \mathbf{X}$.

► **Initialize** with random $\hat{\mathbf{v}} \in \mathbb{R}^d$.

► **Repeat:**

1. $\hat{\mathbf{v}} := \mathbf{X}^\top \mathbf{X} \hat{\mathbf{v}}.$

2. $\hat{\mathbf{v}} := \hat{\mathbf{v}} / \|\hat{\mathbf{v}}\|_2.$

Theorem: For any $\varepsilon \in (0, 1)$, with high probability (over choice of initial $\hat{\mathbf{v}}$),

$$\hat{\mathbf{v}}^\top \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right) \hat{\mathbf{v}} \geq (1 - \varepsilon) \cdot \text{top eigenvalue of } \frac{1}{n} \mathbf{X}^\top \mathbf{X}$$

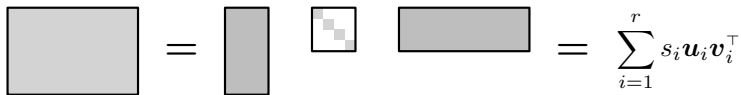
after $O\left(\frac{1}{\varepsilon} \log \frac{d}{\varepsilon}\right)$ iterations.

Similar algorithm can be used to get top k eigenvectors.

SINGULAR VALUE DECOMPOSITION

SINGULAR VALUE DECOMPOSITION

Every matrix $A \in \mathbb{R}^{n \times d}$ has a **singular value decomposition (SVD)**

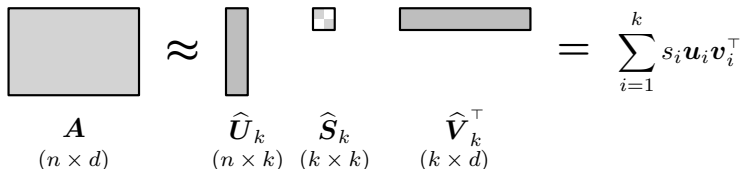

$$\begin{matrix} \boxed{} & = & \boxed{} & \boxed{} & \boxed{} & = & \sum_{i=1}^r s_i \mathbf{u}_i \mathbf{v}_i^\top \\ \mathbf{A} & & \mathbf{U} & \mathbf{S} & \mathbf{V}^\top & & \\ (n \times d) & & (n \times r) & (r \times r) & (r \times d) & & \end{matrix}$$

where

- ▶ $r = \text{rank}(\mathbf{A}) \quad (r \leq \min\{n, d\})$;
- ▶ $\mathbf{U}^\top \mathbf{U} = \mathbf{I}$ (i.e., $\mathbf{U} = [\mathbf{u}_1 | \mathbf{u}_2 | \cdots | \mathbf{u}_r]$ has orthonormal columns)
left singular vectors;
- ▶ $\mathbf{S} = \text{Diag}(s_1, s_2, \dots, s_r)$ where $s_1 \geq s_2 \geq \cdots \geq s_r > 0$
singular values;
- ▶ $\mathbf{V}^\top \mathbf{V} = \mathbf{I}$ (i.e., $\mathbf{V} = [\mathbf{v}_1 | \mathbf{v}_2 | \cdots | \mathbf{v}_r]$ has orthonormal columns)
right singular vectors.

LOW-RANK SVD

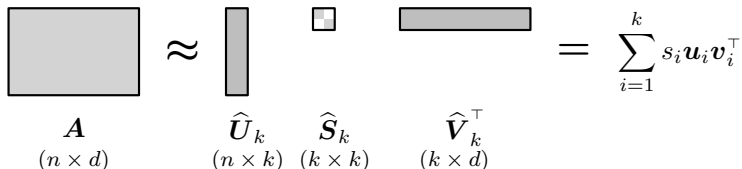
For any $k \leq \text{rank}(\mathbf{A})$, **rank- k SVD approximation**:


$$\begin{array}{ccccccc} \text{[Large Gray Box]} & \approx & \text{[Tall Gray Box]} & \text{[Small Checkerboard Box]} & \text{[Wide Gray Box]} & = & \sum_{i=1}^k s_i \mathbf{u}_i \mathbf{v}_i^\top \\ \mathbf{A} & & \hat{\mathbf{U}}_k & \hat{\mathbf{S}}_k & \hat{\mathbf{V}}_k^\top & & \\ (n \times d) & & (n \times k) & (k \times k) & (k \times d) & & \end{array}$$

(Just retain top k left/right singular vectors and singular values from SVD.)

LOW-RANK SVD

For any $k \leq \text{rank}(\mathbf{A})$, **rank- k SVD approximation**:


$$\mathbf{A} \approx \mathbf{\hat{U}}_k \mathbf{\hat{S}}_k \mathbf{\hat{V}}_k^\top = \sum_{i=1}^k s_i \mathbf{u}_i \mathbf{v}_i^\top$$

$(n \times d) \qquad (n \times k) \quad (k \times k) \quad (k \times d)$

(Just retain top k left/right singular vectors and singular values from SVD.)

Best rank- k approximation:

$$\hat{\mathbf{A}} := \mathbf{\hat{U}}_k \mathbf{\hat{S}}_k \mathbf{\hat{V}}_k^\top = \arg \min_{\substack{\mathbf{M} \in \mathbb{R}^{n \times d}: \\ \text{rank}(\mathbf{M}) \leq k}} \sum_{i=1}^n \sum_{j=1}^d (A_{i,j} - M_{i,j})^2.$$

Minimum value is simply given by

$$\sum_{i=1}^n \sum_{j=1}^d (A_{i,j} - \hat{A}_{i,j})^2 = \sum_{t>k} s_t^2.$$

EXAMPLE: LATENT SEMANTIC ANALYSIS

Represent corpus of documents by counts of words they contain:

	aardvark	abacus	abalone	...
document 1	3	0	0	...
document 2	7	0	4	...
document 3	2	4	0	...
$\vdots$	$\vdots$	$\vdots$	$\vdots$	

- ▶ One column per vocabulary word in $\mathbf{A} \in \mathbb{R}^{n \times d}$
- ▶ One row per document in $\mathbf{A} \in \mathbb{R}^{n \times d}$
- ▶ $A_{i,j}$ = numbers of times word j appears in document i .

EXAMPLE: LATENT SEMANTIC ANALYSIS

Modeling assumption:

- ▶ $k \ll \min\{n, d\}$ “topics”, each represented by a distributions over vocabulary words:

$$\beta_1, \beta_2, \dots, \beta_k \in \mathbb{R}^d$$

(Each $\beta_t = (\beta_{t,1}, \beta_{t,2}, \dots, \beta_{t,d})$ is a probability vector.)

- ▶ Each document i is associated with a topic $t_i \in \{1, 2, \dots, k\}$.

Document i 's count vector (i -th row in $\mathbf{A}$) is drawn from a multinomial distribution with probabilities given by β_{t_i} .

EXAMPLE: LATENT SEMANTIC ANALYSIS

Implication of modeling assumption

In expectation, $\mathbf{A}$ has rank k :


$$\mathbb{E}(\mathbf{A}) = \mathbf{L} \mathbf{B}^{\top}$$

$(n \times d) \qquad (n \times k) \qquad (k \times d)$

- ▶ $L_{i,t_i} = \text{length of document } i$ (other entries are zero).
- ▶ $\beta_t = t\text{-th column of } \mathbf{B}$

EXAMPLE: LATENT SEMANTIC ANALYSIS

Implication of modeling assumption

In expectation, $\mathbf{A}$ has rank k :


$$\begin{matrix} \mathbb{E}(\mathbf{A}) & = & \mathbf{L} & \mathbf{B}^{\top} \\ (n \times d) & & (n \times k) & (k \times d) \end{matrix}$$

- ▶ L_{i,t_i} = length of document i (other entries are zero).
- ▶ β_t = t -th column of $\mathbf{B}$

Observed matrix $\mathbf{A}$:

$$\mathbf{A} = \mathbb{E}(\mathbf{A}) + \mathbf{Zero\ mean\ noise}$$

so $\mathbf{A}$ is generally of rank $\min\{n, d\} \gg k$.

EXAMPLE: LATENT SEMANTIC ANALYSIS

Using SVD: rank- k SVD $\hat{U}_k \hat{S}_k \hat{V}_k^\top$ of A gives approximation to $LB^\top$:

$$\hat{A} := \hat{U}_k \hat{S}_k \hat{V}_k^\top \approx \mathbb{E}(A) = LB^\top.$$

(SVD helps remove some of the effect of the noise.)

EXAMPLE: LATENT SEMANTIC ANALYSIS

Using **SVD**: rank- k SVD $\hat{U}_k \hat{S}_k \hat{V}_k^\top$ of A gives approximation to $LB^\top$:

$$\hat{A} := \hat{U}_k \hat{S}_k \hat{V}_k^\top \approx \mathbb{E}(A) = LB^\top.$$

(SVD helps remove some of the effect of the noise.)

- Each of the n documents can be summarized by k numbers:

$$\hat{A} \hat{V}_k = \hat{U}_k \hat{S}_k \in \mathbb{R}^{n \times k}.$$

EXAMPLE: LATENT SEMANTIC ANALYSIS

Using SVD: rank- k SVD $\hat{U}_k \hat{S}_k \hat{V}_k^\top$ of A gives approximation to $LB^\top$:

$$\hat{A} := \hat{U}_k \hat{S}_k \hat{V}_k^\top \approx \mathbb{E}(A) = LB^\top.$$

(SVD helps remove some of the effect of the noise.)

- Each of the n documents can be summarized by k numbers:

$$\hat{A} \hat{V}_k = \hat{U}_k \hat{S}_k \in \mathbb{R}^{n \times k}.$$

- New document representation very useful for information retrieval.

(Example: cosine similarities between documents become faster to compute and possibly less noisy.)

EXAMPLE: LATENT SEMANTIC ANALYSIS

Using SVD: rank- k SVD $\hat{U}_k \hat{S}_k \hat{V}_k^\top$ of A gives approximation to $LB^\top$:

$$\hat{A} := \hat{U}_k \hat{S}_k \hat{V}_k^\top \approx \mathbb{E}(A) = LB^\top.$$

(SVD helps remove some of the effect of the noise.)

- ▶ Each of the n documents can be summarized by k numbers:

$$\hat{A} \hat{V}_k = \hat{U}_k \hat{S}_k \in \mathbb{R}^{n \times k}.$$

- ▶ New document representation very useful for information retrieval.

(Example: cosine similarities between documents become faster to compute and possibly less noisy.)

- ▶ Actually estimating L and B takes a bit more work.

CONNECTION TO PCA

- ▶ Let X be the design matrix. Suppose we want to compute the rank k PCA for X .

CONNECTION TO PCA

- ▶ Let $\mathbf{X}$ be the design matrix. Suppose we want to compute the rank k PCA for $\mathbf{X}$.
- ▶ Let $\mathbf{X} = \mathbf{U}\mathbf{S}\mathbf{V}^\top$ be the SVD of $\mathbf{X}$.

CONNECTION TO PCA

- ▶ Let $\mathbf{X}$ be the design matrix. Suppose we want to compute the rank k PCA for $\mathbf{X}$.
- ▶ Let $\mathbf{X} = \mathbf{U}\mathbf{S}\mathbf{V}^\top$ be the SVD of $\mathbf{X}$.
- ▶ PCA: Compute top k eigenvalues/eigenvectors of $\frac{1}{n}\mathbf{X}^\top\mathbf{X}$:

$$\frac{1}{n}\mathbf{X}^\top\mathbf{X} = \frac{1}{n}\mathbf{V}\mathbf{S}\mathbf{U}^\top\mathbf{U}\mathbf{S}\mathbf{V}^\top = \frac{1}{n}\mathbf{V}\mathbf{S}^2\mathbf{V}^\top.$$

CONNECTION TO PCA

- ▶ Let $\mathbf{X}$ be the design matrix. Suppose we want to compute the rank k PCA for $\mathbf{X}$.
- ▶ Let $\mathbf{X} = \mathbf{U}\mathbf{S}\mathbf{V}^\top$ be the SVD of $\mathbf{X}$.
- ▶ PCA: Compute top k eigenvalues/eigenvectors of $\frac{1}{n}\mathbf{X}^\top\mathbf{X}$:

$$\frac{1}{n}\mathbf{X}^\top\mathbf{X} = \frac{1}{n}\mathbf{V}\mathbf{S}\mathbf{U}^\top\mathbf{U}\mathbf{S}\mathbf{V}^\top = \frac{1}{n}\mathbf{V}\mathbf{S}^2\mathbf{V}^\top.$$

- ▶ Lo and behold - the SVD yields the eigendecomposition of $\frac{1}{n}\mathbf{X}^\top\mathbf{X}$!
- ▶ Eigenvalues are $\frac{1}{n}$ times the diagonal entries in $\mathbf{S}^2$, and eigenvectors are columns of $\mathbf{V}$.

CONNECTION TO PCA

- ▶ Let $\mathbf{X}$ be the design matrix. Suppose we want to compute the rank k PCA for $\mathbf{X}$.
- ▶ Let $\mathbf{X} = \mathbf{U}\mathbf{S}\mathbf{V}^\top$ be the SVD of $\mathbf{X}$.
- ▶ PCA: Compute top k eigenvalues/eigenvectors of $\frac{1}{n}\mathbf{X}^\top\mathbf{X}$:

$$\frac{1}{n}\mathbf{X}^\top\mathbf{X} = \frac{1}{n}\mathbf{V}\mathbf{S}\mathbf{U}^\top\mathbf{U}\mathbf{S}\mathbf{V}^\top = \frac{1}{n}\mathbf{V}\mathbf{S}^2\mathbf{V}^\top.$$

- ▶ Lo and behold - the SVD yields the eigendecomposition of $\frac{1}{n}\mathbf{X}^\top\mathbf{X}$!
- ▶ Eigenvalues are $\frac{1}{n}$ times the diagonal entries in $\mathbf{S}^2$, and eigenvectors are columns of $\mathbf{V}$.
- ▶ Even better: PCA embedding of a point $\phi(\mathbf{x}) = (\langle \mathbf{v}_1, \mathbf{x} \rangle, \langle \mathbf{v}_2, \mathbf{x} \rangle, \dots, \langle \mathbf{v}_k, \mathbf{x} \rangle)^\top$, is simply $\mathbf{V}_k^\top \mathbf{x}$.
- ▶ So PCA embedding of *all* training points are the columns of $\mathbf{V}_k^\top \mathbf{X}^\top$.

CONNECTION TO PCA

- ▶ Let $\mathbf{X}$ be the design matrix. Suppose we want to compute the rank k PCA for $\mathbf{X}$.
- ▶ Let $\mathbf{X} = \mathbf{U}\mathbf{S}\mathbf{V}^\top$ be the SVD of $\mathbf{X}$.
- ▶ PCA: Compute top k eigenvalues/eigenvectors of $\frac{1}{n}\mathbf{X}^\top\mathbf{X}$:

$$\frac{1}{n}\mathbf{X}^\top\mathbf{X} = \frac{1}{n}\mathbf{V}\mathbf{S}\mathbf{U}^\top\mathbf{U}\mathbf{S}\mathbf{V}^\top = \frac{1}{n}\mathbf{V}\mathbf{S}^2\mathbf{V}^\top.$$

- ▶ Lo and behold - the SVD yields the eigendecomposition of $\frac{1}{n}\mathbf{X}^\top\mathbf{X}$!
- ▶ Eigenvalues are $\frac{1}{n}$ times the diagonal entries in $\mathbf{S}^2$, and eigenvectors are columns of $\mathbf{V}$.
- ▶ Even better: PCA embedding of a point $\phi(\mathbf{x}) = (\langle \mathbf{v}_1, \mathbf{x} \rangle, \langle \mathbf{v}_2, \mathbf{x} \rangle, \dots, \langle \mathbf{v}_k, \mathbf{x} \rangle)^\top$, is simply $\mathbf{V}_k^\top \mathbf{x}$.
- ▶ So PCA embedding of *all* training points are the columns of $\mathbf{V}_k^\top \mathbf{X}^\top$.
- ▶ But $\mathbf{V}_k^\top = \mathbf{S}_k^{-1} \mathbf{U}_k^\top \mathbf{X}$, so we can also write the PCA embedding as $\mathbf{S}_k^{-1} \mathbf{U}_k^\top \mathbf{X} \mathbf{X}^\top$. Useful for kernelization (hw!)

- ▶ **PCA**: directions of maximum variance in data $\equiv$ subspace that minimizes residual squared error.
- ▶ **Computation**: power method
- ▶ **SVD**: general decomposition for arbitrary rectangular matrices
Low-rank SVD: best low-rank approximation of a matrix
- ▶ **PCA/SVD**: often useful when low-rank structure is expected (e.g., probabilistic modeling).