

# Dimensionality Reduction

## PCA

Aarti Singh

Machine Learning 10-315  
Nov 10, 2021

Slides Courtesy: Tom Mitchell, Eric Xing, Lawrence Saul



**MACHINE LEARNING** DEPARTMENT



# High-Dimensional data

- High-Dimensions = Lot of Features

## Document classification

Features per document =  
thousands of words/unigrams  
millions of bigrams, contextual  
information



## Surveys - Netflix

480189 users x 17770 movies

|        | movie 1 | movie 2 | movie 3 | movie 4 | movie 5 | movie 6 |
|--------|---------|---------|---------|---------|---------|---------|
| Tom    | 5       | ?       | ?       | 1       | 3       | ?       |
| George | ?       | ?       | 3       | 1       | 2       | 5       |
| Susan  | 4       | 3       | 1       | ?       | 5       | 1       |
| Beth   | 4       | 3       | ?       | 2       | 4       | 2       |

# High-Dimensional data

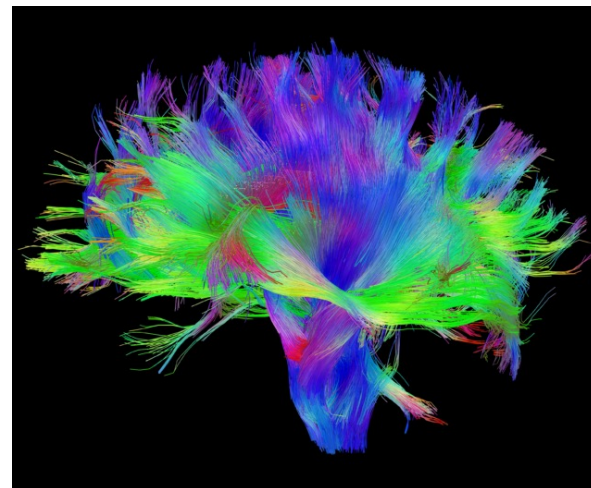
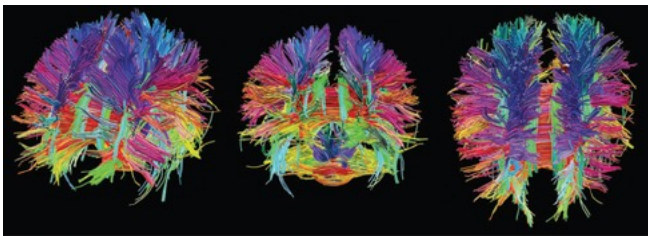
- High-Dimensions = Lot of Features

High resolution images

millions of pixels

Diffusion scans of Brain

300,000 brain fibers

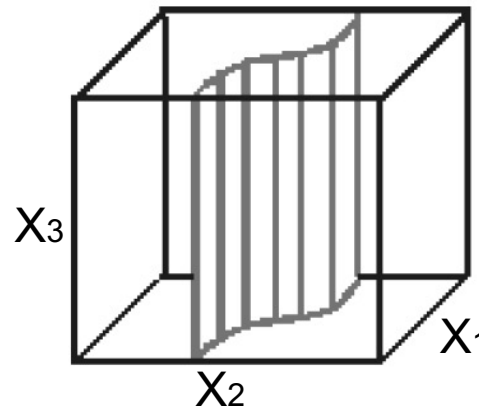


# Curse of Dimensionality

- Why are more features bad?
  - Redundant features (not all words are useful to classify a document)  
more noise added than signal
  - Hard to interpret and visualize
  - Hard to store and process data (computationally challenging)
  - Complexity of decision rule tends to grow with # features. Hard to learn complex rules as it needs more data (statistically challenging)

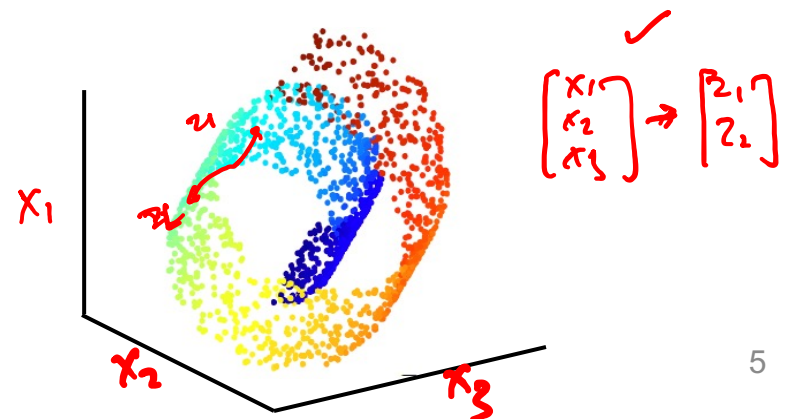
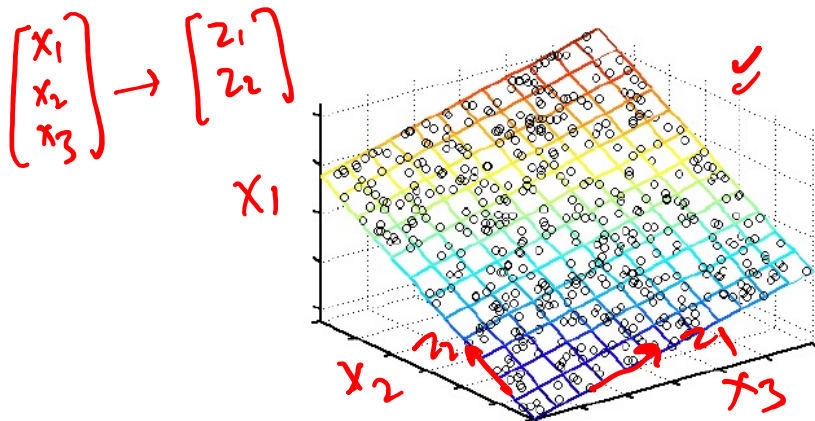
# Dimensionality Reduction

- **Feature Selection** – Only a few features are relevant to the learning task



$X_3$  - Irrelevant

- **Latent features** – Some linear/nonlinear combination of features provides a more efficient representation than observed features



# Feature Selection

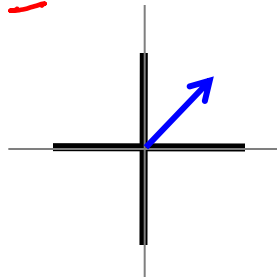
lasso

- One Approach: **Regularization (MAP)**

Integrate feature selection into learning objective by penalizing number of features with non-zero weights

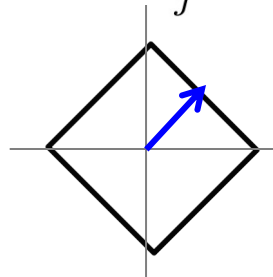
$$\hat{W} = \arg \min_W \underbrace{\sum_{i=1}^n -\log P(Y_i|X_i; W)}_{\text{-ve log likelihood}} + \underbrace{\lambda \|W\|}_{\text{penalty}}$$

$$\|W\|_0 = \#\{W_j > 0\}$$



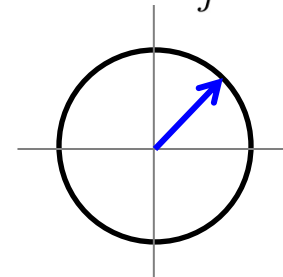
Minimizes # features  
chosen

$$\|W\|_1 = \sum_j |W_j|$$



Convex  
compromise

$$\|W\|_2 = \sum_j W_j^2$$



Small weights of  
features chosen

# Latent Features

Combinations of observed features provide more efficient representation, and capture underlying relations that govern the data

E.g. Ego, personality and intelligence are hidden attributes that characterize human behavior instead of survey questions

Topics (sports, science, news, etc.) instead of documents ✓

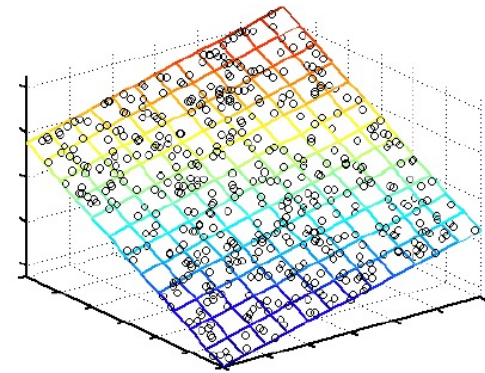
Often may not have physical meaning

- Linear

**Principal Component Analysis (PCA)** ✓

Factor Analysis ✓

Independent Component Analysis (ICA) ✓

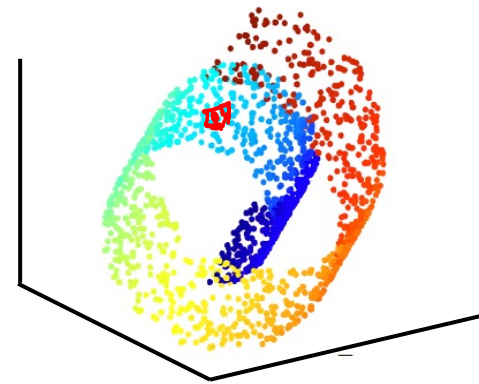


- Nonlinear

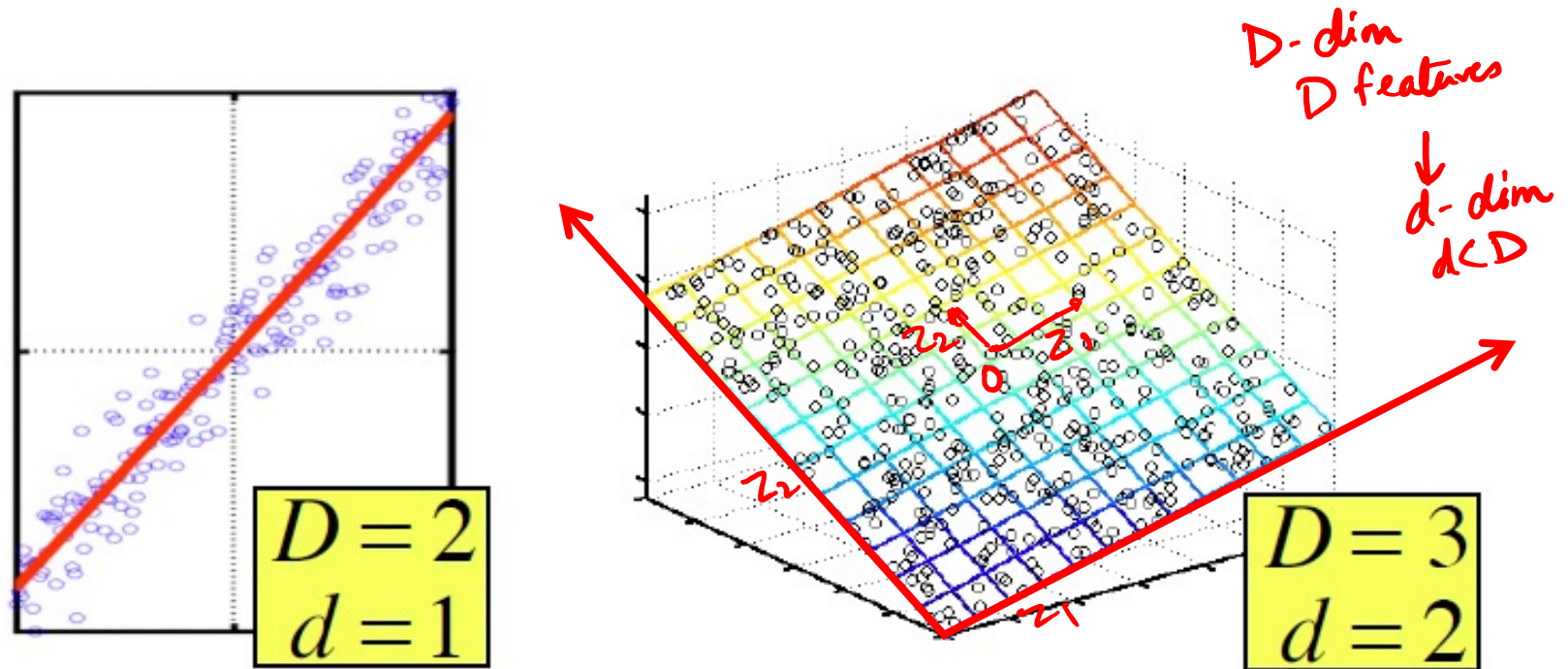
**Kernel PCA** ✓

Laplacian Eigenmaps, ISOMAP, LLE ←

Autoencoders ✓



# Principal Component Analysis (PCA)



When data lies on or near a low  $d$ -dimensional linear subspace, axes of this subspace are an effective representation of the data

Identifying the axes is known as Principal Components Analysis, and can be obtained by Eigen or Singular value decomposition

# Data for PCA

Data  $X$  =  $[x_1, x_2, \dots, x_n]$  where each data point  $x_i$  is D-dimensional vector

$X$  is  $D \times n$  matrix

Assume data are centered i.e. sample mean  $\frac{1}{n} \sum_{i=1}^n x_i = 0$  ✓

What if data is not centered?

Subtract off sample mean from each data point

$$\tilde{x}_i = x_i - \frac{1}{n} \sum_{i=1}^n x_i$$

$$\frac{1}{n} \sum_{i=1}^n \tilde{x}_i = 0$$

Since data matrix is centered, sample covariance matrix can be written as

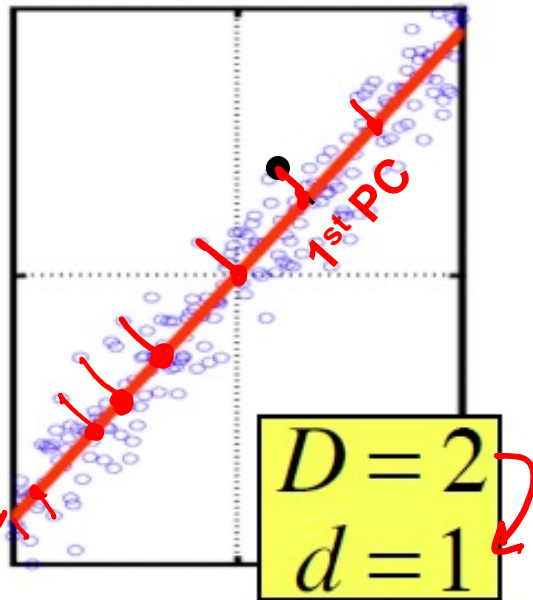
$$\text{Var}(X) = E[(X - EX)^T]$$

$$S = \frac{1}{n} X X^T$$

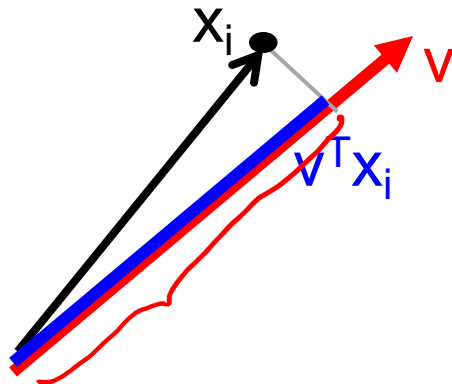
$D \times D$

$$\frac{\tilde{X} \tilde{X}^T}{n}$$

# Principal Component Analysis (PCA)



$\vec{v}$   
 $\|\vec{v}\|=1$



Principal Components (PC) are orthogonal directions that capture most of the variance in the data

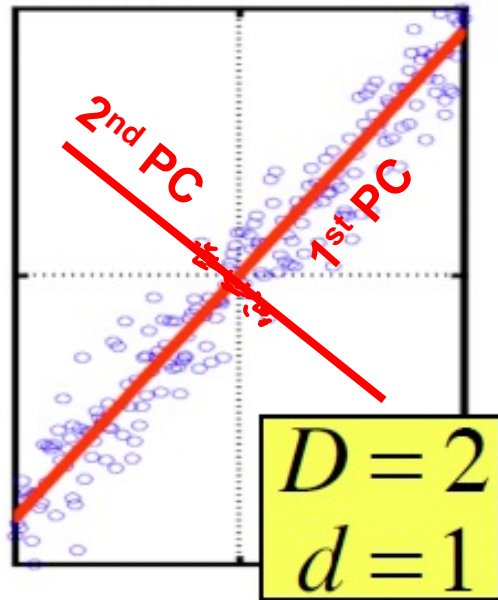
1<sup>st</sup> PC – direction of greatest variability in data ✓

Projection of data points along 1<sup>st</sup> PC discriminate the data most along any one direction

Take a data point  $x_i$  (D-dimensional vector)

Projection of  $x_i$  onto the 1<sup>st</sup> PC  $v$  is  $v^T x_i$

# Principal Component Analysis (PCA)



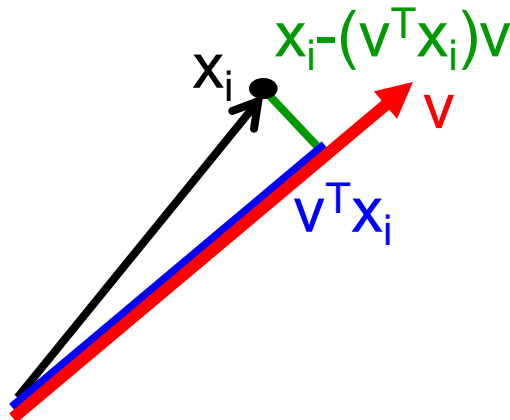
Principal Components (PC) are orthogonal unit norm directions that capture most of the variance in the data

1<sup>st</sup> PC – direction of greatest variability in data  $v_1$   $\|v_1\|=1$

2<sup>nd</sup> PC – Next orthogonal (uncorrelated) direction of greatest variability  $v_2$   $v_2^T v_1 = 0$   $\|v_2\|=1$

(remove all variability in first direction, then find next direction of greatest variability)

And so on ...



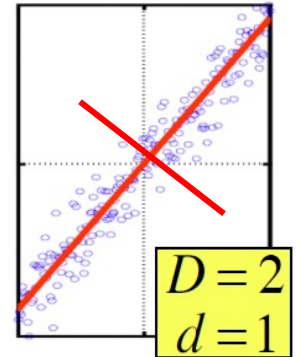
$$D \rightarrow d < D$$

# Principal Component Analysis (PCA)

Let  $\underline{v}_1, \underline{v}_2, \dots, \underline{v}_d$  denote the principal components

Orthogonal and unit norm  $\underline{v}_i^T \underline{v}_j = 0 \quad i \neq j$   
 $\underline{v}_i^T \underline{v}_i = 1 \quad \equiv \quad \|\underline{v}_i\| = 1$

Find vector that maximizes sample variance of projection

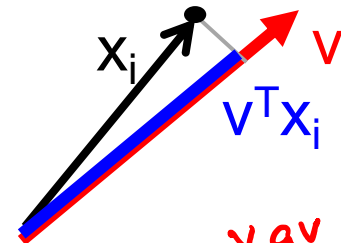


$$\frac{1}{n} \sum_{i=1}^n (\underline{v}^T \underline{x}_i)^2 = \frac{\underline{v}^T \underline{X} \underline{X}^T \underline{v}}{n}$$

*Handwritten notes:  $1 \times D \quad D \times D \quad D \times 1$*

$$\left[ \max_{\underline{v}} \underline{v}^T \underline{X} \underline{X}^T \underline{v} \quad \text{s.t.} \quad \underline{v}^T \underline{v} = 1 \right]$$

*Handwritten notes:  $\underline{v}^T \underline{v} = 1$  is underlined in red.*



$$\underline{v}^T \underline{v} = \underline{v}^T \underline{v}^2$$

$$\underline{a}^T \underline{x} \underline{x}^T \underline{a} \geq 0$$



Poll:

➤ Is this a convex optimization problem?

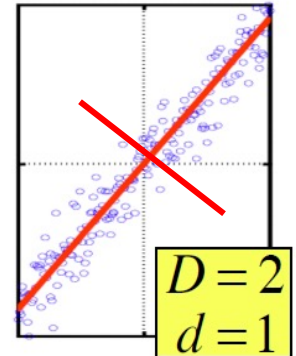
# Principal Component Analysis (PCA)

Let  $v_1, v_2, \dots, v_d$  denote the principal components

Orthogonal and unit norm  $v_i^T v_j = 0 \quad i \neq j$

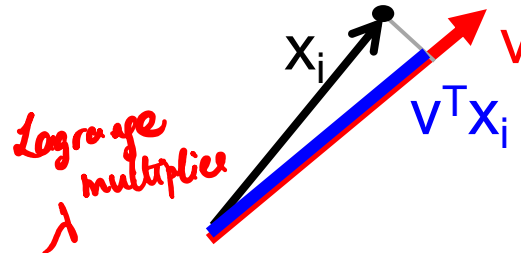
$$v_i^T v_i = 1$$

Find vector that maximizes sample variance of projection



$$\frac{1}{n} \sum_{i=1}^n (v^T x_i)^2 = \frac{v^T X X^T v}{n} \quad \leftarrow \frac{\lambda}{n}$$

$$\max_v v^T X X^T v \quad \text{s.t.} \quad v^T v = 1$$



Lagrangian:  $\max_v v^T X X^T v - \lambda(v^T v - 1)$  Wrap constraints into the objective function

$$\frac{\partial}{\partial v} = 0 \quad 2 \times \frac{\partial}{\partial v} v^T X X^T v - 2 \lambda v = 0$$

$$(X X^T - \lambda I) v = 0$$

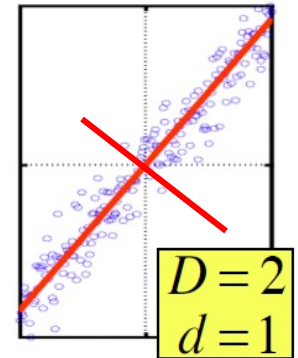
$$\Rightarrow \boxed{(X X^T) v = \lambda v} \quad \checkmark$$

# Principal Component Analysis (PCA)

1) center data

2)  $(\underline{XX^T})\underline{v} = \lambda \underline{v}$

Therefore,  $\underline{v}$  is the eigenvector of sample covariance matrix  $\underline{XX^T}$



Sample variance of projection  $= \underline{v}^T \underline{XX^T} \underline{v} = \lambda \underline{v}^T \underline{v} = \lambda$

Thus, the eigenvalue  $\lambda$  denotes the amount of variability captured along that dimension (aka amount of energy along that dimension).

Eigenvalues  $\underline{\lambda_1} > \underline{\lambda_2} > \underline{\lambda_3} > \dots$

The 1<sup>st</sup> Principal component  $\underline{v_1}$  is the eigenvector of the sample covariance matrix  $\underline{XX^T}$  associated with the largest eigenvalue  $\lambda_1$

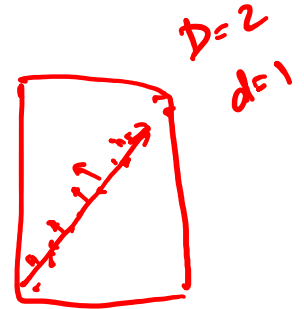
The 2<sup>nd</sup> Principal component  $\underline{v_2}$  is the eigenvector of the sample covariance matrix  $\underline{XX^T}$  associated with the second largest eigenvalue  $\lambda_2$

And so on ...

# Another interpretation

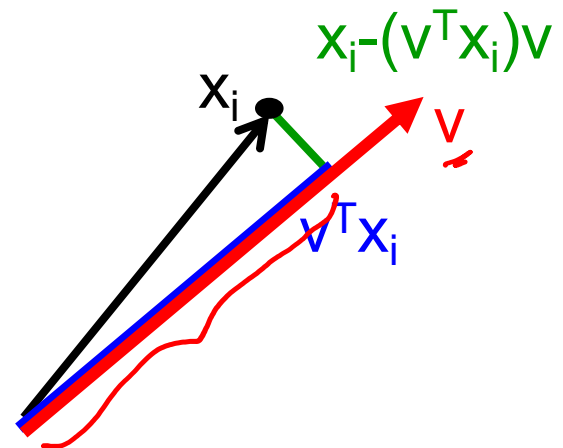
**Maximum Variance Subspace:** PCA finds vectors  $v$  such that projections on to the vectors capture maximum variance in the data

$$\frac{1}{n} \sum_{i=1}^n (v^T x_i)^2 = v^T X X^T v \quad \checkmark$$



**Minimum Reconstruction Error:** PCA finds vectors  $v$  such that projection on to the vectors yields minimum MSE reconstruction

$$\frac{1}{n} \sum_{i=1}^n \| \underbrace{x_i}_{\hat{x}_i} - \underbrace{(v^T x_i)v}_{\hat{x}_i} \|^2 \quad \checkmark$$



$$\|z\|^2 = z^T z$$

$$\frac{1}{n} \sum_{i=1}^n \|x_i - (v^T x_i) v\|^2$$

$$= \frac{1}{n} \sum_{i=1}^n (x_i - (v^T x_i) v)^T (x_i - (v^T x_i) v)$$

$$= \frac{1}{n} \sum_{i=1}^n \left( x_i^T x_i + \overbrace{v^T (v^T x_i)^2 v} - 2 (v^T x_i) v^T x_i \right)$$

$$= \frac{1}{n} \sum_{i=1}^n \left( x_i^T x_i - (v^T x_i)^2 \right)$$

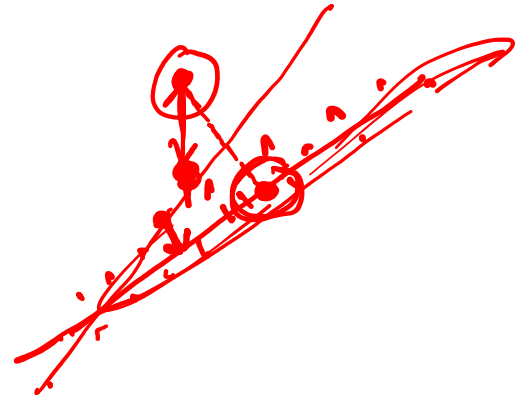
$$\arg \min_v \frac{1}{n} \sum_{i=1}^n (v^T x_i)^2$$

$$= \arg \max_v \frac{1}{n} \sum_{i=1}^n (v^T x_i)^2$$

$$= v^T X X^T v$$



ICA



$$X X^T \rightarrow V_1 \dots V_D$$

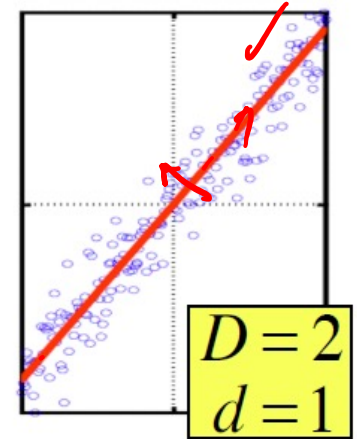
$D \times D$

# Dimensionality Reduction using PCA

The eigenvalue  $\lambda$  denotes the amount of variability captured along that dimension.

Zero eigenvalues indicate no variability along those directions => data lies exactly on a linear subspace

Only keep data projections onto principal components with non-zero eigenvalues, say  $v_1, \dots, v_d$  where  $d = \text{rank}(X X^T)$



**Original Representation**  
data point

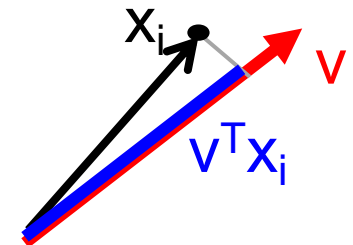
$$x_i = [x_i^1, x_i^2, \dots, x_i^D]^T$$

(D-dimensional vector)

**Transformed representation**  
projections

$$[v_1^T x_i, v_2^T x_i, \dots, v_d^T x_i]$$

(d-dimensional vector)

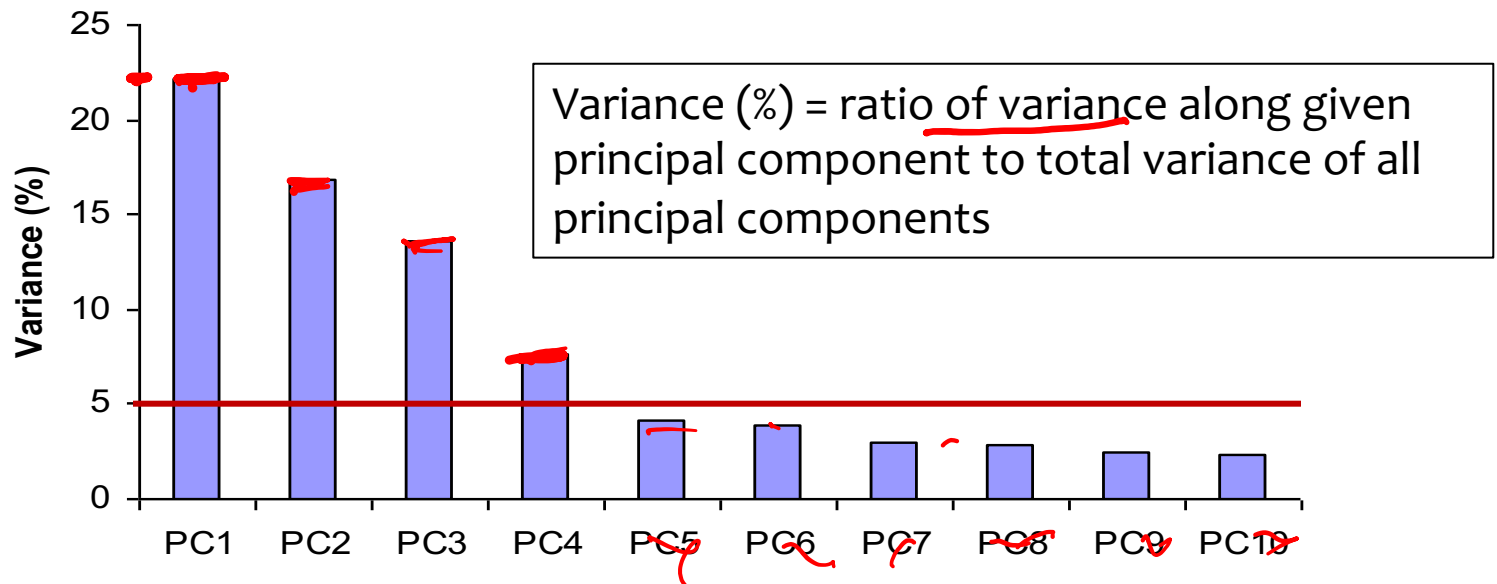


# Dimensionality Reduction using PCA

In high-dimensional problem, data usually lies near a linear subspace, as noise introduces small variability

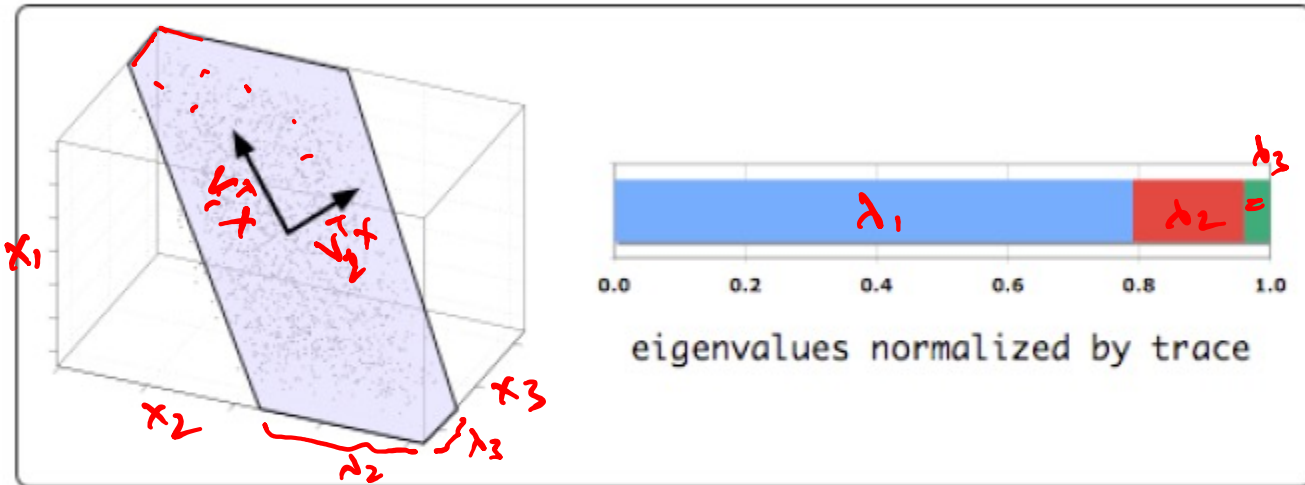
Only keep data projections onto principal components with **large** eigenvalues

Can *ignore* the components of lesser significance.



You might **lose some information**, but if the eigenvalues are small, you don't lose much

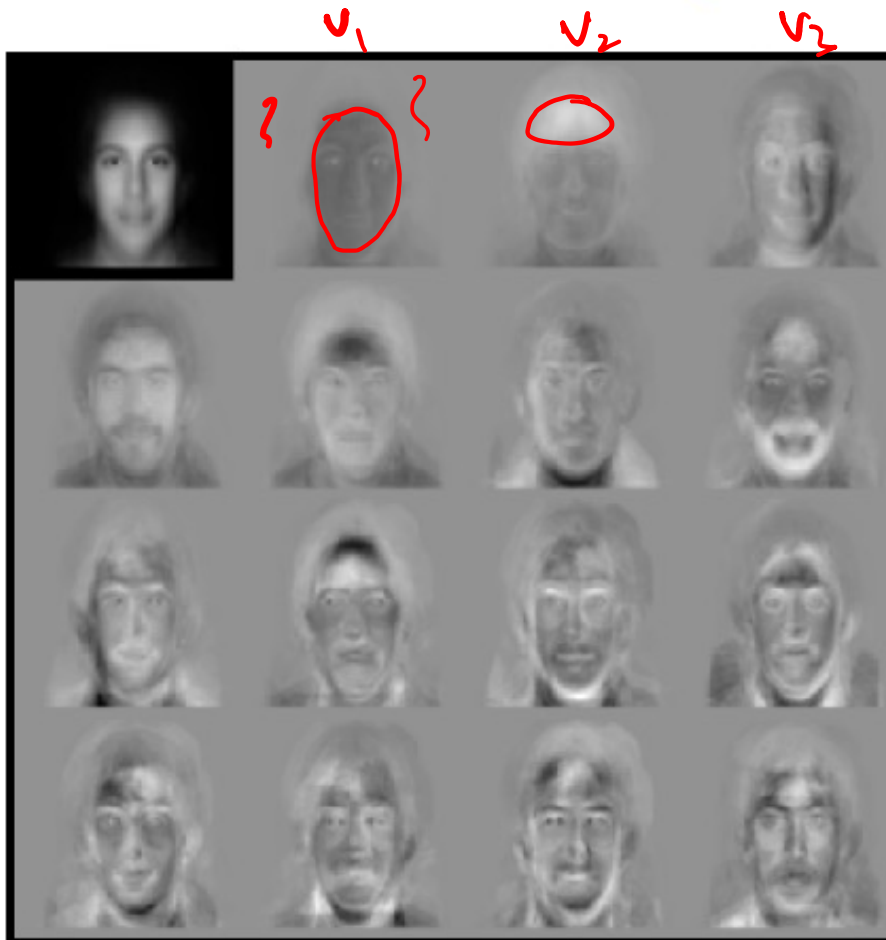
# Example of PCA



**Eigenvectors and eigenvalues of covariance matrix for  $n=1600$  inputs in  $d=3$  dimensions.**

# Example: faces

$x_i = \begin{bmatrix} \vdots \end{bmatrix}$



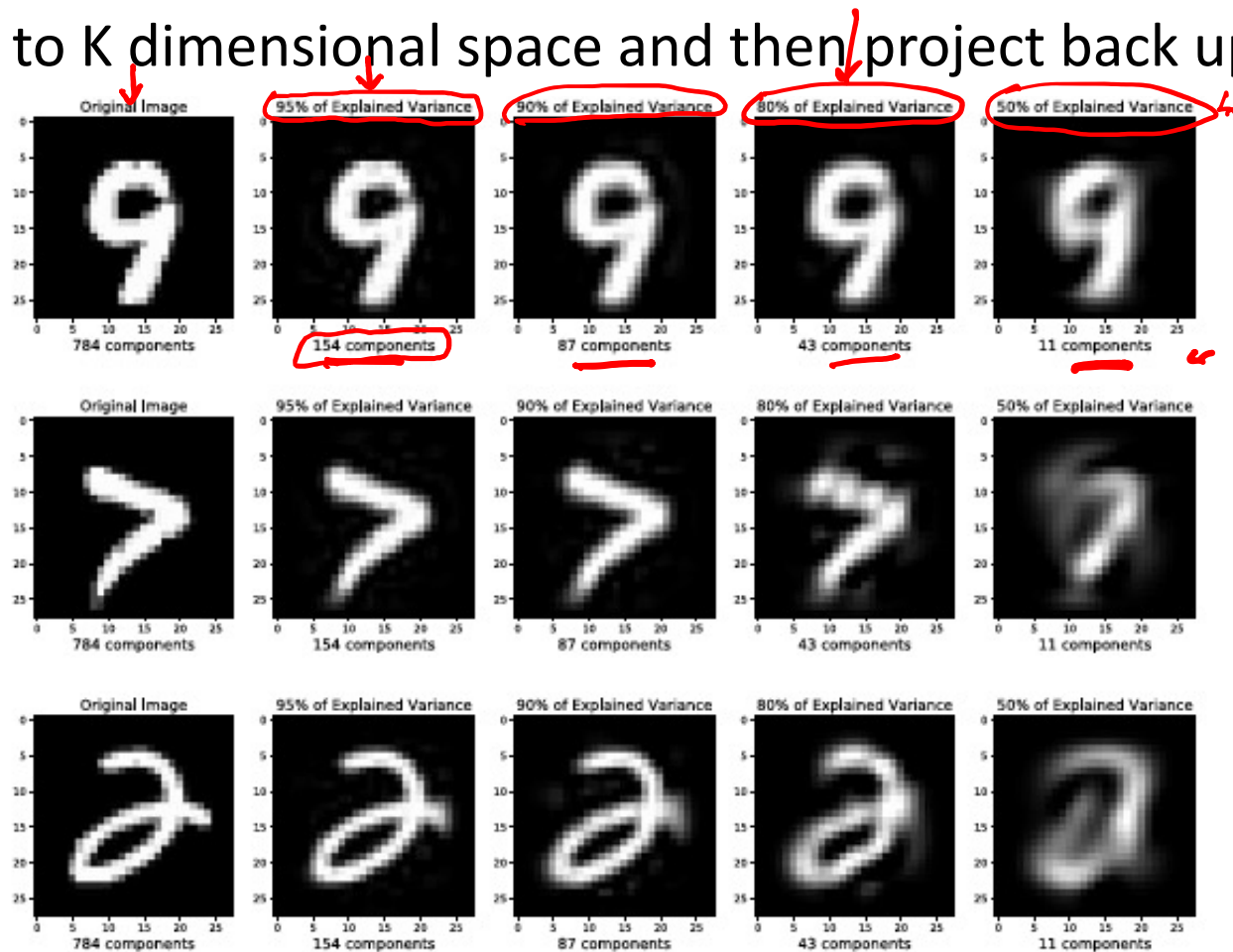
**Eigenfaces**  
from 7562  
images:

top left image  
is linear  
combination  
of rest.

Sirovich & Kirby (1987)  
Turk & Pentland (1991)

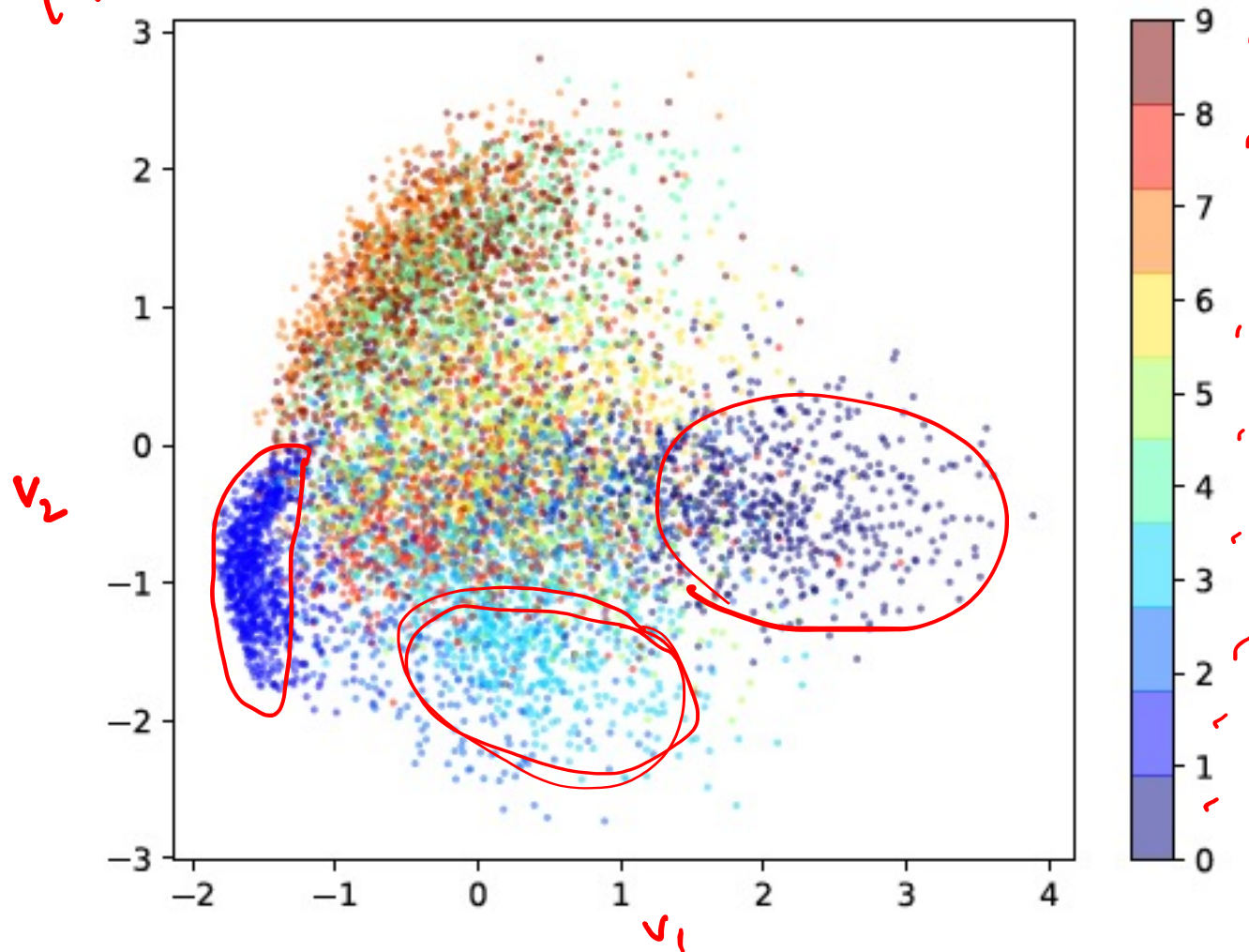
# Example: MNIST digits

- 28x28 images = 784 PCA vectors
- Project to K dimensional space and then project back up

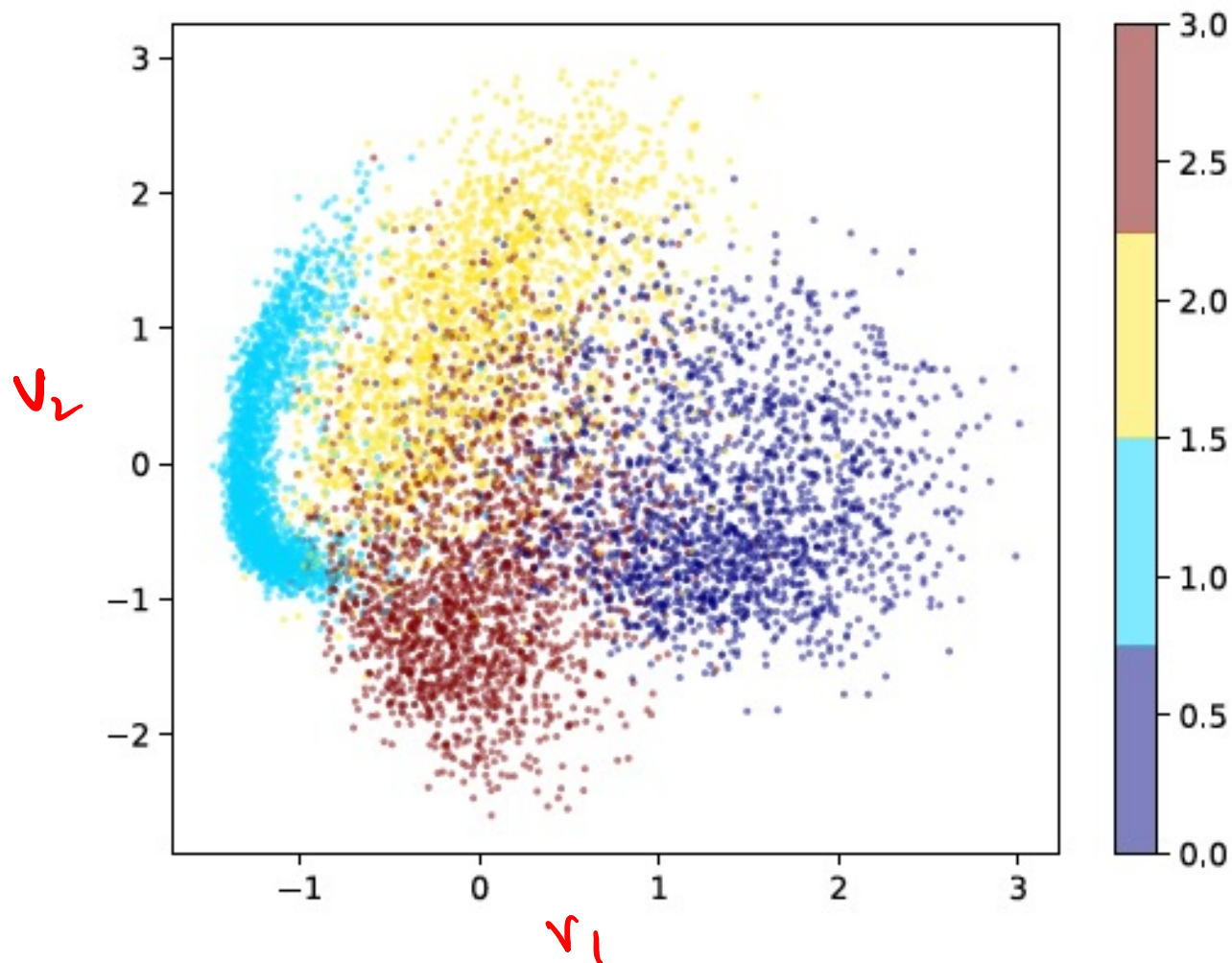


# Projecting MNIST digits

$$\begin{bmatrix} x_1 \\ \vdots \\ x_{784} \end{bmatrix} \rightarrow \begin{bmatrix} v_1^T x \\ v_2^T x \end{bmatrix}$$



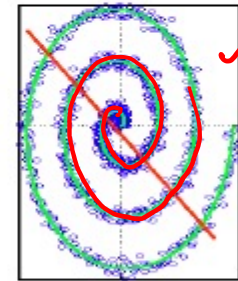
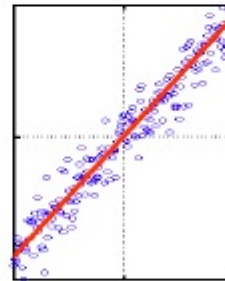
# Projecting MNIST digits



# Properties of PCA

- **Strengths**

- Eigenvector method
- No tuning parameters
- Non-iterative
- No local optima



- **Weaknesses**

- Limited to second order statistics
- Limited to linear projections

$X^T X$

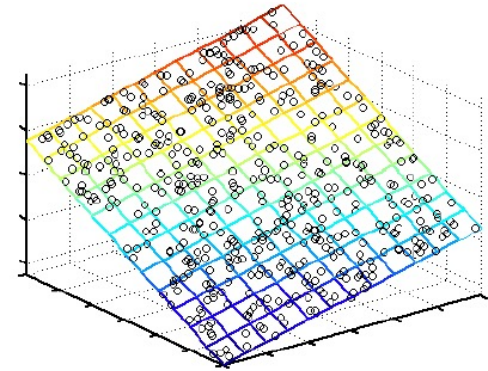
# Unsupervised Dimensionality Reduction

- Linear

**Principal Component Analysis (PCA)** ↙

Factor Analysis

Independent Component Analysis (ICA)



- Nonlinear

**Kernel PCA** ✓

Laplacian Eigenmaps, ISOMAP, LLE

Autoencoders

