

Foundations of Machine Learning

AI2000 and AI5000

FoML-29
PCA

Dr. Konda Reddy Mopuri

Department of AI, IIT Hyderabad
July-Nov 2025



భారతీయ సాంకేతిక విజ్ఞాన సంస్థ హైదరాబాద్
भारतीय प्रौद्योगिकी संस्थान हैदराबाद
Indian Institute of Technology Hyderabad



So far in FoML

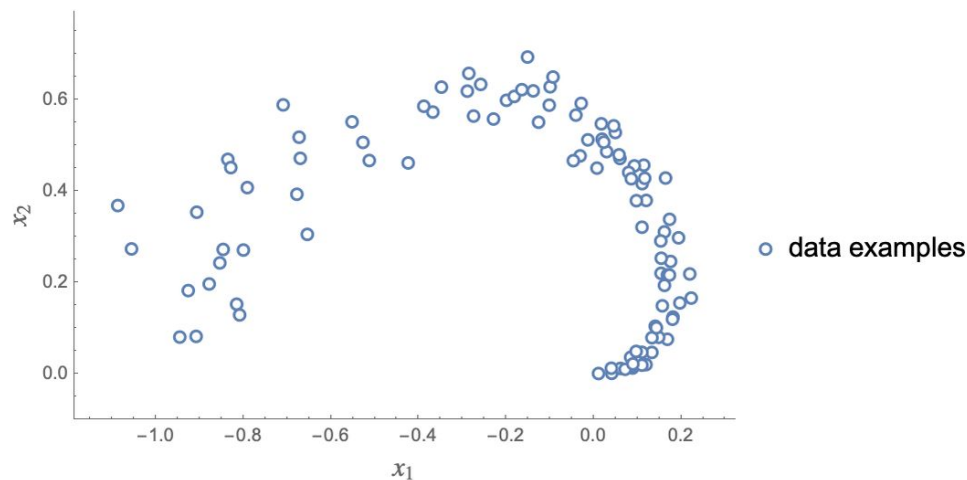
- Intro to ML and Probability refresher
- MLE, MAP, and fully Bayesian treatment
- Supervised learning
 - a. Linear Regression with basis functions
 - b. Bias-Variance Decomposition
 - c. Decision Theory - three broad classification strategies
 - d. Neural Networks
- Unsupervised learning
 - a. K-Means, Hierarchical, and GMM for clustering



For today

- PCA - Principal Component Analysis (Pearson, 1901) & (Hotelling, 1933)

Manifold coordinates as Latent variables



$$\{x_1, x_2\} = \{t \cos(3 t), t \sin(3 t)\}$$



Example - facial image data

- Possible degrees of freedom
 - Skull size
 - Skin color
 - Eye color
 - Facial attributes
 - Horizontal orientation
 - Vertical orientation
 - Mood (e.g., happy)
 - etc.



Example - facial image data

- Possible degrees of freedom
 - Skull size
 - Skin color
 - Eye color
 - Facial attributes
 - Horizontal orientation
 - Vertical orientation
 - Mood (e.g., happy)
 - etc.

Latent subspace will be a nonlinear transformation of image data



PCA

- Linear latent subspaces



What is PCA?

- Tool to summarize a large set of variables (dimensions) with a smaller set

What is PCA?

- Tool to summarize a large set of variables (dimensions) with a smaller set
 - Of 'representative' variables that explain the 'variability' in the original set



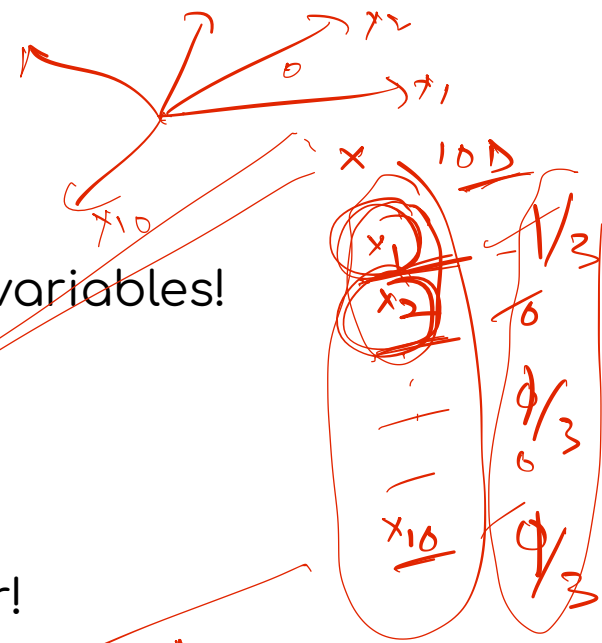
What is PCA?

- Tool to summarize a large set of variables (dimensions) with a smaller set
 - Of 'representative' variables that explain the 'variability' in the original set

$$\begin{matrix} 1 & 1 \\ X \end{matrix} \xrightarrow{\text{Linear}} \begin{matrix} 1 & 1 \\ Z \end{matrix}$$

What is PCA?

- Smaller set need not be a subset of original variables!
 - Rather, combinations of original variables
- They may not mean the same as originals
 - lost interpretability!
- New variables are independent of each other!

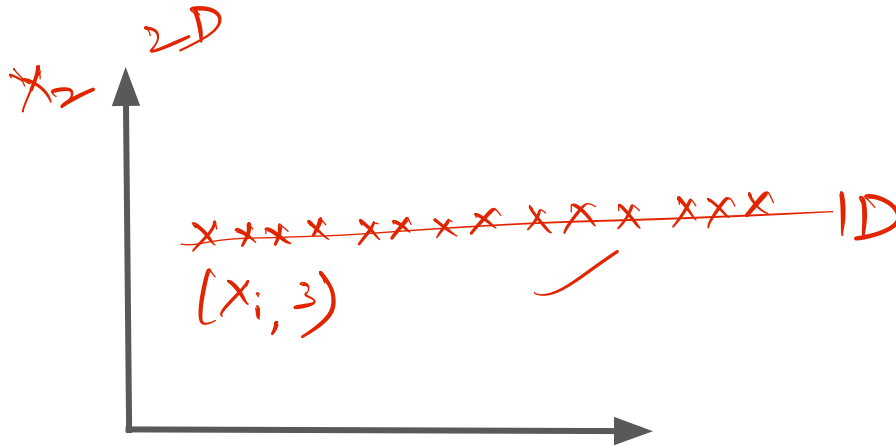


$$Z = W^T X$$



What is PCA?

- Gives the directions along which the data are highly 'variable'
 - Projects linearly such that the variance in the projected space is maximal



1-dim

$$\begin{pmatrix} x_1 \\ -1 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$



PCA

- Data $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$

$$\mathbf{x}_i \in \mathbb{R}^D$$

- Aim: project data onto M dimensional space ($M < D$) maximizing the variance of the projected data

PCA

- Mean

$$\bar{\mathbf{x}} = \frac{1}{N} \sum_{i=1}^N \mathbf{x}_i$$

- Covariance

$$\mathbf{S} = \frac{1}{N} \sum_{i=1}^N (\mathbf{x}_i - \bar{\mathbf{x}})(\mathbf{x}_i - \bar{\mathbf{x}})^T$$

Symmetric



PSD

$$\left\{ \begin{array}{l} \lambda_i \geq 0 \\ \text{lin ind Eig Vectors} \end{array} \right.$$



1D representation using PCA

- Project data onto a direction where most of the variance is preserved
- Projecting onto u_1 gives a scalar $\rightarrow$ 1D representation

$$\mathbf{z}_i = \mathbf{u}_1^T \mathbf{x}_i$$

Handwritten diagram illustrating the projection of data points $x_1, x_2, \dots, x_N$ onto the direction u_1 to produce 1D representations $z_1, z_2, \dots, z_N$. The diagram shows a vertical list of $x_1, x_2, \dots, x_N$ on the left, enclosed in a large curly bracket. To the right, a vertical list of $u_1^T x_1, u_1^T x_2, \dots, u_1^T x_N$ is shown, with arrows pointing from each x_i to its corresponding projection $u_1^T x_i$. Below the diagram, the text "1D representations" is written.



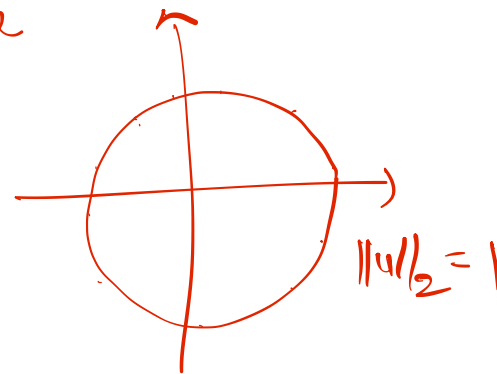
1D representation using PCA

- Direction of u_1 is important $\rightarrow$ consider unit vector in that direction

$$\|u_1\|_2 = 1$$



↓
But not its magnitude
w/o this constraint,
we can't optimize



1D representation using PCA

- Consider the variance in the new subspace

$$\begin{aligned}\text{Var}[\mathbf{z}] &= \frac{1}{N} \sum_{i=1}^N (\underline{z}_i - \bar{z})^2 = \frac{1}{N} \sum_{i=1}^N (\underline{u}_1^T \underline{x}_i - \underline{u}_1^T \bar{\mathbf{x}})^2 \\ &= \frac{1}{N} \sum_{i=1}^N [\underline{u}_1^T (\underline{x}_i - \bar{\mathbf{x}})]^2 = \frac{1}{N} \sum_{i=1}^N \underline{u}_1^T (\underline{x}_i - \bar{\mathbf{x}}) (\underline{x}_i - \bar{\mathbf{x}})^T \underline{u}_1 \\ &= \frac{1}{N} \underline{u}_1^T \underbrace{\sum_{i=1}^N (\underline{x}_i - \bar{\mathbf{x}}) (\underline{x}_i - \bar{\mathbf{x}})^T}_S \underline{u}_1 = \underline{u}_1^T S \underline{u}_1\end{aligned}$$



1D representation using PCA

1st principal component -
PC1

↓
 u_1

- Let's find the direction (u_1) that maximizes the variance in z_i

$$\arg \max_{u_1} u_1^T S u_1 \quad \text{such that} \quad u_1^T u_1 = 1$$

using Lagrange multiplier

$$\arg \max_{u_1, \lambda} \quad \frac{u_1^T S u_1 + \lambda(1 - u_1^T u_1)}{2} = 0$$

$$\nabla f(x) + \lambda \nabla g(x) = 0$$
$$f(x) + \lambda g(x)$$

$$2S u_1 - 2\lambda u_1 = 0$$

$$S u_1 = \lambda u_1$$

u_1 is eig. vector
 λ is corresponding eig. value of S



PCA via maximum variance

$$\frac{(x_i - z_i u_1)}{\|x_i - z_i u_1\|}$$

- Repeat the procedure for the next $M-1$ orthogonal vectors
 - Maximize the variance by projecting onto a direction orthogonal to the found ones
 - These are the next $M-1$ eigenvectors of the covariance matrix (S)

$$\underline{u_1^T S u_1 = u_1^T \lambda u_1 = \lambda u_1^T u_1 = \lambda}$$

The variance captured by a PC is its eigen value (λ_i)
(of the projected data)

PCA - Eigen decomposition

- For the symmetric PSD matrix $\mathbf{S} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^T$
- Eigenvectors are orthonormal (contained in $\mathbf{U}$)
- Eigenvalues are non-negative (contained in $\mathbf{\Lambda}$)



PCA - Eigen decomposition

- Variance is $\text{Tr}(S)$

of the projected data

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_D$$

$\rightarrow \text{Tr}[\text{Cov}[Z]] = \text{variance of}$

data along each of the new dims

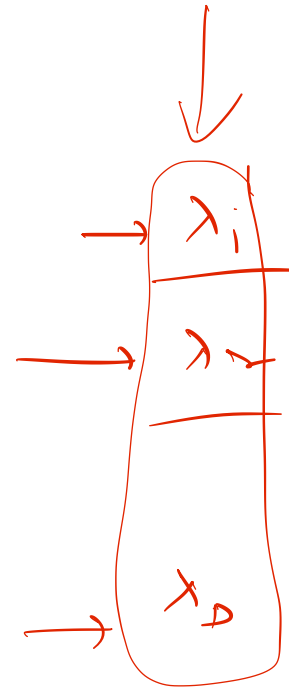
$$= \lambda_1 + \lambda_2 + \dots + \lambda_D$$

$$\underline{\underline{S = U \Lambda U^T}}$$

$$\rightarrow [u_1^T S u_1]$$

$$\rightarrow [u_2^T S u_2]$$

$$[u_D^T S u_D]$$



PCA - some notes



PCA - Scaling the features

- Sensitive to the scales of the features/variables
 - Features with greater range dominate the process of finding the PCs
- Perform standardization to prevent this

mean subtraction and variance normalization



PCA - Proportion of the Variance Explained

- How much of the information is lost by projecting onto PCs?

PCA - Proportion of the Variance Explained

- Total variance P = no. of dimensions in the original space

$$\sum_{j=1}^p \text{Var}(X_j) = \sum_{j=1}^p \frac{1}{n} \sum_{i=1}^n x_{ij}^2$$

↙
assume mean subtracted

PCA - Proportion of the Variance Explained

- Total variance
- Variance explained by the 'm'th PC

$$\sum_{j=1}^p \text{Var}(X_j) = \sum_{j=1}^p \frac{1}{n} \sum_{i=1}^n x_{ij}^2$$

$$\frac{1}{n} \sum_{i=1}^n z_{im}^2$$

linear transformation
doesn't disturb the mean
centering

PCA - Proportion of the Variance Explained

- Total variance
- Variance explained by the 'm'th PC

$$\sum_{j=1}^p \text{Var}(X_j) = \sum_{j=1}^p \frac{1}{n} \sum_{i=1}^n x_{ij}^2$$

$$\frac{1}{n} \sum_{i=1}^n z_{im}^2$$

PVE of the 'm'th PC =

↑

$$\frac{\lambda_m}{\sum_{i=1}^D \lambda_i}$$

up to m ✓

$$\frac{\sum_{i=1}^m \lambda_i}{\sum_{i=1}^D \lambda_i}$$



PCA - How many PCs to consider?

- For an $n \times p$ data matrix
 - $\min(n-1, p)$ PCs are possible
 - Why?

mean centering makes
the rows dependent,
hence $n-1$

$$X = \begin{bmatrix} \text{---} \\ \text{---} \\ \text{---} \\ \text{---} \end{bmatrix} \begin{matrix} 10 \\ \text{---} \\ \text{---} \\ \text{---} \end{matrix}$$

500 $n \times D$
 n - # Samples
 D - dims

$n > D$ generally
 $n < D$ rarely



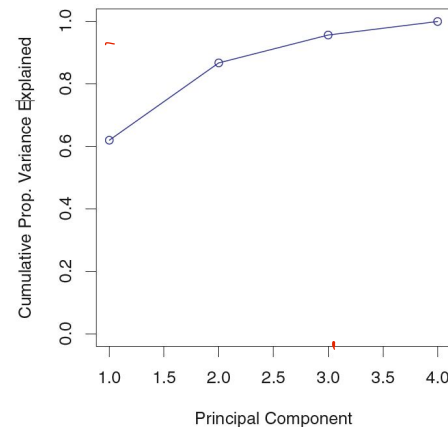
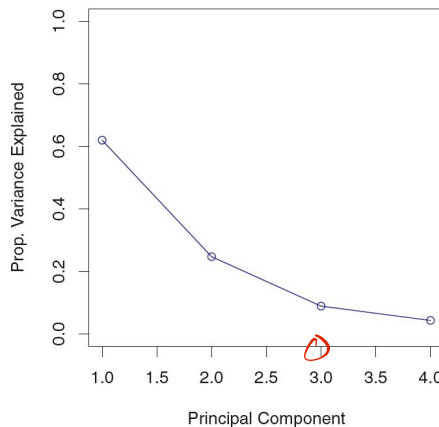
PCA - How many PCs to consider?

- For an $n \times p$ data matrix
 - $\min(n-1, p)$ PCs are possible
- Not all of them may be interesting

(leads to acceptable
loss of information)

PCA - How many PCs to consider?

- Generally, we want the 'smallest' number of them → good understanding of the data
- → scree plot & elbow



PCA

(not feature selection,
rather feature extraction)

- Doesn't discard the redundant variables
 - Finds new variables (linear combinations of the ' p ' variables) that summarize the data well
 - The 'best' variables (among the all possible linear combinations)
 - Resulting new features are uncorrelated (covariance matrix will be diagonal)

' p ' and ' D ' both indicate the original
dimension of the data

Applications of PCA

- Dimensionality reduction → tackles curse of dimensionality
- Less compute requirement
- Less prone to overfitting
- Useful preprocessing

Next

- PCA continued

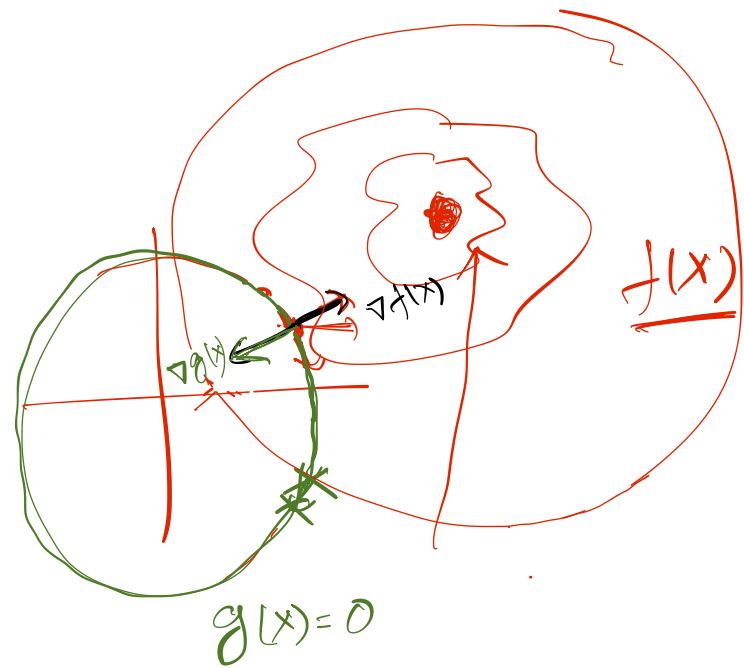
on the constraint set (green level set)
 $\nabla g(x)$ is $\perp$ to the level set surface

$$g(x+\epsilon) = g(x) + \epsilon^T \nabla g(x)$$

$g(x+\epsilon) = g(x)$ on the level set

$$\Rightarrow \epsilon^T \nabla g(x) = 0 \quad \nabla g(x) \perp \epsilon$$

ϵ is the movement along the level set



$\nabla f(x)$ will also be $\perp$ to the level set at x^* . Otherwise we can increase $f(x)$ by moving along the $g(x)$