



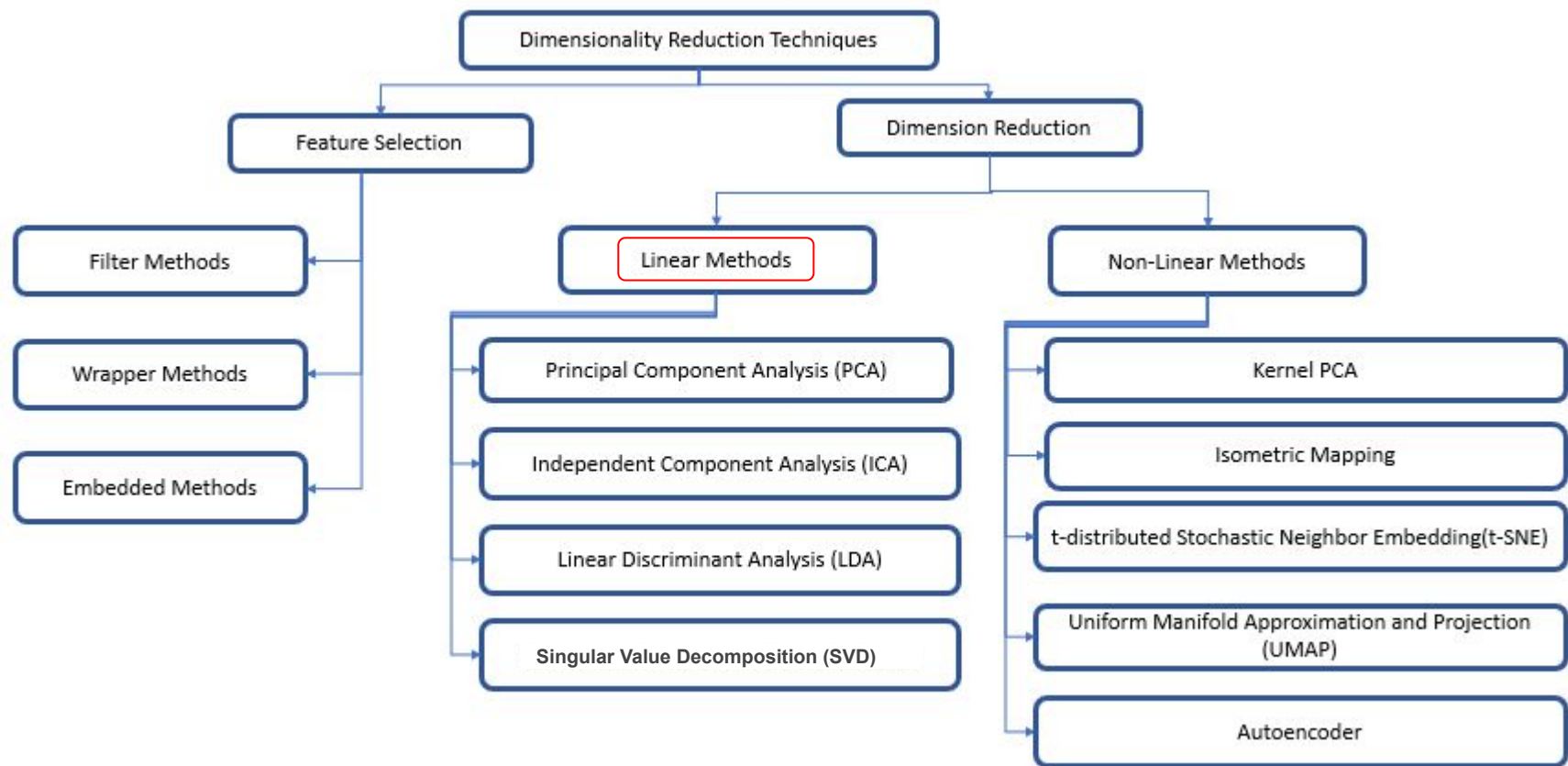
Εργαστήριο Συστημάτων Τεχνητής Νοημοσύνης και Μάθησης
Τομέας Τεχνολογίας Πληροφορικής και Υπολογιστών
Σχολή Ηλεκτρολόγων Μηχανικών και Μηχανικών Υπολογιστών

AILS

ΜΗΧΑΝΙΚΗ ΜΑΘΗΣΗ

Dimensionality Reduction

Τζούβελη Παρασκευή



Dimensionality Reduction

Taking a picture



Taking a picture



Dimensionality Reduction



Dimensionality Reduction



Ποια είναι η ιδανική γραμμή για προβολή των δεδομένων, όπου αυτά θα είναι όσο πιο διακριτά γίνεται;

Principal Component Analysis (PCA)

Στόχος: Εύρεση προβολών των σημείων δεδομένων $\mathbf{x}_n$ όπου:

- διατηρούν τη μέγιστη ομοιότητα με τα αρχικά δεδομένα και ελαχιστοποιούν το σφάλμα ανακατασκευής
- πετυχαίνουν χαμηλότερη διάσταση, μειώνοντας την πολυπλοκότητα και διατηρώντας τη δομή των δεδομένων.

π.χ. Housing Dataset

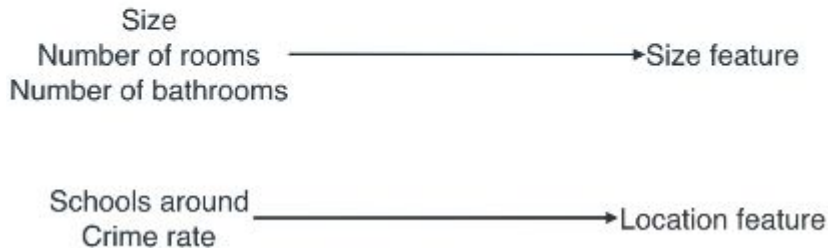
Dataset' features

Size
Number of rooms
Number of
bathrooms
Schools around
Crime rate

Διάνυσμα 5 χαρακτηριστικών

Μείωση διαστάσεων

Διάνυσμα 2 χαρακτηριστικών



Principal Component Analysis (PCA)

Προβολή από τον χώρο των 2 χαρακτηριστικών (2D) στο 1D

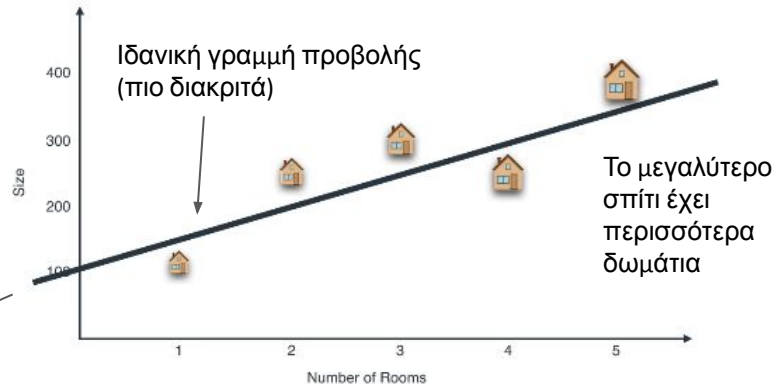
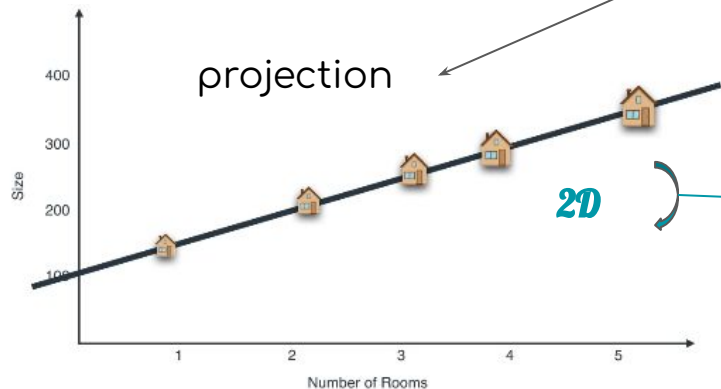
Housing Data

Size
Number of rooms
~~Number of bathrooms~~

Size feature

Schools around
Crime rate

Location feature



Principal Component Analysis (PCA): μετρικές

$$\sigma^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{N}$$

where:

x_i = Each value in the data set

$\bar{x}$ = Mean of all values in the data set

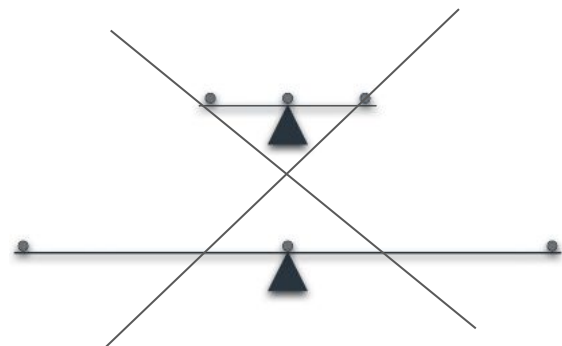
N = Number of values in the data set

Mean

σημείο
ισορροπίας



$$\text{Mean} = \frac{1+2+6}{3} = 3$$

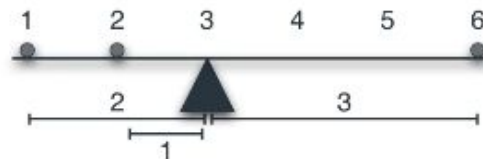


Ίδιο mean value, →
δεν είναι βολική
μετρική

Variance

$$\text{Variance} = \frac{1^2 + 0^2 + 1^2}{3} = 2/3$$

$$\text{Variance} = \frac{5^2 + 0^2 + 5^2}{3} = 50/3$$



$$\text{Variance} = \frac{2^2 + 1^2 + 3^2}{3} = 14/3$$

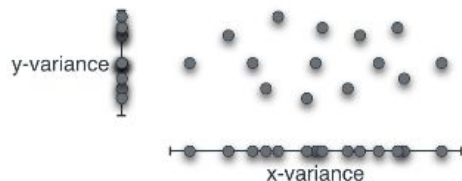
Υπολογισμός:
παίρνω την
απόσταση κάθε
σημείου από το
κέντρο.

Principal Component Analysis (PCA): μετρικές

Εξετάζουμε τη διακύμανση στο χώρο

→ εξετάζουμε και τις δύο διαστάσεις

Variance?



$$\text{Var}(X) = \mathbb{E}[(X - \mu)^2]$$

where:

- X : Random variable.
- $\mu = \mathbb{E}[X]$: Mean (expected value) of X .
- $\mathbb{E}$: Expectation operator.

Alternatively, it can be expressed as:

$$\text{Var}(X) = \mathbb{E}[X^2] - (\mathbb{E}[X])^2$$

$$\text{var}[f] = \mathbb{E} \left[(f(x) - \mathbb{E}[f(x)])^2 \right] = \mathbb{E}[f(x)^2] - \mathbb{E}[f(x)]^2$$

Variance?



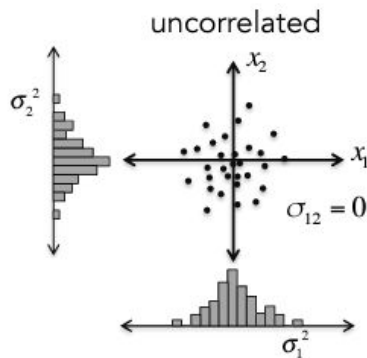
$$\text{x-variance} = \frac{2^2 + 0^2 + 2^2}{3} = 8/3$$

$$\text{y-variance} = \frac{1^2 + 0^2 + 1^2}{3} = 2/3$$

Principal Component Analysis (PCA): μετρικές

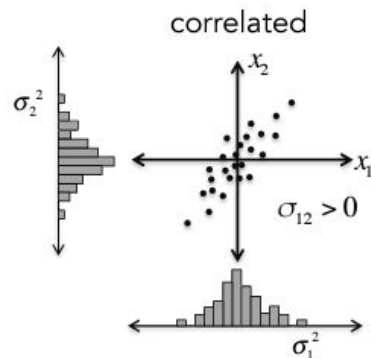
Αν, όμως, έχω δύο διαφορετικά σύνολα μπορεί να έχουν την ίδια διακύμανση στις προβολές x_1 και $x_2 \rightarrow$ μη διαχωρίσιμα

Άρα χρειάζεται άλλη μετρική για να περιγράψουμε την μεταξύ τους σχέση.



covariance

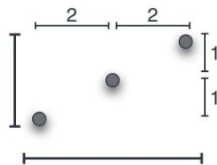
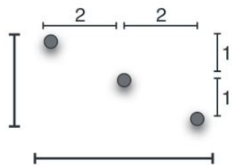
$$\sigma_{12} = \frac{1}{m} \sum_{j=1}^m (x_1^{(j)} - \mu_1)(x_2^{(j)} - \mu_2)$$



correlation

$$\rho = \frac{\sigma_{12}}{\sqrt{\sigma_1^2 \sigma_2^2}}$$

Principal Component Analysis (PCA): μετρικές



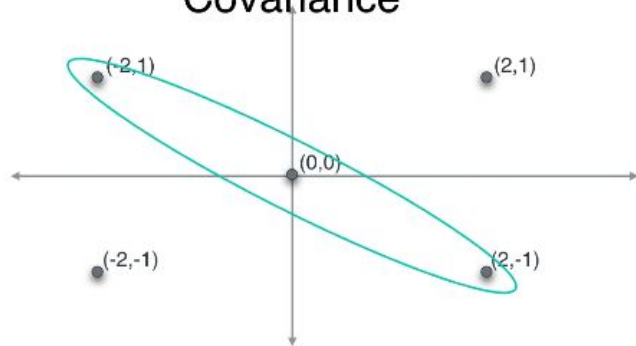
$$\text{cov}[x, y] = \mathbb{E}_{x,y} [\{x - \mathbb{E}[x]\} \{y - \mathbb{E}[y]\}]$$

$$= \mathbb{E}_{x,y} [xy] - \mathbb{E}[x]\mathbb{E}[y]$$

$$\text{cov}[\mathbf{x}, \mathbf{y}] = \mathbb{E}_{\mathbf{x}, \mathbf{y}} [\{\mathbf{x} - \mathbb{E}[\mathbf{x}]\} \{\mathbf{y}^T - \mathbb{E}[\mathbf{y}^T]\}]$$

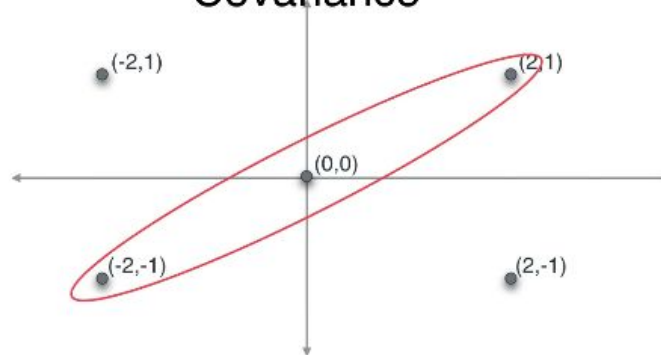
$$= \mathbb{E}_{\mathbf{x}, \mathbf{y}} [\mathbf{x}\mathbf{y}^T] - \mathbb{E}[\mathbf{x}]\mathbb{E}[\mathbf{y}^T]$$

Covariance

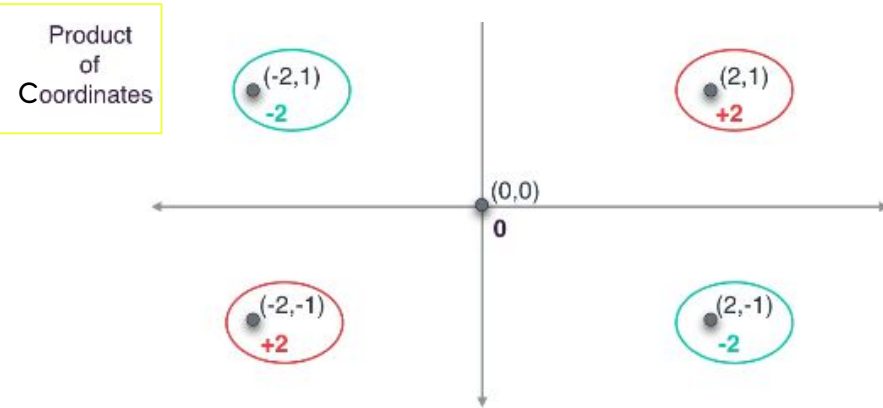


Διαφορά;

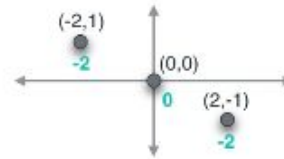
Covariance



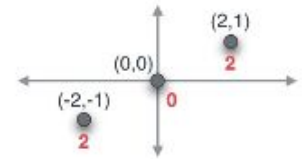
Principal Component Analysis (PCA): μετρικές



Covariance



$$\text{covariance} = \frac{(-2) + 0 + (-2)}{3} = -4/3$$



$$\text{covariance} = \frac{2 + 0 + 2}{3} = 4/3$$

$$\text{Cov}(X, Y) = \mathbb{E}[(X - \mu_X)(Y - \mu_Y)]$$

where:

- $\mu_X = \mathbb{E}[X]$: Mean of X .
- $\mu_Y = \mathbb{E}[Y]$: Mean of Y .
- $\mathbb{E}$: Expectation operator.

Alternatively, it can be expressed as:

$$\text{Cov}(X, Y) = \mathbb{E}[XY] - \mathbb{E}[X]\mathbb{E}[Y]$$

If X and Y are independent, then $\text{Cov}(X, Y) = 0$.

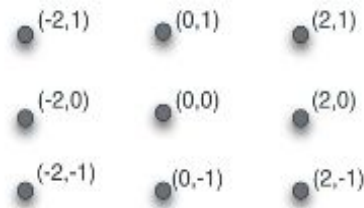
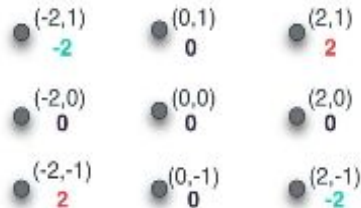
$$\text{Cov}(X, Y) = \frac{\sum (X_i - \bar{X})(Y_j - \bar{Y})}{n}$$

Principal Component Analysis (PCA): μετρικές

Ας υπολογίσουμε το covariance των σημείων:

Είναι 0 και αυτό φαίνεται λογικό αφού δεν φαίνεται να υπάρχει κάποια θετική ή αρνητική συνδιακύμανση

Covariance



$$\text{covariance} = \frac{-2 + 0 + 2 + 0 + 0 + 0 + 0 + 2 + 0 + -2}{9} = 0$$

Principal Component Analysis (PCA): μετρικές

Covariance



negative
covariance

Αν το x αυξάνεται,
τότε το y μειώνεται



covariance zero
(or very small)

Το x ανεξάρτητο του y

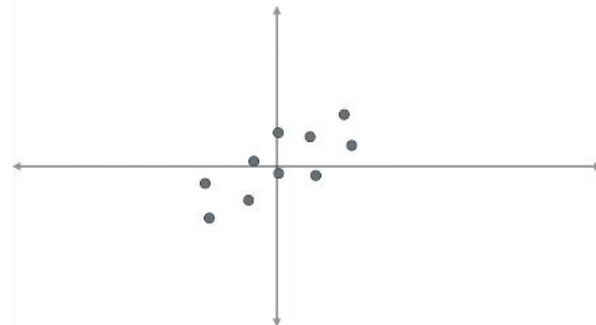
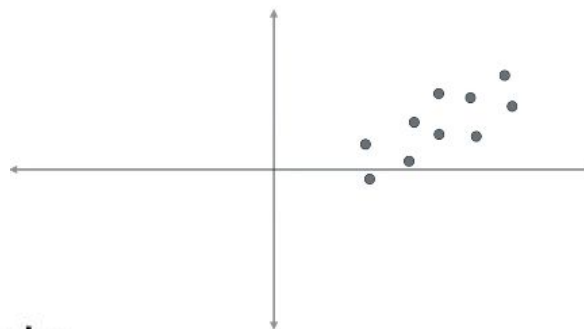


positive
covariance

Τα x και y αυξομειώνονται μαζί

Principal Component Analysis (PCA): Covariance Matrix

Πώς θα βρούμε την τέλεια προβολή; Θα ορίσουμε σύστημα αξόνων → θα βρούμε το σημείο όπου τα δεδομένα ισορροπούν



Covariance matrix

$$\Sigma = \begin{pmatrix} \text{Var}(X) & \text{Cov}(X,Y) \\ \text{Cov}(X,Y) & \text{Var}(Y) \end{pmatrix}$$

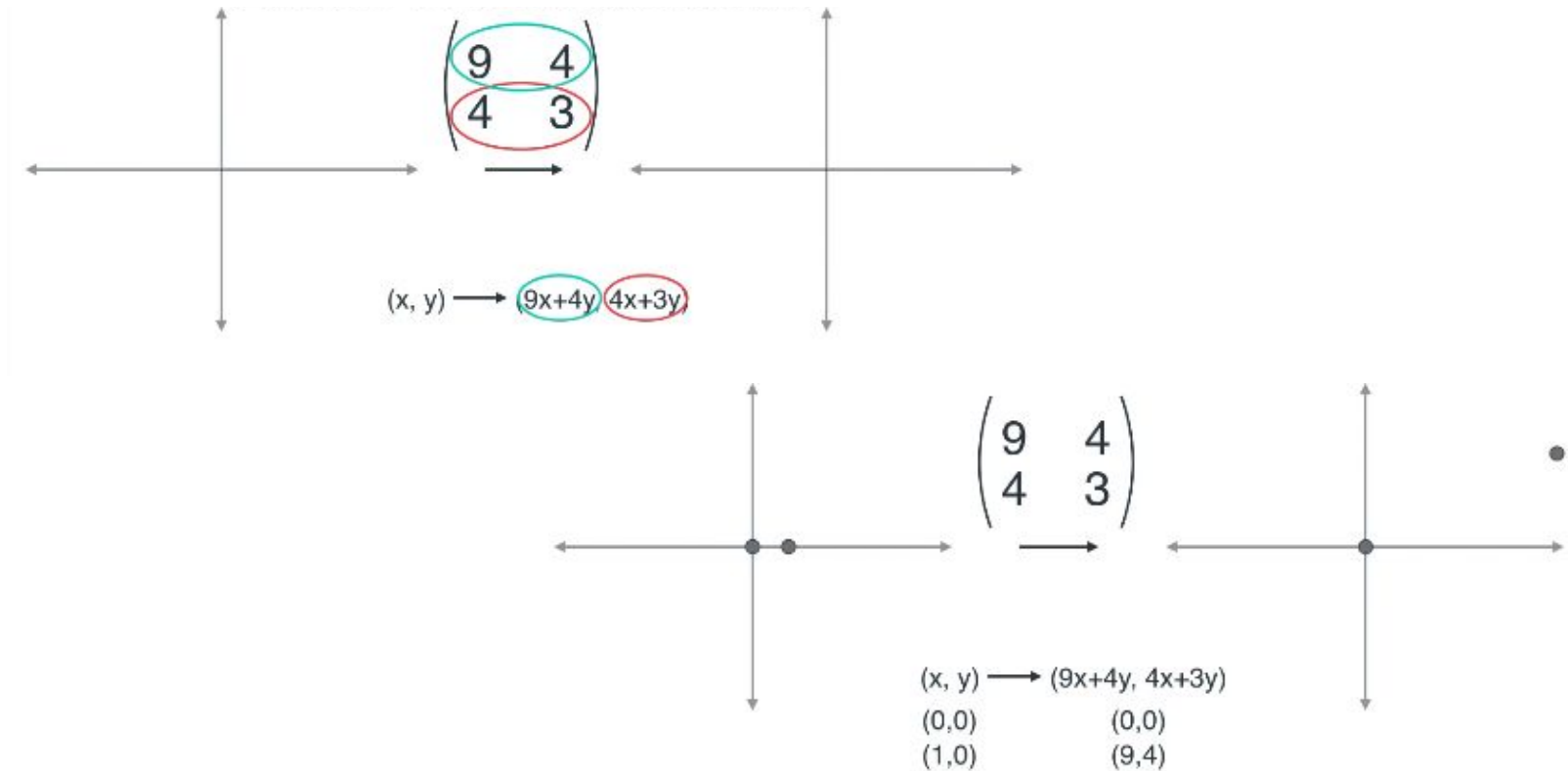
$$\text{Έστω } \Sigma = \begin{pmatrix} 9 & 4 \\ 4 & 3 \end{pmatrix}$$

Υπενθύμιση: ο Covariance matrix είναι θετικά ημιορισμένος (έλεγχος: συμμετρικός: $C=C^T$, ιδιοτιμές $\lambda_i \geq 0$ και για κάθε μη μηδενικό διάνυσμα $x \in \mathbb{R}^{n \times n}$, $x^T C x \geq 0$)

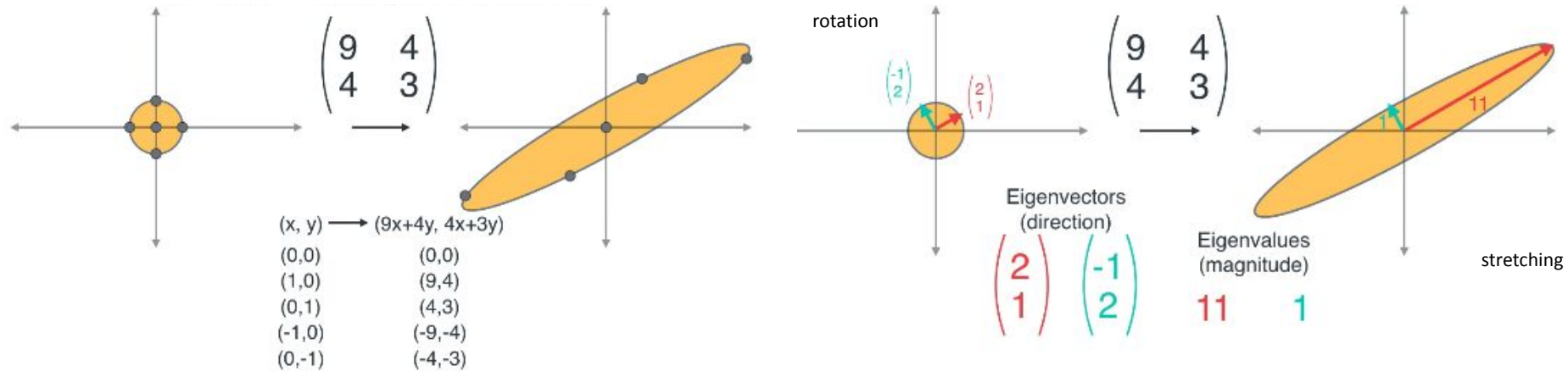
→ εξασφαλίζει ότι τα δεδομένα μπορούν να αναλυθούν σε έναν υποχώρο με βάση τη μη αρνητική διακύμανσή τους.

Παρατηρούμε ότι παρουσιάζει μεγάλη διακύμανση στο x άξονα σε αντίθεση με τον y άξονα

Principal Component Analysis (PCA): Linear Transformation



Principal Component Analysis (PCA): Linear Transformation



Eigenvalues
(magnitude)

$$\begin{vmatrix} x-9 & -4 \\ -4 & x-3 \end{vmatrix} = (x-9)(x-3) - (-4)(-4) = x^2 - 12x + 11$$

$$= (x-11)(x-1) \rightarrow \mathbf{11} \quad \mathbf{1}$$

Eigenvectors
(direction)

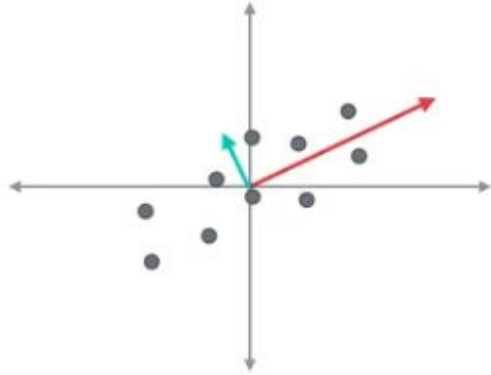
$$\begin{pmatrix} 9 & 4 \\ 4 & 3 \end{pmatrix} \begin{pmatrix} u \\ v \end{pmatrix} = \mathbf{11} \begin{pmatrix} u \\ v \end{pmatrix}$$

$$\begin{pmatrix} u \\ v \end{pmatrix} = \mathbf{\begin{pmatrix} 2 \\ 1 \end{pmatrix}}$$

$$\begin{pmatrix} 9 & 4 \\ 4 & 3 \end{pmatrix} \begin{pmatrix} u \\ v \end{pmatrix} = \mathbf{1} \begin{pmatrix} u \\ v \end{pmatrix}$$

$$\begin{pmatrix} u \\ v \end{pmatrix} = \mathbf{\begin{pmatrix} -1 \\ 2 \end{pmatrix}}$$

Principal Component Analysis (PCA)



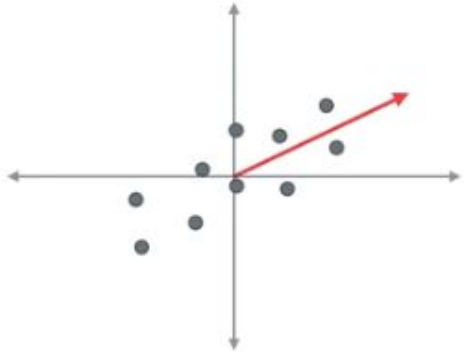
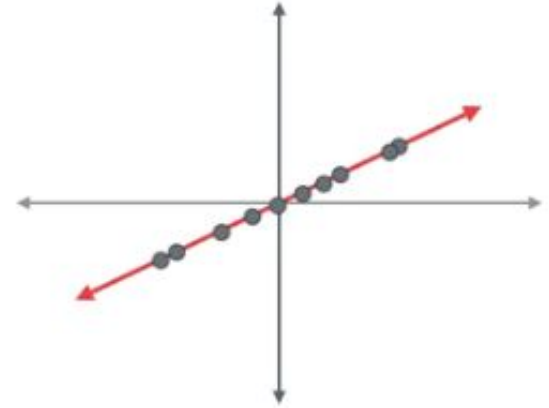
$$\Sigma = \begin{pmatrix} 9 & 4 \\ 4 & 3 \end{pmatrix}$$

$$\begin{pmatrix} 2 \\ 1 \end{pmatrix} \quad \begin{pmatrix} -1 \\ 2 \end{pmatrix}$$

Eigenvectors
(direction)

$$11 \quad 1$$

Eigenvalues
(magnitude)



$$\Sigma = \begin{pmatrix} 9 & 4 \\ 4 & 3 \end{pmatrix}$$

$$\begin{pmatrix} 2 \\ 1 \end{pmatrix}$$

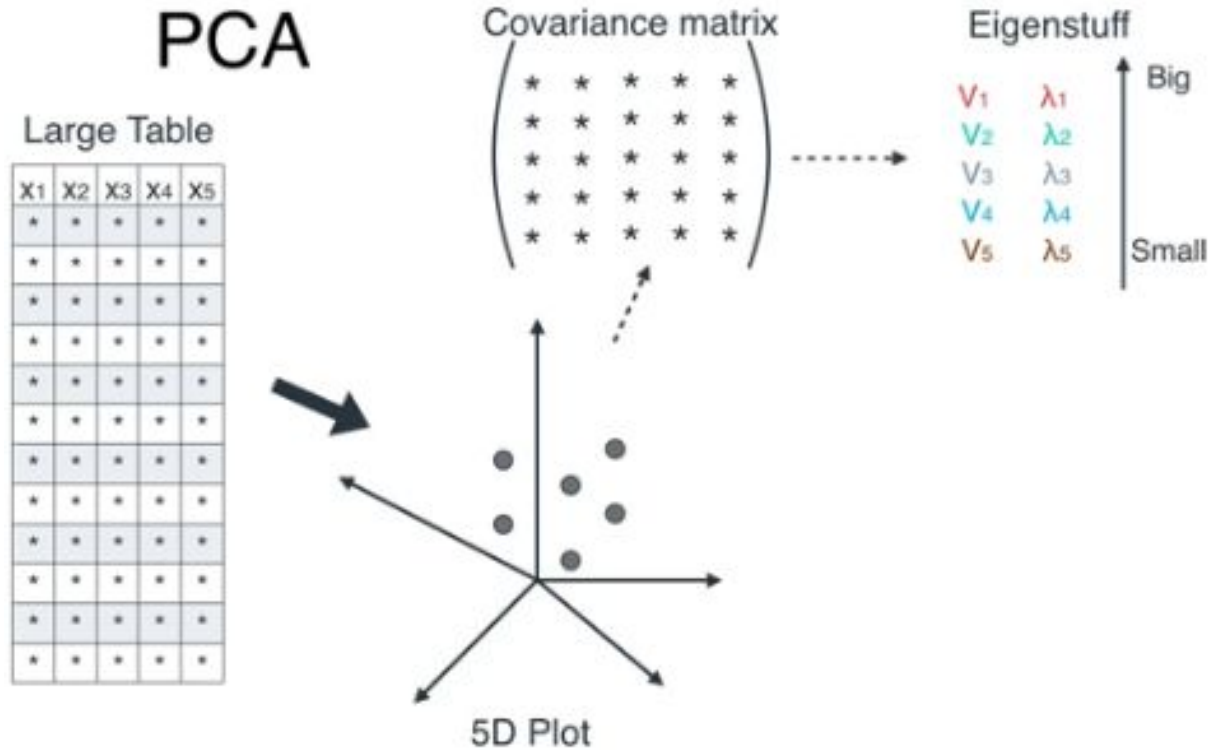
Eigenvectors
(direction)

$$11$$

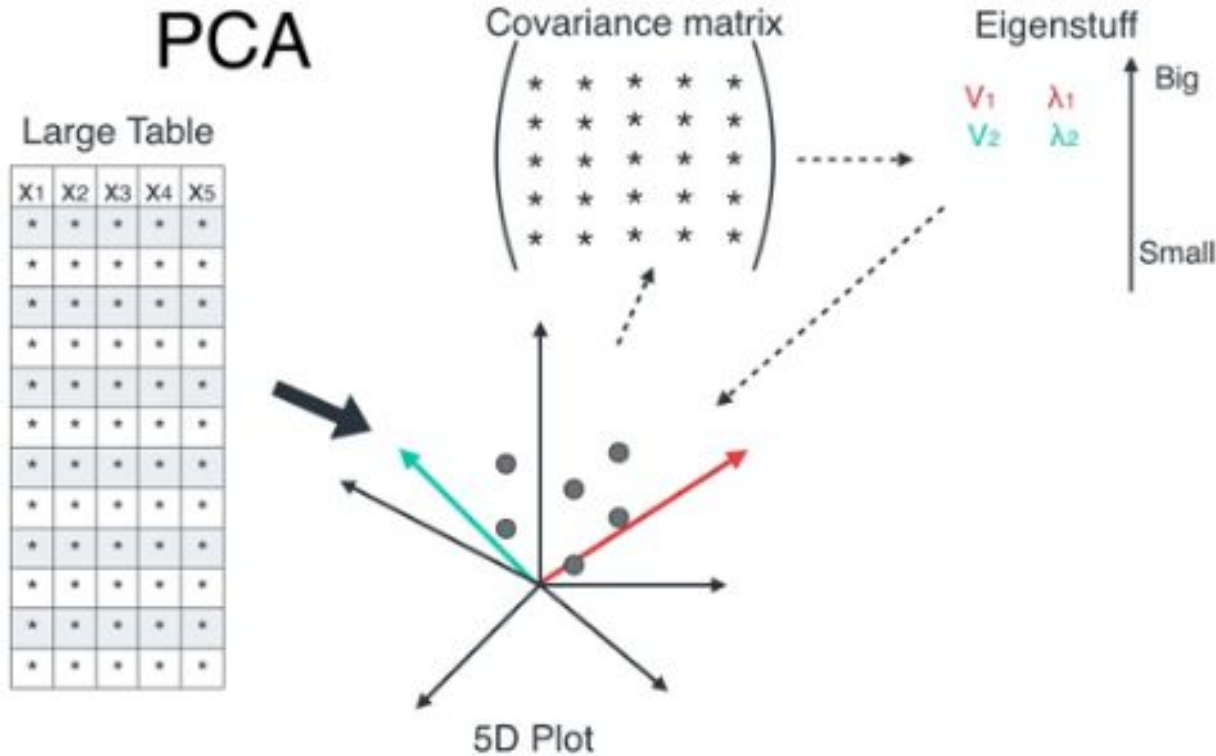
Eigenvalues
(magnitude)



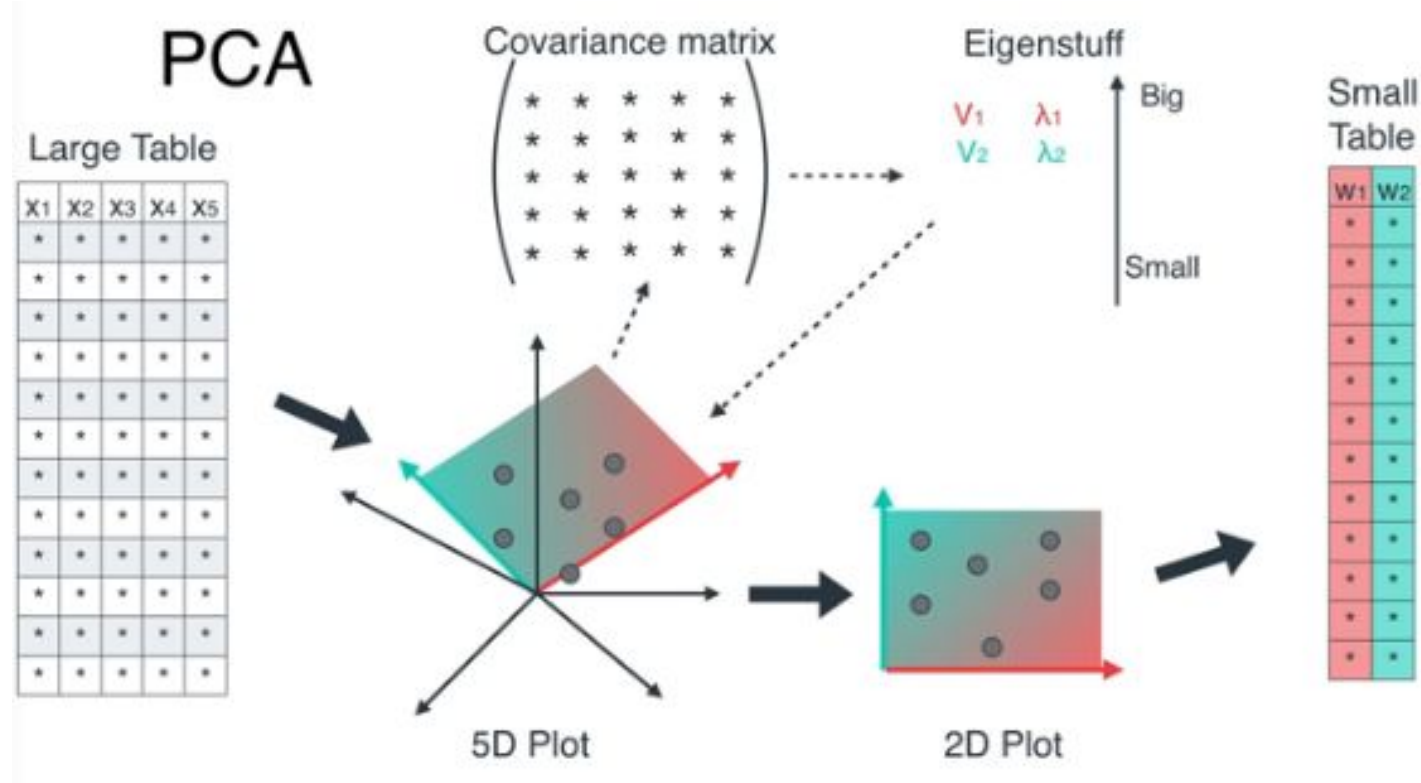
Principal Component Analysis (PCA): 5D features → 2D features



Principal Component Analysis (PCA): 5D features → 2D features

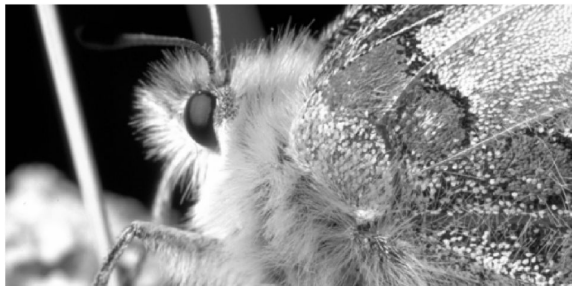


Principal Component Analysis (PCA) 5D features → 2D features

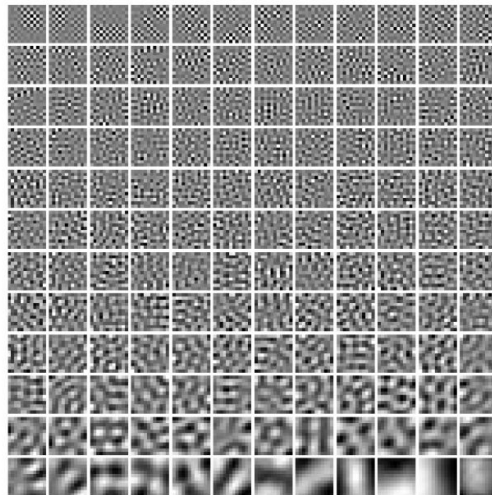


Εφαρμογές PCA : Συμπίεση εικόνας

PCA example — original



PCA example — PCA basis



PCA example — 90% compressed



PCA example — 50% compressed



Principal Component Analysis (PCA) : Projection onto a subspace

Πρόβλημα: χωρητικότητα δεδομένων

Λύση: Συμπίεσης δεδομένων με τρόπο που απαιτεί λιγότερη μνήμη, αλλά χωρίς να χάνει πολύ σε ακρίβεια.

Στόχος: αναπαράσταση σε μικρότερες διαστάσεις, χωρίς να χάνουμε πολύ σε ακρίβεια

Set-up: given a dataset $\mathcal{D} = \{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(N)}\} \subset \mathbb{R}^D$

Set $\boldsymbol{\mu}$ to the mean of the data, $\boldsymbol{\mu} = \frac{1}{N} \sum_{i=1}^N \mathbf{x}^{(i)}$

Goal: find a K -dimensional subspace $\mathcal{S} \subset \mathbb{R}^D$ such that $\mathbf{x}^{(n)} - \boldsymbol{\mu}$ is “well-represented” by its projection onto $\mathcal{S}$

Recall: The **projection** of a point $\mathbf{x}$ onto $\mathcal{S}$ is the point in $\mathcal{S}$ closest to $\mathbf{x}$.

Principal Component Analysis (PCA): Projection onto a subspace

Let $\{\mathbf{u}_k\}_{k=1}^K$ be an orthonormal basis of the subspace $\mathcal{S}$

Approximate each data point $\mathbf{x}$ as:

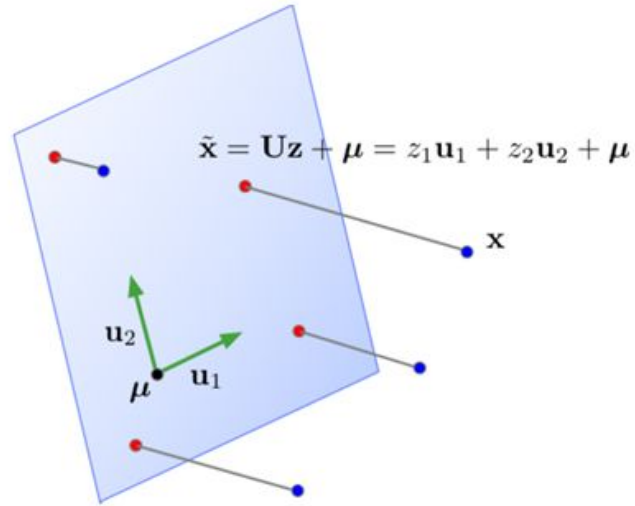
$$\begin{aligned}\tilde{\mathbf{x}} &= \boldsymbol{\mu} + \text{Proj}_{\mathcal{S}}(\mathbf{x} - \boldsymbol{\mu}) \\ &= \boldsymbol{\mu} + \sum_{k=1}^K z_k \mathbf{u}_k\end{aligned}$$

From linear algebra: $z_k = \mathbf{u}_k^T (\mathbf{x} - \boldsymbol{\mu})$

Let $\mathbf{U}$ be a matrix with columns $\{\mathbf{u}_k\}_{k=1}^K$ then $\mathbf{z} = \mathbf{U}^T (\mathbf{x} - \boldsymbol{\mu})$

Also: $\tilde{\mathbf{x}} = \boldsymbol{\mu} + \mathbf{U}\mathbf{z}$

Principal Component Analysis (PCA) : Projection onto a Subspace



$$\mathbf{z} = \mathbf{U}^\top (\mathbf{x} - \boldsymbol{\mu})$$

In machine learning, $\tilde{\mathbf{x}}$ is also called the **reconstruction** of $\mathbf{x}$.
 $\mathbf{z}$ is its **representation**, or **code**.

Principal Component Analysis (PCA) : Learning a Subspace

How to choose a good subspace $\mathcal{S}$?

- Need to choose $D \times K$ matrix $\mathbf{U}$ with orthonormal columns.

Two criteria:

- Minimize the **reconstruction error**

$$\min \frac{1}{N} \sum_{i=1}^N \|\mathbf{x}^{(i)} - \tilde{\mathbf{x}}^{(i)}\|^2$$

- Maximize the variance of the code vectors

$$\begin{aligned} \max \sum_j \text{Var}(z_j) &= \frac{1}{N} \sum_j \sum_i (z_j^{(i)} - \bar{z}_j)^2 \\ &= \frac{1}{N} \sum_i \|\mathbf{z}^{(i)} - \bar{\mathbf{z}}\|^2 \\ &= \frac{1}{N} \sum_i \|\mathbf{z}^{(i)}\|^2 \end{aligned}$$

Exercise: show $\bar{\mathbf{z}} = \mathbf{0}$

- Note: here, $\bar{\mathbf{z}}$ denotes the mean, not a derivative.

Principal Component Analysis (PCA) : Learning a Subspace

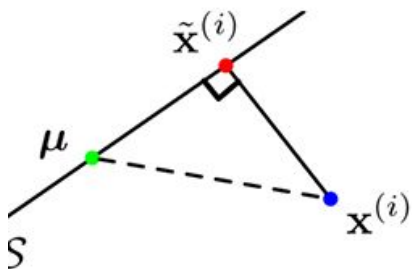
These two criteria are equivalent! I.e., we'll show

$$\frac{1}{N} \sum_{i=1}^N \|\mathbf{x}^{(i)} - \tilde{\mathbf{x}}^{(i)}\|^2 = \text{const} - \frac{1}{N} \sum_i \|\mathbf{z}^{(i)}\|^2$$

Observation: by unitarity,

$$\|\tilde{\mathbf{x}}^{(i)} - \boldsymbol{\mu}\| = \|\mathbf{U}\mathbf{z}^{(i)}\| = \|\mathbf{z}^{(i)}\|$$

By the Pythagorean Theorem,



$$\underbrace{\frac{1}{N} \sum_{i=1}^N \|\tilde{\mathbf{x}}^{(i)} - \boldsymbol{\mu}\|^2}_{\text{projected variance}} + \underbrace{\frac{1}{N} \sum_{i=1}^N \|\mathbf{x}^{(i)} - \tilde{\mathbf{x}}^{(i)}\|^2}_{\text{reconstruction error}} \\ = \underbrace{\frac{1}{N} \sum_{i=1}^N \|\mathbf{x}^{(i)} - \boldsymbol{\mu}\|^2}_{\text{constant}}$$

Principal Component Analysis (PCA)

Choosing a subspace to maximize the projected variance, or minimize the reconstruction error, is called **principal component analysis (PCA)**.

Recall:

- **Spectral Decomposition**: a symmetric matrix $\mathbf{A}$ has a full set of eigenvectors, which can be chosen to be orthogonal. This gives a decomposition

$$\mathbf{A} = \mathbf{Q}\mathbf{\Lambda}\mathbf{Q}^T,$$

where $\mathbf{Q}$ is orthogonal and $\mathbf{\Lambda}$ is diagonal. The columns of $\mathbf{Q}$ are eigenvectors, and the diagonal entries λ_j of $\mathbf{\Lambda}$ are the corresponding eigenvalues.

- I.e., symmetric matrices are diagonal in some basis.
- A symmetric matrix $\mathbf{A}$ is positive semidefinite iff each $\lambda_j \geq 0$.

Principal Component Analysis (PCA)

Consider the **empirical covariance matrix**:

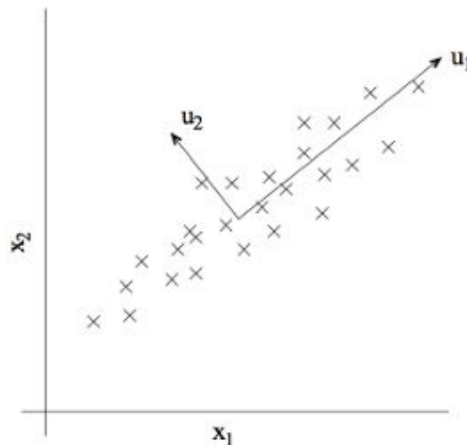
$$\Sigma = \frac{1}{N} \sum_{i=1}^N (\mathbf{x}^{(i)} - \mu)(\mathbf{x}^{(i)} - \mu)^\top$$

Recall: Covariance matrices are symmetric and positive semidefinite.

The optimal PCA subspace is spanned by the top K eigenvectors of Σ .

- More precisely, choose the first K of any orthonormal eigenbasis for Σ .
- The general case is tricky, but we'll show this for $K = 1$.

These eigenvectors are called **principal components**, analogous to the principal axes of an ellipse.



Principal Component Analysis (PCA)

For $K = 1$, we are fitting a unit vector $\mathbf{u}$, and the code is a scalar $z = \mathbf{u}^\top (\mathbf{x} - \boldsymbol{\mu})$.

$$\begin{aligned}\frac{1}{N} \sum_i [z^{(i)}]^2 &= \frac{1}{N} \sum_i (\mathbf{u}^\top (\mathbf{x}^{(i)} - \boldsymbol{\mu}))^2 \\&= \frac{1}{N} \sum_{i=1}^N \mathbf{u}^\top (\mathbf{x}^{(i)} - \boldsymbol{\mu}) (\mathbf{x}^{(i)} - \boldsymbol{\mu})^\top \mathbf{u} \\&= \mathbf{u}^\top \left[\frac{1}{N} \sum_{i=1}^N (\mathbf{x}^{(i)} - \boldsymbol{\mu}) (\mathbf{x}^{(i)} - \boldsymbol{\mu})^\top \right] \mathbf{u} \\&= \mathbf{u}^\top \boldsymbol{\Sigma} \mathbf{u} \\&= \mathbf{u}^\top \mathbf{Q} \boldsymbol{\Lambda} \mathbf{Q}^\top \mathbf{u} \\&= \mathbf{a}^\top \boldsymbol{\Lambda} \mathbf{a} \\&= \sum_{j=1}^D \lambda_j a_j^2\end{aligned}$$

Spectral Decomposition

for $\mathbf{a} = \mathbf{Q}^\top \mathbf{u}$

Principal Component Analysis (PCA)

Maximize $\mathbf{a}^\top \mathbf{\Lambda} \mathbf{a} = \sum_{j=1}^D \lambda_j a_j^2$ for $\mathbf{a} = \mathbf{Q}^\top \mathbf{u}$.

- This is a change-of-basis to the eigenbasis of $\mathbf{\Sigma}$.

Assume the λ_i are in sorted order. For simplicity, assume they are all distinct.

Observation: since $\mathbf{u}$ is a unit vector, then by unitarity, $\mathbf{a}$ is also a unit vector. I.e., $\sum_j a_j^2 = 1$.

By inspection, set $a_1 = \pm 1$ and $a_j = 0$ for $j \neq 1$.

Hence, $\mathbf{u} = \mathbf{Q} \mathbf{a} = \mathbf{q}_1$ (the top eigenvector).

A similar argument shows that the k th principal component is the k th eigenvector of $\mathbf{\Sigma}$. If you're interested, look up the [Courant-Fischer Theorem](#).

Principal Component Analysis (PCA)

Interesting fact: the dimensions of $\mathbf{z}$ are decorrelated. For now, let Cov denote the empirical covariance.

$$\begin{aligned}\text{Cov}(\mathbf{z}) &= \text{Cov}(\mathbf{U}^\top (\mathbf{x} - \boldsymbol{\mu})) \\ &= \mathbf{U}^\top \text{Cov}(\mathbf{x}) \mathbf{U} \\ &= \mathbf{U}^\top \boldsymbol{\Sigma} \mathbf{U} \\ &= \mathbf{U}^\top \mathbf{Q} \boldsymbol{\Lambda} \mathbf{Q}^\top \mathbf{U} \\ &= \begin{pmatrix} \mathbf{I} & \mathbf{0} \end{pmatrix} \boldsymbol{\Lambda} \begin{pmatrix} \mathbf{I} \\ \mathbf{0} \end{pmatrix} && \text{by orthogonality} \\ &= \text{top left } K \times K \text{ block of } \boldsymbol{\Lambda}\end{aligned}$$

If the covariance matrix is diagonal, this means the features are uncorrelated.

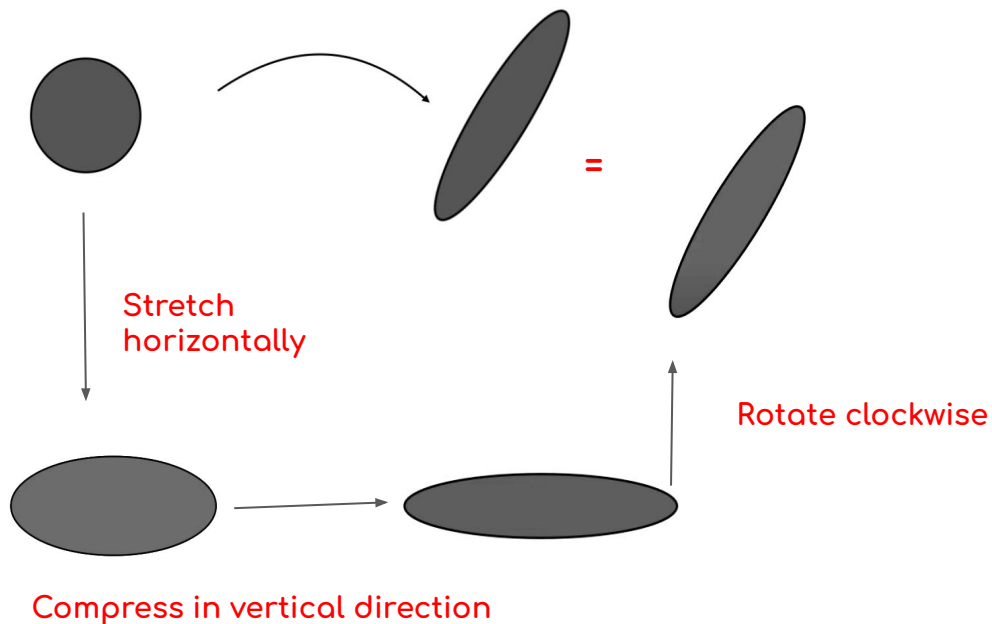
This is why PCA was originally invented (in 1901!).

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Μετασχηματισμοί: μόνο ως προς οριζόντια και κάθετα κατεύθυνση, όχι υπό άλλη γωνία

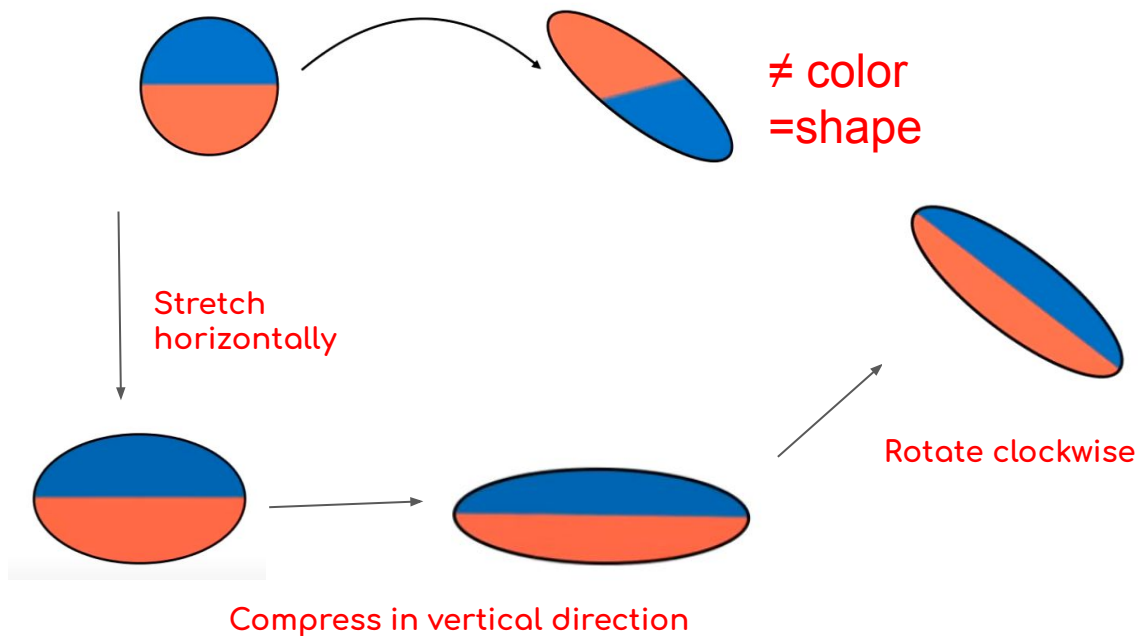
?

Puzzle (easy)



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Puzzle (hard)

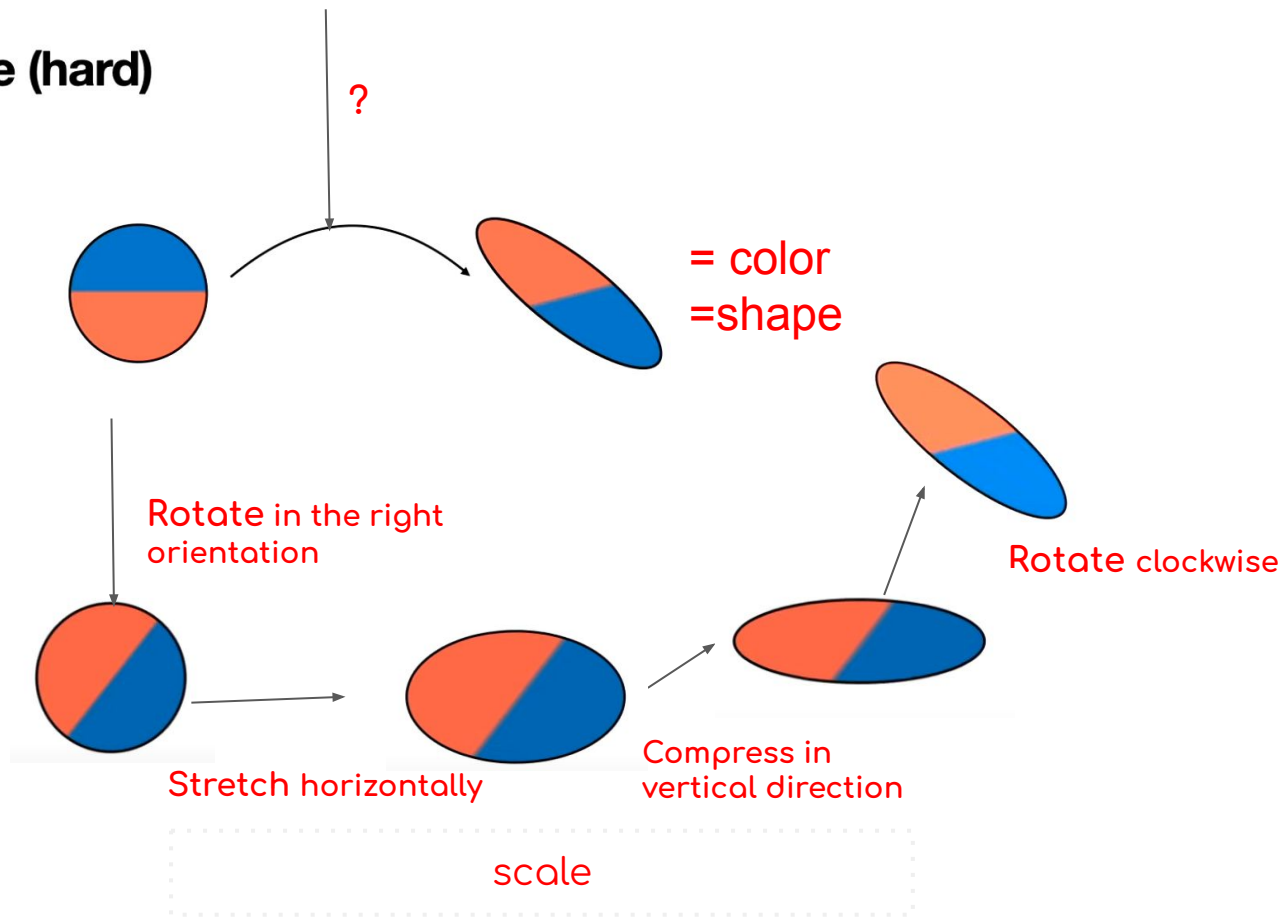


Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Puzzle (hard)

Λύση

Με περιστροφή,
κλιμάκωση και
περιστροφή
μπορούμε να
μιμηθούμε κάθε
γραμμικό
μετασχηματισμό



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Γραμμικός μετασχηματισμός

$$A = \begin{bmatrix} 3 & 0 \\ 4 & 5 \end{bmatrix}$$

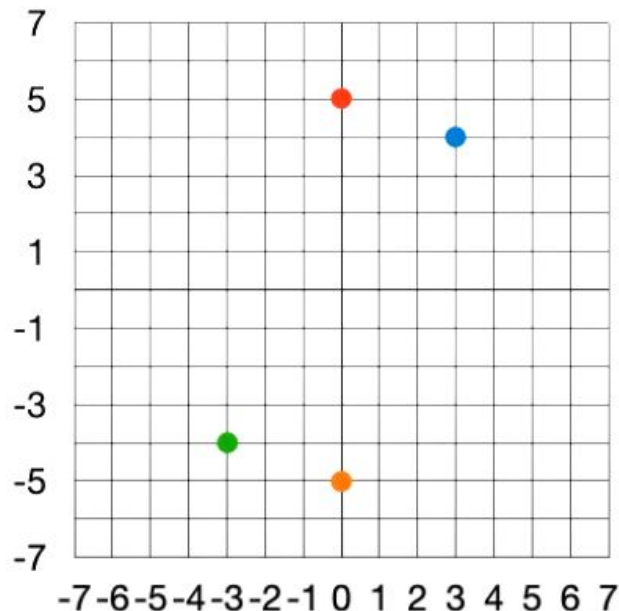
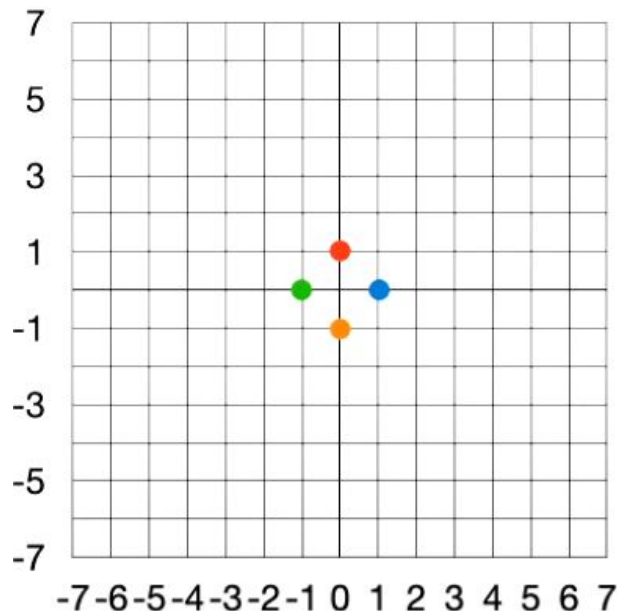
$$(p, q) \quad (3p+0q, 4p+5q)$$

$$(1, 0) \quad (3, 4)$$

$$(0, 1) \quad (0, 5)$$

$$(-1, 0) \quad (-3, -4)$$

$$(0, -1) \quad (0, -5)$$

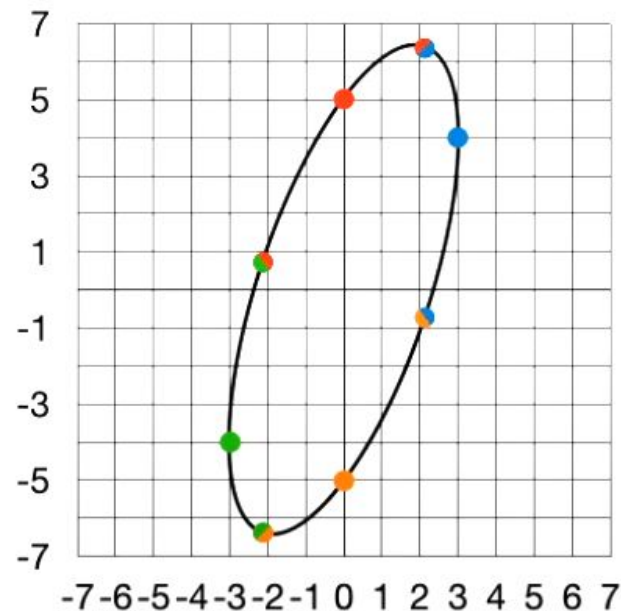
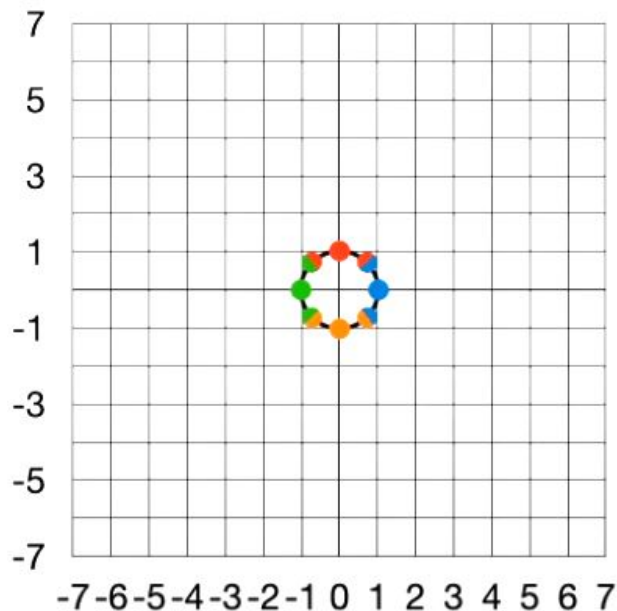


Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Γραμμικός μετασχηματισμός

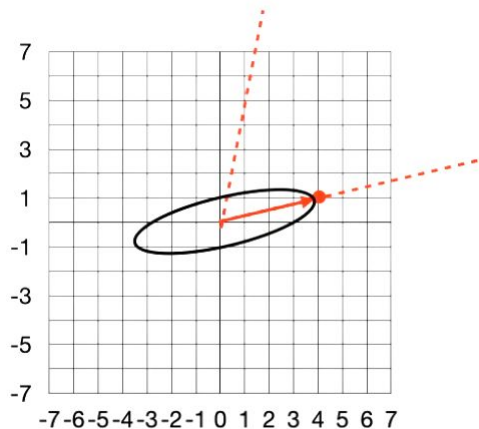
$$A = \begin{bmatrix} 3 & 0 \\ 4 & 5 \end{bmatrix}$$

(p,q)	$(3p+0q, 4p+5q)$
$(1,0)$	$(3, 4)$
$(0,1)$	$(0, 5)$
$(-1,0)$	$(-3, -4)$
$(0,-1)$	$(0, -5)$

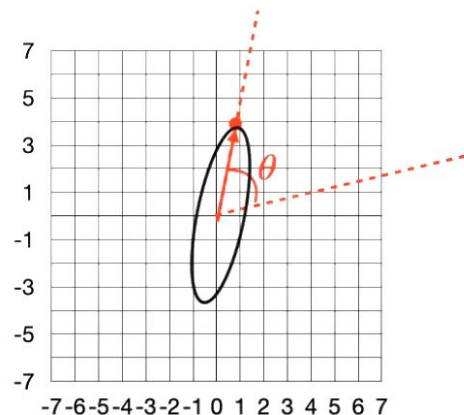


Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Πίνακας περιστροφής



$$\begin{bmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{bmatrix}$$

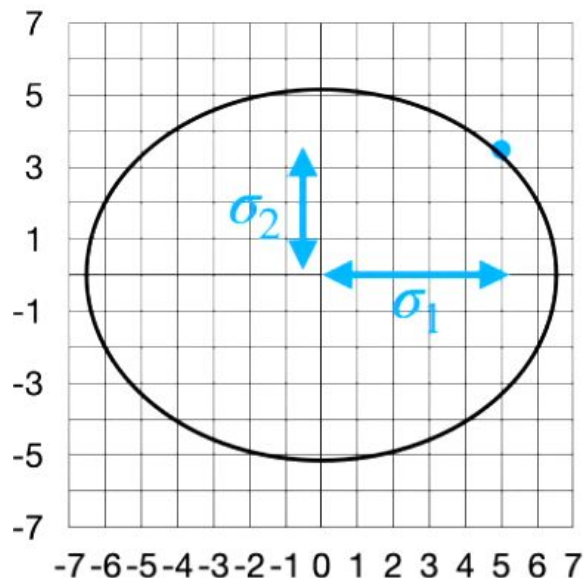


$$\begin{bmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{bmatrix}$$

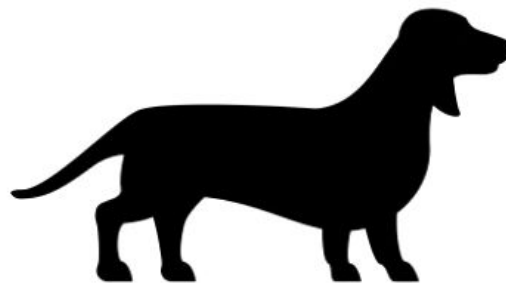


Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

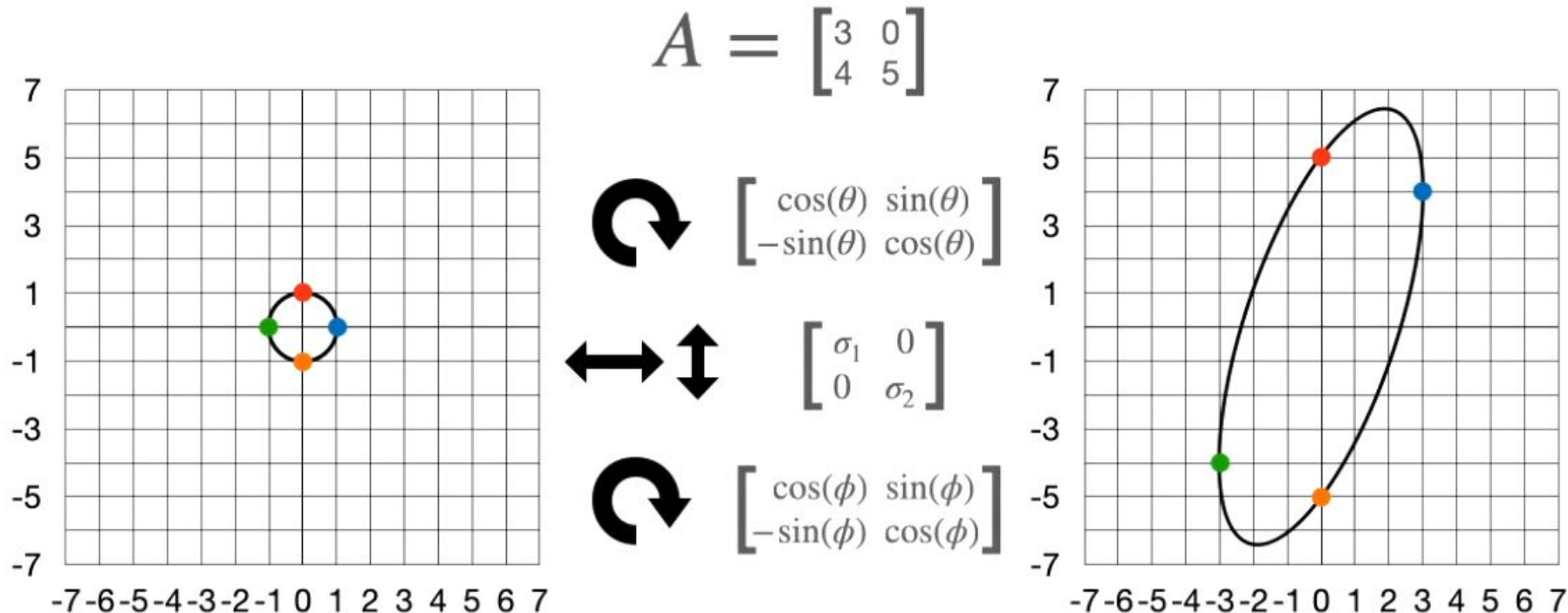
Πίνακας κλιμάκωσης



$$\begin{bmatrix} \sigma_1 & 0 \\ 0 & \sigma_2 \end{bmatrix}$$



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

 **WolframAlpha** computational intelligence.

singular value decomposition $\begin{bmatrix} 3 & 0 \\ 4 & 5 \end{bmatrix}$

Extended Keyboard Upload

Input:

singular value decomposition $\begin{pmatrix} 3 & 0 \\ 4 & 5 \end{pmatrix}$

Result:

$M = U \cdot \Sigma \cdot V^\dagger$

where

$M = \begin{pmatrix} 3 & 0 \\ 4 & 5 \end{pmatrix}$

$U = \begin{pmatrix} \frac{1}{\sqrt{10}} & -\frac{3}{\sqrt{10}} \\ \frac{3}{\sqrt{10}} & \frac{1}{\sqrt{10}} \end{pmatrix}$

$\Sigma = \begin{pmatrix} 3\sqrt{5} & 0 \\ 0 & \sqrt{5} \end{pmatrix}$

$V = \begin{pmatrix} \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{pmatrix}$

$$A = U \Sigma V^\dagger$$

$$\begin{bmatrix} 3 & 0 \\ 4 & 5 \end{bmatrix} = \begin{bmatrix} \cos(\theta) & \sin(\theta) \\ -\sin(\theta) & \cos(\theta) \end{bmatrix} \begin{bmatrix} \sigma_1 & 0 \\ 0 & \sigma_2 \end{bmatrix} \begin{bmatrix} \cos(\phi) & \sin(\phi) \\ -\sin(\phi) & \cos(\phi) \end{bmatrix}$$

```
from numpy.linalg import svd
A = np.array([[3,0],[4,5]])
svd(A)

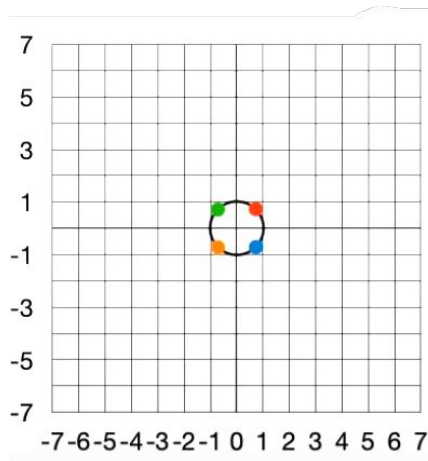
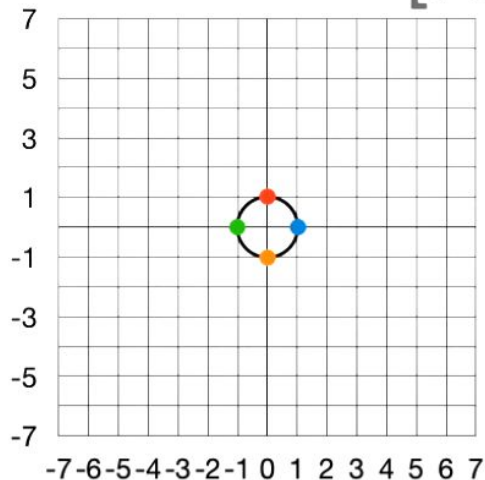
(array([[ -0.31622777, -0.9486833 ],
        [ -0.9486833 ,  0.31622777 ]]),
 array([6.70820393, 2.23606798]),
 array([[ -0.70710678, -0.70710678 ],
        [ -0.70710678,  0.70710678 ]]))
```

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

$$A = U\Sigma V^\dagger$$

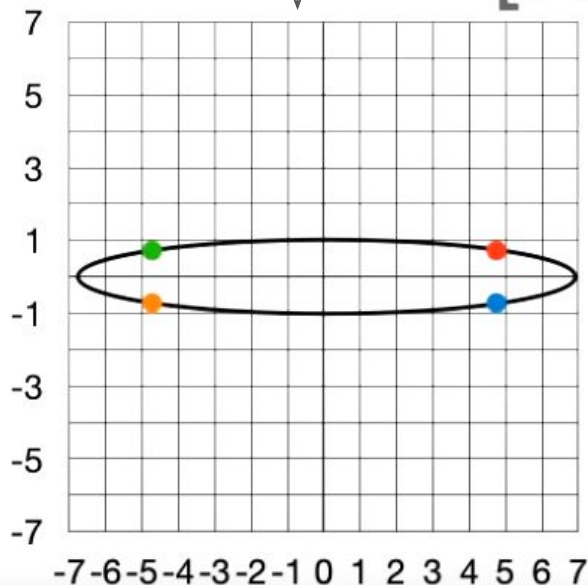
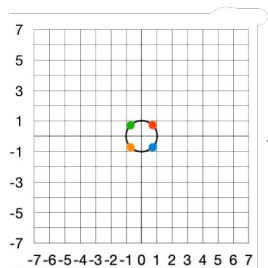
$$\begin{bmatrix} 3 & 0 \\ 4 & 5 \end{bmatrix} = \begin{bmatrix} 0.316 & -0.949 \\ 0.949 & 0.316 \end{bmatrix} \begin{bmatrix} 6.708 & 0 \\ 0 & 2.236 \end{bmatrix} \begin{bmatrix} 0.7071 & 0.7071 \\ -0.7071 & 0.7071 \end{bmatrix}$$

$$\begin{bmatrix} 1/\sqrt{10} & -3/\sqrt{10} \\ 3/\sqrt{10} & 1/\sqrt{10} \end{bmatrix} \begin{bmatrix} 3\sqrt{5} & 0 \\ 0 & \sqrt{5} \end{bmatrix} \begin{bmatrix} 1/\sqrt{2} & 1/\sqrt{2} \\ -1/\sqrt{2} & 1/\sqrt{2} \end{bmatrix}$$



Rotation of $\theta = -\frac{\pi}{4} = -45^\circ$

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



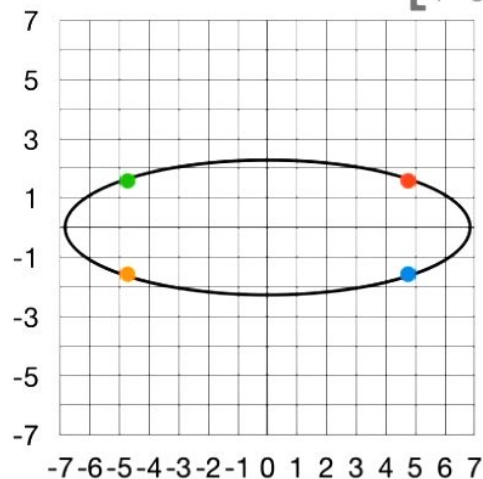
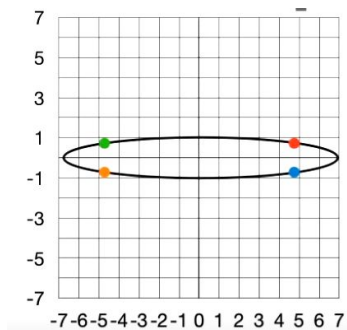
$$A = U\Sigma V^\dagger$$

$$\begin{bmatrix} 3 & 0 \\ 4 & 5 \end{bmatrix} = \begin{bmatrix} 0.316 & -0.949 \\ 0.949 & 0.316 \end{bmatrix} \begin{bmatrix} 6.708 & 0 \\ 0 & 2.236 \end{bmatrix} \begin{bmatrix} 0.7071 & 0.7071 \\ -0.7071 & 0.7071 \end{bmatrix}$$

$$\begin{bmatrix} 1/\sqrt{10} & -3/\sqrt{10} \\ 3/\sqrt{10} & 1/\sqrt{10} \end{bmatrix} \begin{bmatrix} 3\sqrt{5} & 0 \\ 0 & \sqrt{5} \end{bmatrix} \begin{bmatrix} 1/\sqrt{2} & 1/\sqrt{2} \\ -1/\sqrt{2} & 1/\sqrt{2} \end{bmatrix}$$

Horizontal scaling by $3\sqrt{5}$

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

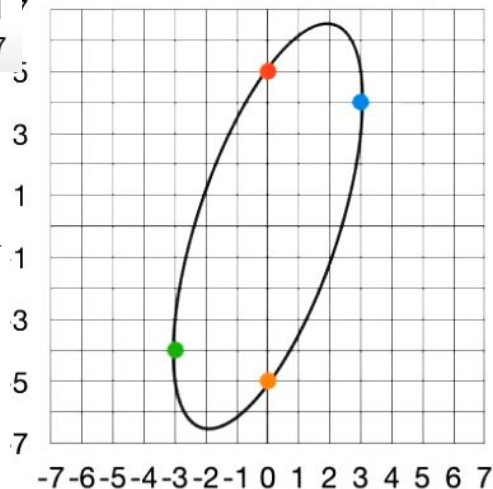
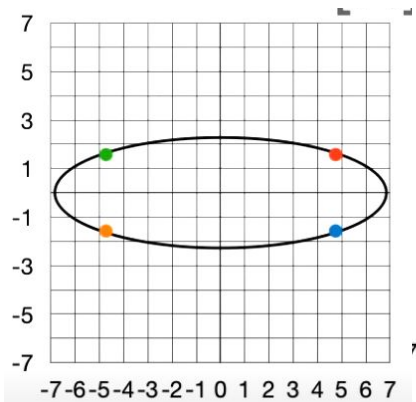


$$\begin{bmatrix} 3 & 0 \\ 4 & 5 \end{bmatrix} = \begin{bmatrix} 0.316 & -0.949 \\ 0.949 & 0.316 \end{bmatrix} \begin{bmatrix} 6.708 & 0 \\ 0 & 2.236 \end{bmatrix} \begin{bmatrix} 0.7071 & 0.7071 \\ -0.7071 & 0.7071 \end{bmatrix}$$

$$\begin{bmatrix} 1/\sqrt{10} & -3/\sqrt{10} \\ 3/\sqrt{10} & 1/\sqrt{10} \end{bmatrix} \begin{bmatrix} 3\sqrt{5} & 0 \\ 0 & \sqrt{5} \end{bmatrix} \begin{bmatrix} 1/\sqrt{2} & 1/\sqrt{2} \\ -1/\sqrt{2} & 1/\sqrt{2} \end{bmatrix}$$

Vertical scaling by $\sqrt{5}$

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

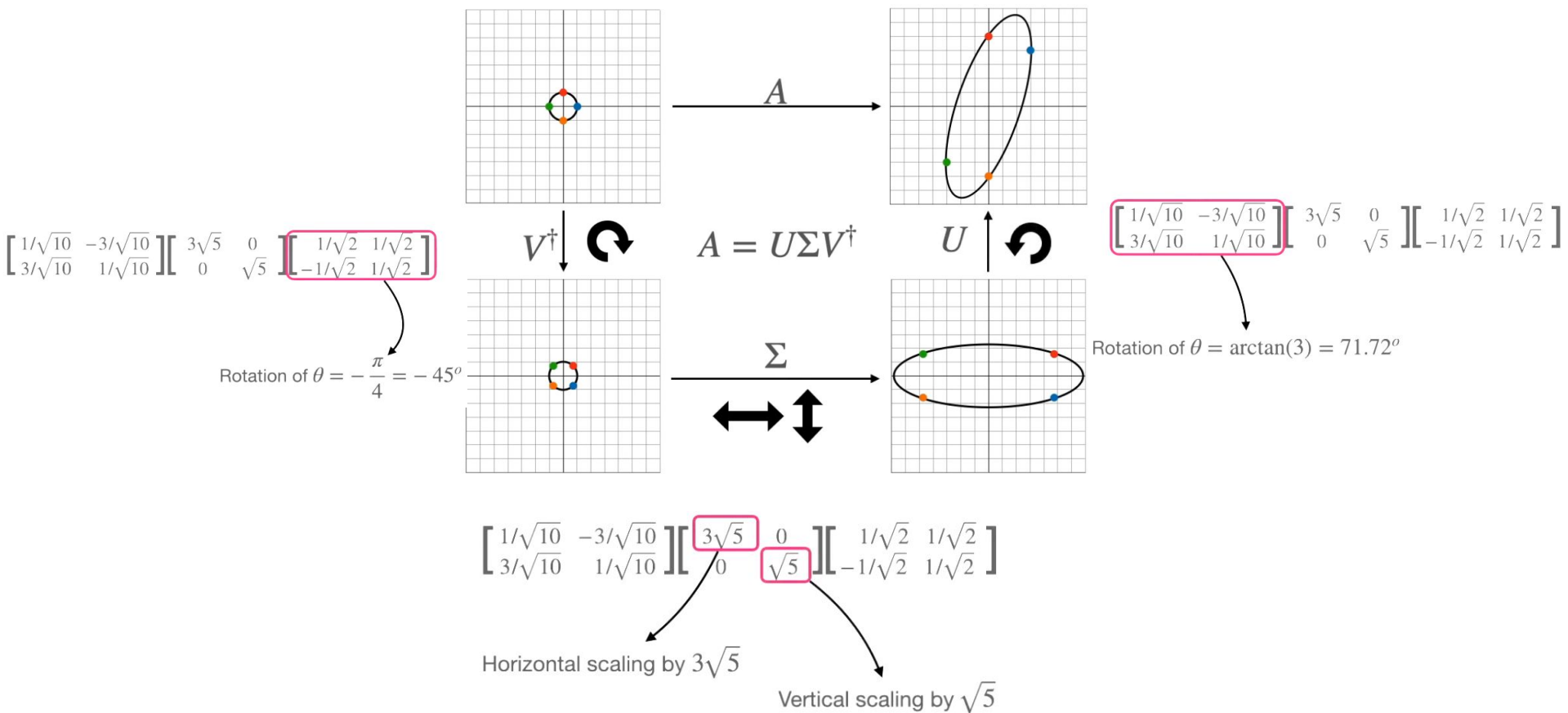


$$\begin{bmatrix} 3 & 0 \\ 4 & 5 \end{bmatrix} = \begin{bmatrix} 0.316 & -0.949 \\ 0.949 & 0.316 \end{bmatrix} \begin{bmatrix} 6.708 & 0 \\ 0 & 2.236 \end{bmatrix} \begin{bmatrix} 0.7071 & 0.7071 \\ -0.7071 & 0.7071 \end{bmatrix}$$

$$\begin{bmatrix} 1/\sqrt{10} & -3/\sqrt{10} \\ 3/\sqrt{10} & 1/\sqrt{10} \end{bmatrix} \begin{bmatrix} 3\sqrt{5} & 0 \\ 0 & \sqrt{5} \end{bmatrix} \begin{bmatrix} 1/\sqrt{2} & 1/\sqrt{2} \\ -1/\sqrt{2} & 1/\sqrt{2} \end{bmatrix}$$

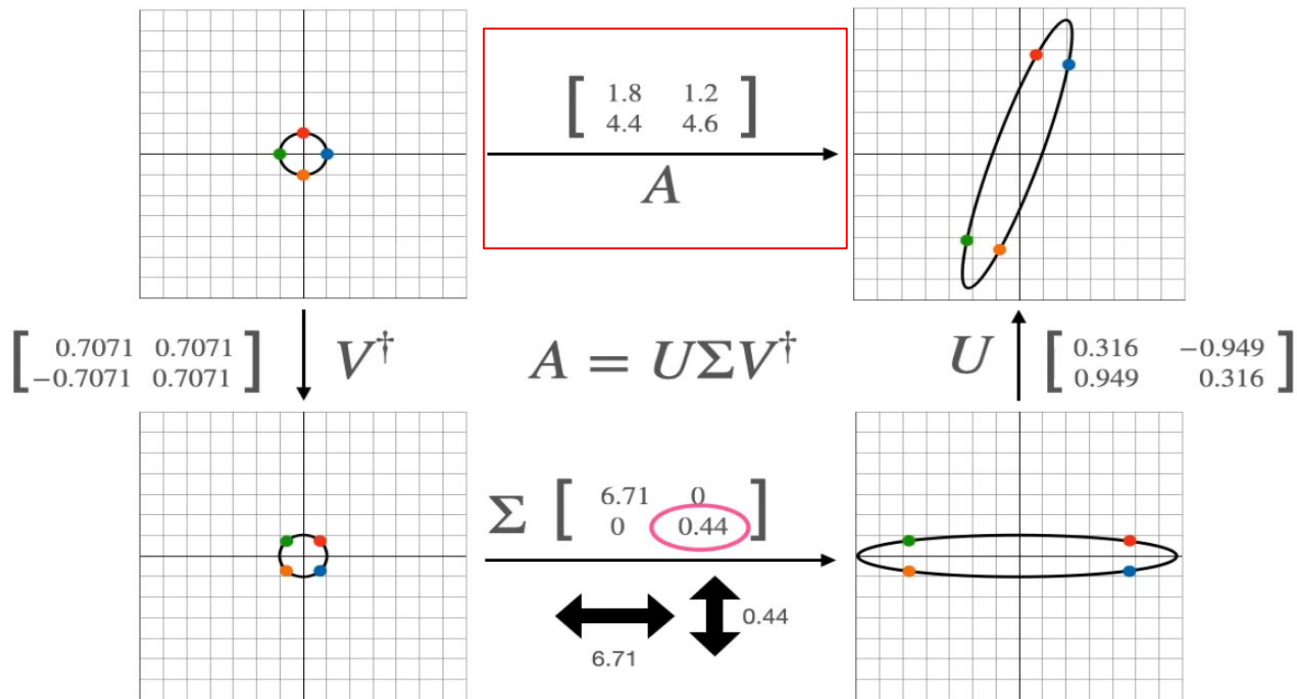
Rotation of $\theta = \arctan(3) = 71.72^\circ$

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

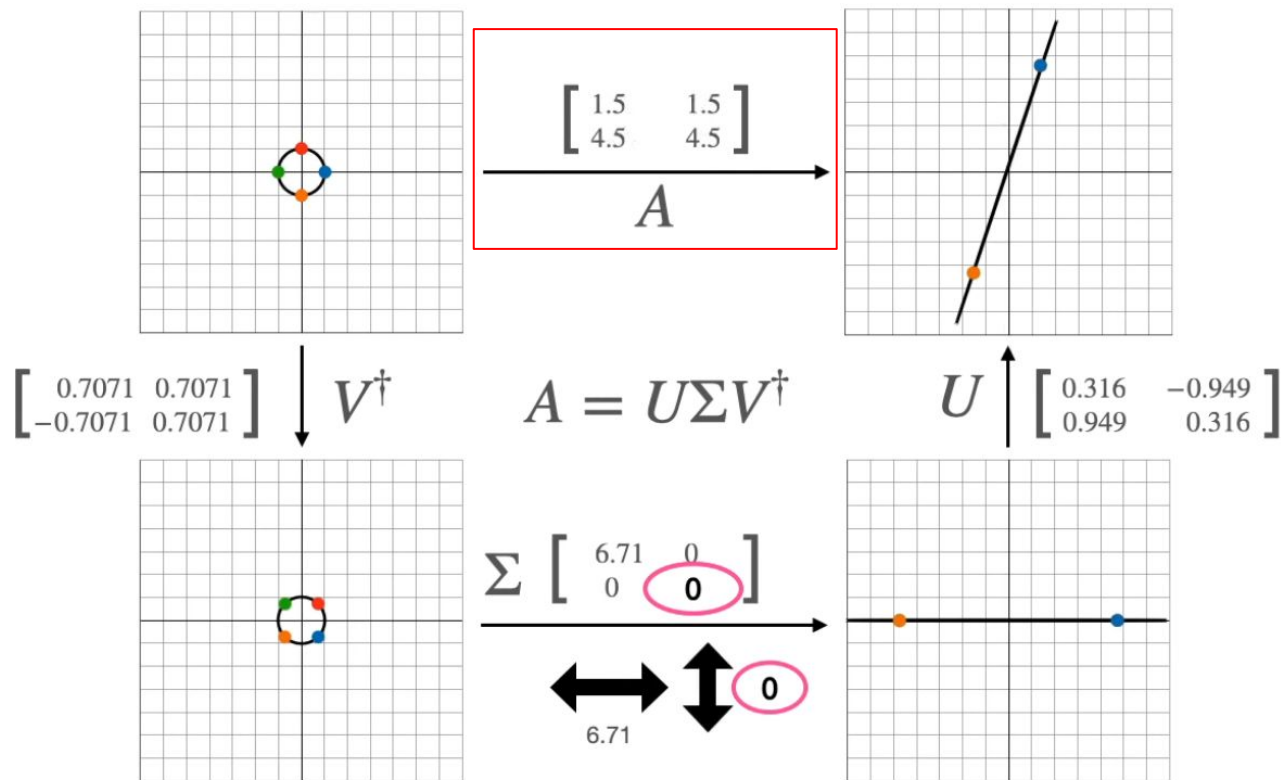


Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

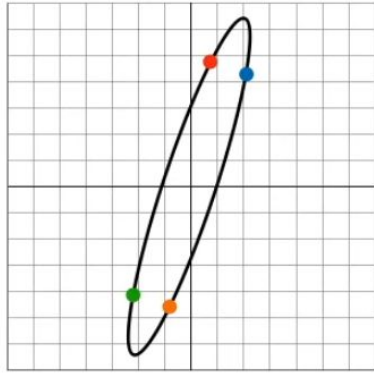
Μείωση διαστάσεων



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

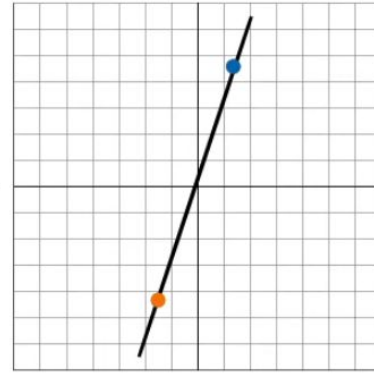


Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



$$\begin{bmatrix} 1.8 & 1.2 \\ 4.4 & 4.6 \end{bmatrix}$$

Rank 2

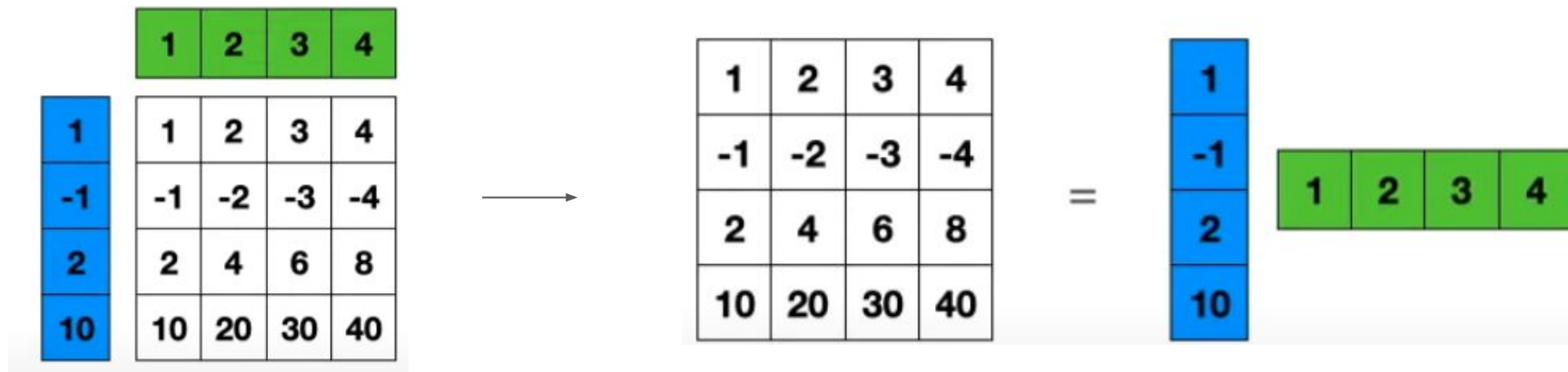


$$\begin{bmatrix} 1.5 & 1.5 \\ 4.5 & 4.5 \end{bmatrix}$$

Rank 1

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Πίνακας με rank=1

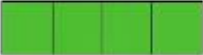


Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

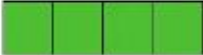
Πίνακας με $\text{rank} > 1$: προσέγγιση όπως λειτουργεί για $\text{rank} = 1$

3	1	4	1
5	9	2	6
5	3	5	8
9	7	9	3

=



+

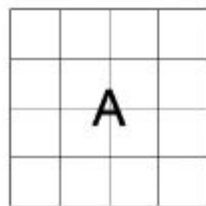


+

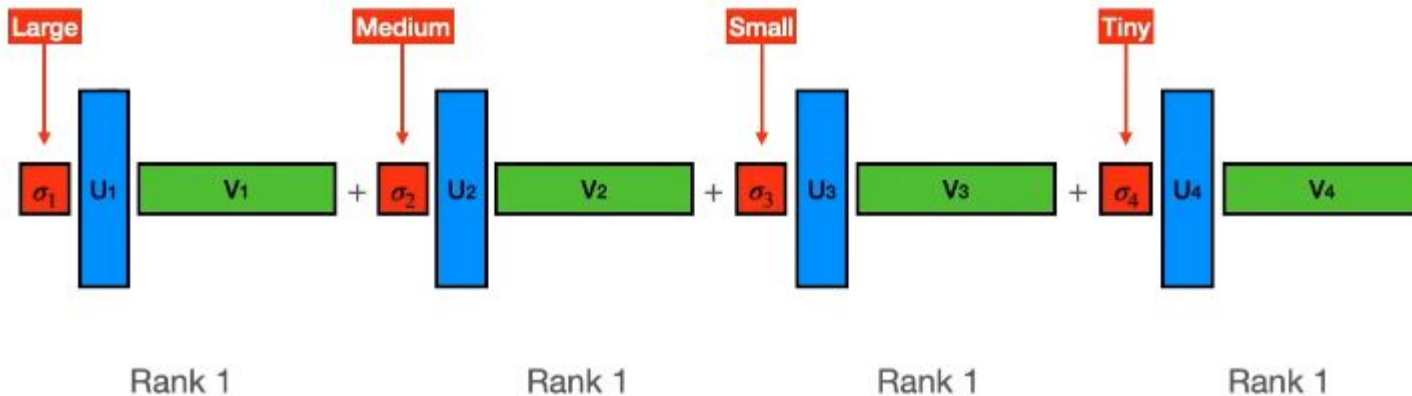
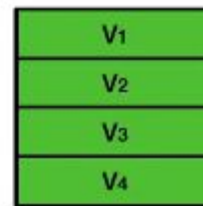
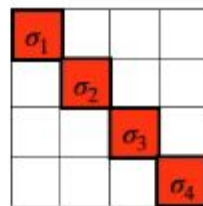
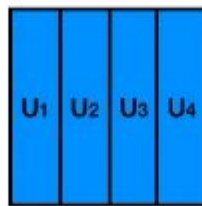
...

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

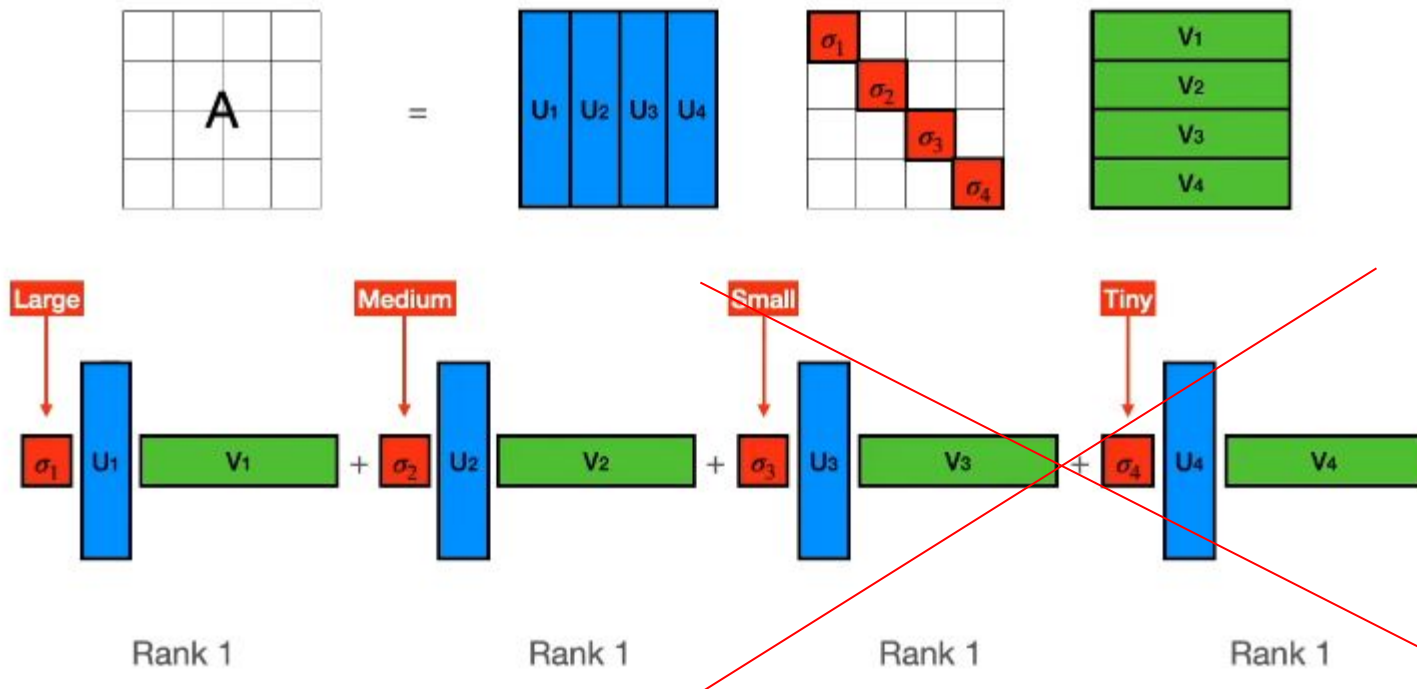
$$A = U \Sigma V^T$$



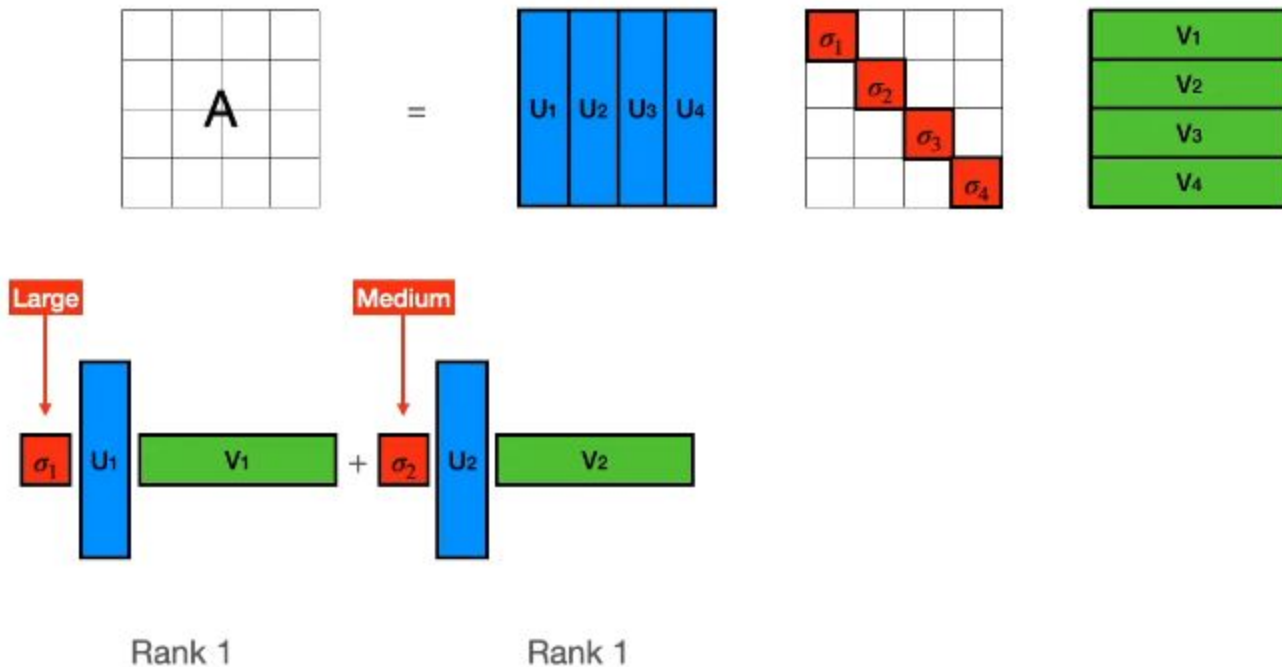
=



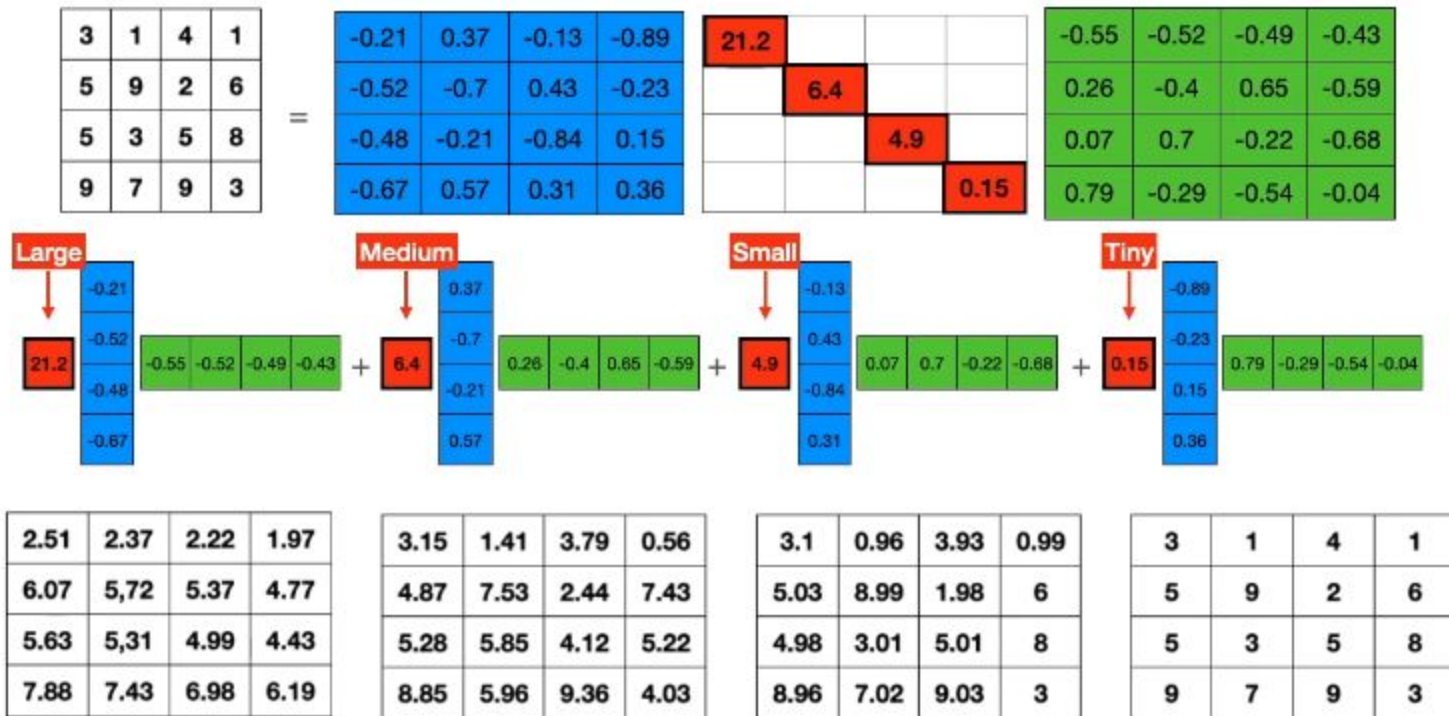
Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



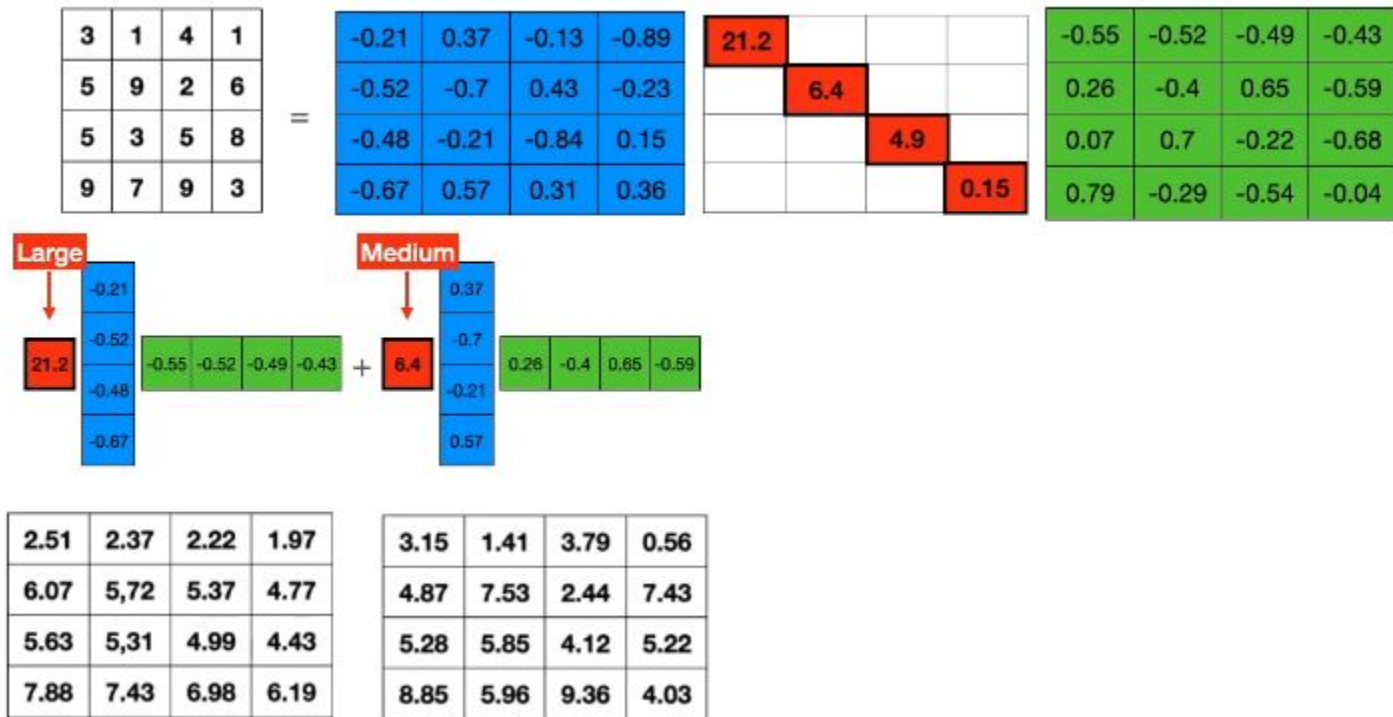
Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

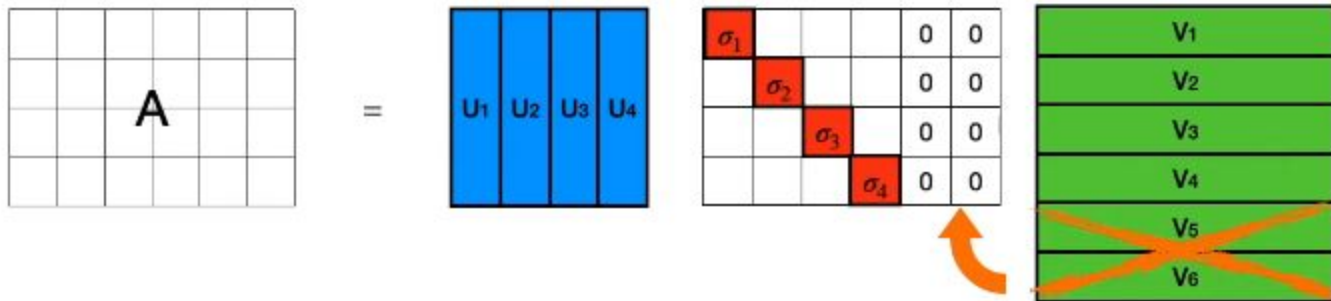


Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Μη τετραγωνικός πίνακας;



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Ανάλυση διανυσμάτων πάνω σε ορθογώνιους άξονες

- Αν ένας πίνακας δεν είναι τετραγωνικός, η ανάλυση ιδιοτιμών δεν ορίζεται.
- Αντ' αυτού, χρησιμοποιούμε την ανάλυση σε ιδιάζουσες τιμές

Κάθε πίνακας A ($m \times n$) μπορεί να γραφεί ως $A = U\Sigma V^T$

Σημαντική εφαρμογή: Μπορούμε να γενικεύσουμε, εν μέρει, την αντιστροφή
ενός μη τετραγωνικού πίνακα

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

$$A = U\Sigma V^T$$

Όπου **U**: Ορθογώνιος πίνακας $m \times m$ → Περιστροφή (Rotation)

Σ : Διαγώνιος πίνακας διάστασης $m \times n$ → Έκταση (Stretch)

(όχι απαραίτητα τετραγωνικός)

V: Ορθογώνιος πίνακας $n \times n$ → Περιστροφή (Rotation)

$$A = \begin{bmatrix} u_1 \\ u_2 \end{bmatrix} \begin{bmatrix} \sigma_1 & 0 \\ 0 & \sigma_2 \end{bmatrix} \begin{bmatrix} v_1^T \\ v_2^T \end{bmatrix}$$

The diagram illustrates the decomposition of matrix A into three components. Arrows point from the three matrices in the equation to their respective labels below:

- Left singular vector (pointing to $\begin{bmatrix} u_1 \\ u_2 \end{bmatrix}$)
- singular values (pointing to $\begin{bmatrix} \sigma_1 & 0 \\ 0 & \sigma_2 \end{bmatrix}$)
- Right singular vector (pointing to $\begin{bmatrix} v_1^T \\ v_2^T \end{bmatrix}$)

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

Τι είναι τα U, Σ, V^T

Ίδιος πίνακας

$$A = V \Lambda V^{-1}$$

Διαφορετικοί Πίνακες

$$A = U \Sigma V^T$$

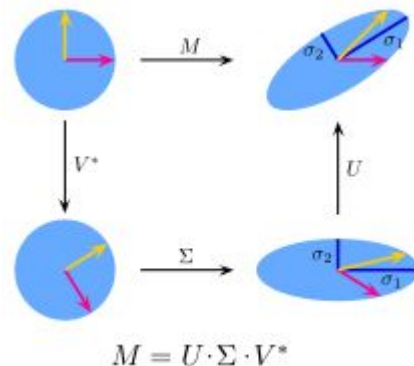
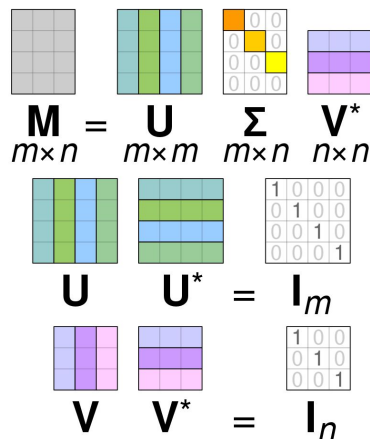
Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

$$A = U\Sigma V^T$$

$$\boxed{A^T A} = (V\Sigma^T U^T)U\Sigma V^T = V\Sigma^T (U^T U)\Sigma V^T = V(\Sigma^T \Sigma)V^T$$

$$\overline{AA^T} = U\Sigma V^T (V\Sigma^T U^T) = U\Sigma (V^T V)\Sigma^T U^T = U(\Sigma\Sigma^T)U^T$$

λ , για $A^T A$
 σ^2 , για A



Βήμα προς βήμα με python

Deep Learning Book Series · 2.8 Singular Value Decomposition

A : πίνακας που μπορούμε να τον “δούμε” ως γραμμικό μετασχηματισμό που μπορεί να αναλυθεί σε 3 υπό- μετασχηματισμούς (αντιστοιχούν σε 3 πίνακες) :

1. Περιστροφή,
2. Έκταση,
3. Περιστροφή.

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



www.github.com/luisguiserrano/singular_value_decomposition

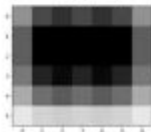
0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

```
[[-0.36  0.   -0.73 -0.05  0.56  0.13]
 [-0.54  0.35  0.27 -0.08 -0.16  0.69]
 [-0.54  0.35  0.27 -0.08  0.16 -0.69]
 [-0.45 -0.35 -0.27  0.52 -0.56 -0.13]
 [-0.28 -0.71  0.18 -0.62  0.   0. ]
 [-0.08 -0.35  0.46  0.57  0.56  0.13]]
```

```
[ [4.74 0.  0.  0.  0.  0.  0. ]
 [0.  1.41 0.  0.  0.  0.  0. ]
 [0.  0.  1.41 0.  0.  0.  0. ]
 [0.  0.  0.  0.73 0.  0.  0. ]
 [0.  0.  0.  0.  0.  0.  0. ]
 [0.  0.  0.  0.  0.  0.  0. ]]
```

```
[[-0.23 -0.4  -0.46 -0.4  -0.46 -0.4  -0.23]
 [ 0.5  0.25 -0.25 -0.5  -0.25  0.25  0.5 ]
 [ 0.39 -0.32 -0.19  0.65 -0.19 -0.32  0.39]
 [-0.22  0.42 -0.44  0.42 -0.44  0.42 -0.22]
 [ 0.56 -0.43  0.03  0.  -0.03  0.43 -0.56]
 [-0.42 -0.55 -0.16  0.  0.16  0.55  0.42]
 [-0.12 -0.11  0.69  0.  -0.69  0.11  0.12]]
```

4.74



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

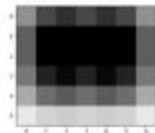


www.github.com/luisguiserrano/singular_value_decomposition

0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

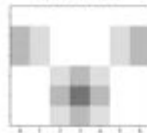
```
[[-0.36 -0. -0.73 -0.05 0.56 0.13]
 [-0.54 0.35 0.27 -0.08 -0.16 0.69]
 [-0.54 0.35 0.27 -0.08 0.16 -0.69]
 [-0.45 -0.35 -0.27 0.52 -0.56 -0.13]
 [-0.28 -0.71 0.18 -0.62 -0. -0. ]
 [-0.08 -0.35 0.46 0.57 0.56 0.13]]
```

4.74



```
[[4.74 0. 0. 0. 0. 0. 0. ]
 [0. 1.41 0. 0. 0. 0. 0. ]
 [0. 0. 1.41 0. 0. 0. 0. ]
 [0. 0. 0. 0.73 0. 0. 0. ]
 [0. 0. 0. 0. 0. 0. 0. ]
 [0. 0. 0. 0. 0. 0. 0. ]]
```

1.41



```
[[-0.23 -0.4 -0.46 -0.4 -0.46 -0.4 -0.23]
 [0.5 0.25 -0.25 -0.5 -0.25 0.25 0.5 ]
 [0.39 -0.32 -0.19 0.65 -0.19 -0.32 0.39]
 [-0.22 0.42 -0.44 0.42 -0.44 0.42 -0.22]
 [0.56 -0.43 0.03 0. -0.03 0.43 -0.56]
 [-0.42 -0.55 -0.16 0. 0.16 0.55 0.42]
 [-0.12 -0.11 0.69 -0. -0.69 0.11 0.12]]
```

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)

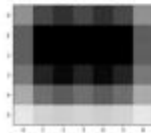


www.github.com/luisguiserrano/singular_value_decomposition

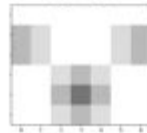
0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

```
[[-0.36 -0. -0.73 -0.05 0.56 0.13]
 [-0.54 0.35 0.27 -0.08 -0.16 0.69]
 [-0.54 0.35 0.27 -0.08 0.16 -0.69]
 [-0.45 -0.35 -0.27 0.52 -0.56 -0.13]
 [-0.28 -0.71 0.18 -0.62 -0. -0. ]
 [-0.08 -0.35 0.46 0.57 0.56 0.13]]
```

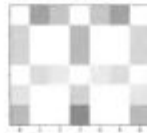
4.74



1.41



1.41



```
[[4.74 0. 0. 0. 0. 0. 0. ]
 [0. 1.41 0. 0. 0. 0. 0. ]
 [0. 0. 1.41 0. 0. 0. 0. ]
 [0. 0. 0. 0.73 0. 0. 0. ]
 [0. 0. 0. 0. 0. 0. 0. ]
 [0. 0. 0. 0. 0. 0. 0. ]]
```

```
[[-0.23 -0.4 -0.46 -0.4 -0.46 -0.4 -0.23]
 [ 0.5 0.25 -0.25 -0.5 -0.25 0.25 0.5 ]
 [ 0.39 -0.32 -0.19 0.65 -0.19 -0.32 0.39]
 [-0.22 0.42 -0.44 0.42 -0.44 0.42 -0.22]
 [ 0.56 -0.43 0.03 0. -0.03 0.43 -0.56]
 [-0.42 -0.55 -0.16 0. 0.16 0.55 0.42]
 [-0.12 -0.11 0.69 -0. -0.69 0.11 0.12]]
```

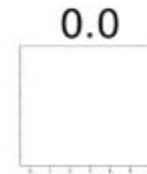
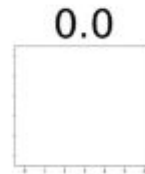
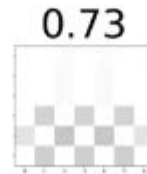
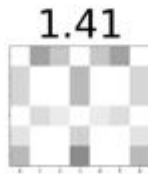
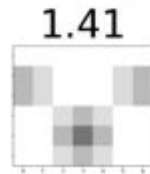
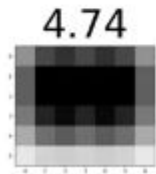
Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



www.github.com/luisguiserrano/singular_value_decomposition

0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

```
[[-0.36 -0.   -0.73 -0.05  0.56  0.13]
 [-0.54  0.35  0.27 -0.08 -0.16  0.69]
 [-0.54  0.35  0.27 -0.08  0.16 -0.69]
 [-0.45 -0.35 -0.27  0.52 -0.56 -0.13]
 [-0.28 -0.71  0.18 -0.62 -0.   -0. ]
 [-0.08 -0.35  0.46  0.57  0.56  0.13]]
```



```
[[4.74 0.   0.   0.   0.   0.   0. ]
 [0.  1.41 0.   0.   0.   0.   0. ]
 [0.   0.  1.41 0.   0.   0.   0. ]
 [0.   0.   0.  0.73 0.   0.   0. ]
 [0.   0.   0.   0.   0.   0.   0. ]
 [0.   0.   0.   0.   0.   0.   0. ]]
```

```
[[-0.23 -0.4  -0.46 -0.4  -0.46 -0.4  -0.23]
 [ 0.5   0.25 -0.25 -0.5  -0.25  0.25  0.5 ]
 [ 0.39 -0.32 -0.19  0.65 -0.19 -0.32  0.39]
 [-0.22  0.42 -0.44  0.42 -0.44  0.42 -0.22]
 [ 0.56 -0.43  0.03  0.   -0.03  0.43 -0.56]
 [-0.42 -0.55 -0.16  0.   0.16  0.55  0.42]
 [-0.12 -0.11  0.69 -0.   -0.69  0.11  0.12]]
```

Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



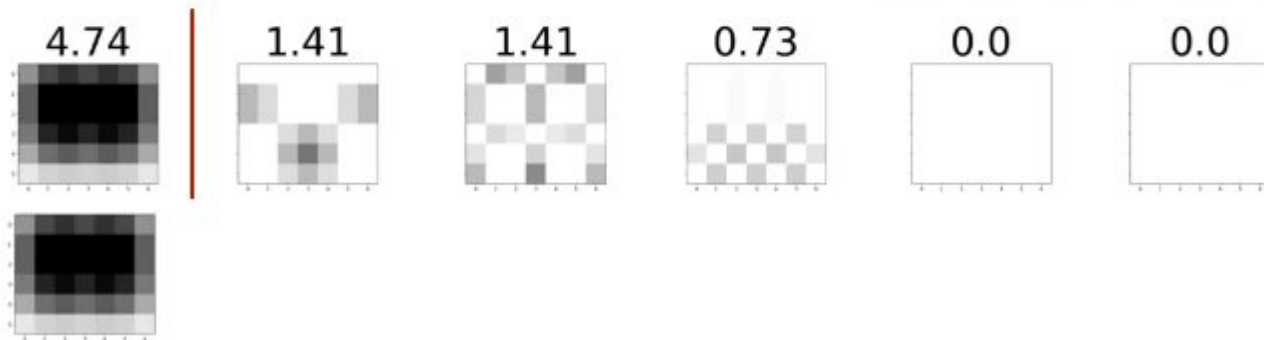
www.github.com/luisguiserrano/singular_value_decomposition

0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

```
[[-0.36  0.   -0.73 -0.05  0.56  0.13]
 [-0.54  0.35  0.27 -0.08 -0.16  0.69]
 [-0.54  0.35  0.27 -0.08  0.16 -0.69]
 [-0.45 -0.35 -0.27  0.52 -0.56 -0.13]
 [-0.28 -0.71  0.18 -0.62  0.   0.   ]
 [-0.08 -0.35  0.46  0.57  0.56  0.13]]
```

```
[[4.74 0.   0.   0.   0.   0.   0. ]
 [0.  1.41 0.   0.   0.   0.   0. ]
 [0.  0.  1.41 0.   0.   0.   0. ]
 [0.  0.  0.  0.73 0.   0.   0. ]
 [0.  0.  0.  0.   0.   0.   0. ]
 [0.  0.  0.  0.   0.   0.   0. ]]
```

```
[[-0.23 -0.4  -0.46 -0.4  -0.46 -0.4  -0.23]
 [ 0.5  0.25 -0.25 -0.5  -0.25  0.25  0.5 ]
 [ 0.39 -0.32 -0.19  0.65 -0.19 -0.32  0.39]
 [-0.22  0.42 -0.44  0.42 -0.44  0.42 -0.22]
 [ 0.56 -0.43  0.03  0.   -0.03  0.43 -0.56]
 [-0.42 -0.55 -0.16  0.   0.16  0.55  0.42]
 [-0.12 -0.11  0.69  0.   -0.69  0.11  0.12]]
```



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



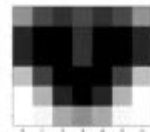
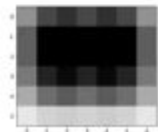
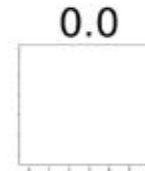
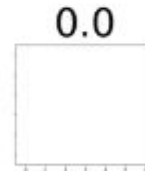
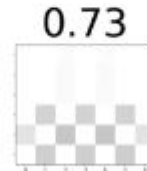
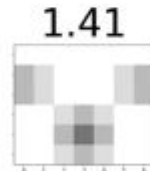
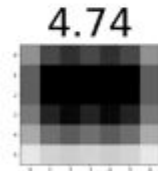
www.github.com/luisguiserrano/singular_value_decomposition

0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

```
[[-0.36  0.   -0.73 -0.05  0.56  0.13]
 [-0.54  0.35  0.27 -0.08 -0.16  0.69]
 [-0.54  0.35  0.27 -0.08  0.16 -0.69]
 [-0.45 -0.35 -0.27  0.52 -0.56 -0.13]
 [-0.28 -0.71  0.18 -0.62 -0.   -0. ]
 [-0.08 -0.35  0.46  0.57  0.56  0.13]]
```

```
[[4.74 0.   0.   0.   0.   0.   0. ]
 [0.   1.41 0.   0.   0.   0.   0. ]
 [0.   0.   1.41 0.   0.   0.   0. ]
 [0.   0.   0.   0.73 0.   0.   0. ]
 [0.   0.   0.   0.   0.   0.   0. ]
 [0.   0.   0.   0.   0.   0.   0. ]]
```

```
[[-0.23 -0.4  -0.46 -0.4  -0.46 -0.4  -0.23]
 [ 0.5   0.25 -0.25 -0.5  -0.25  0.25  0.5 ]
 [ 0.39 -0.32 -0.19  0.65 -0.19 -0.32  0.39]
 [-0.22  0.42 -0.44  0.42 -0.44  0.42 -0.22]
 [ 0.56 -0.43  0.03  0.   -0.03  0.43 -0.56]
 [-0.42 -0.55 -0.16  0.   0.16  0.55  0.42]
 [-0.12 -0.11  0.69 -0.   -0.69  0.11  0.12]]
```



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



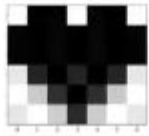
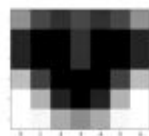
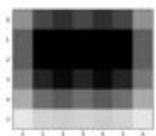
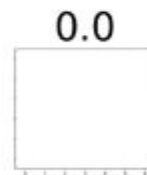
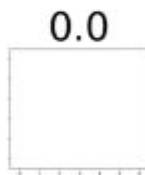
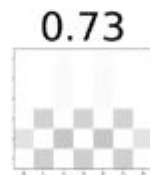
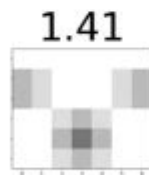
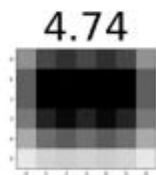
www.github.com/luiguiserrano/singular_value_decomposition

0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

```
[[-0.36 -0.   -0.73 -0.05  0.56  0.13]
 [-0.54  0.35  0.27 -0.08 -0.16  0.69]
 [-0.54  0.35  0.27 -0.08  0.16 -0.69]
 [-0.45 -0.35 -0.27  0.52 -0.56 -0.13]
 [-0.28 -0.71  0.18 -0.62 -0.   -0. ]
 [-0.08 -0.35  0.46  0.57  0.56  0.13]]
```

```
[[4.74 0.   0.   0.   0.   0.   0. ]
 [0.   1.41 0.   0.   0.   0.   0. ]
 [0.   0.   1.41 0.   0.   0.   0. ]
 [0.   0.   0.   0.73 0.   0.   0. ]
 [0.   0.   0.   0.   0.   0.   0. ]
 [0.   0.   0.   0.   0.   0.   0. ]]
```

```
[[-0.23 -0.4  -0.46 -0.4  -0.46 -0.4  -0.23]
 [ 0.5   0.25 -0.25 -0.5  -0.25  0.25  0.5 ]
 [ 0.39 -0.32 -0.19  0.65 -0.19 -0.32  0.39]
 [-0.22  0.42 -0.44  0.42 -0.44  0.42 -0.22]
 [ 0.56 -0.43  0.03  0.   -0.03  0.43 -0.56]
 [-0.42 -0.55 -0.16  0.   0.16  0.55  0.42]
 [-0.12 -0.11  0.69 -0.   -0.69  0.11  0.12]]
```



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



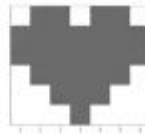
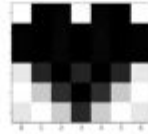
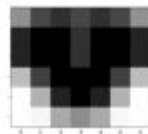
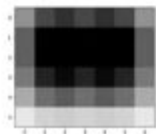
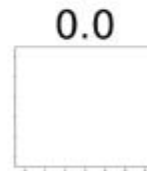
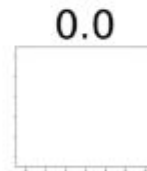
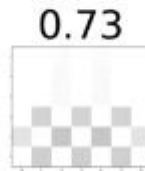
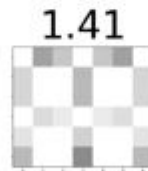
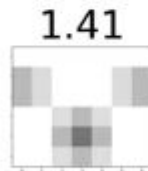
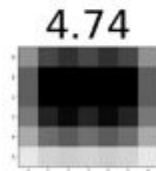
www.github.com/luisguiserrano/singular_value_decomposition

0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

```
[[-0.36 -0. -0.73 -0.05 0.56 0.13]
 [-0.54 0.35 0.27 -0.08 -0.16 0.69]
 [-0.54 0.35 0.27 -0.08 0.16 -0.69]
 [-0.45 -0.35 -0.27 0.52 -0.56 -0.13]
 [-0.28 -0.71 0.18 -0.62 -0. -0. ]
 [-0.08 -0.35 0.46 0.57 0.56 0.13]]
```

```
[[4.74 0. 0. 0. 0. 0. 0. ]
 [0. 1.41 0. 0. 0. 0. 0. ]
 [0. 0. 1.41 0. 0. 0. 0. ]
 [0. 0. 0. 0.73 0. 0. 0. ]
 [0. 0. 0. 0. 0. 0. 0. ]
 [0. 0. 0. 0. 0. 0. 0. ]]
```

```
[[-0.23 -0.4 -0.46 -0.4 -0.46 -0.4 -0.23]
 [ 0.5 0.25 -0.25 -0.5 -0.25 0.25 0.5 ]
 [ 0.39 -0.32 -0.19 0.65 -0.19 -0.32 0.39]
 [-0.22 0.42 -0.44 0.42 -0.44 0.42 -0.22]
 [ 0.56 -0.43 0.03 0. -0.03 0.43 -0.56]
 [-0.42 -0.55 -0.16 0. 0.16 0.55 0.42]
 [-0.12 -0.11 0.69 -0. -0.69 0.11 0.12]]
```



Ανάλυση σε Ιδιάζουσες Τιμές (Singular Value Decomposition)



www.github.com/luisguiserrano/singular_value_decomposition

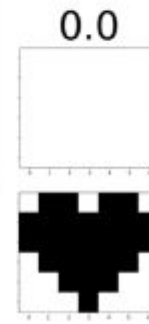
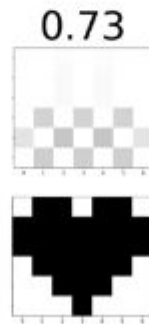
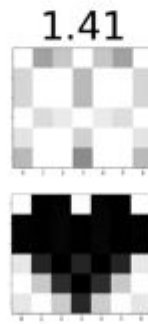
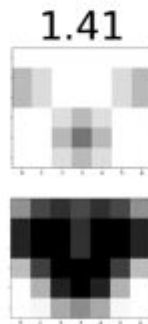
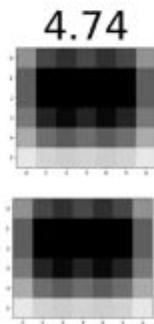
0	1	1	0	1	1	0
1	1	1	1	1	1	1
1	1	1	1	1	1	1
0	1	1	1	1	1	0
0	0	1	1	1	0	0
0	0	0	1	0	0	0

Rank 4

```
[[-0.36 -0.   -0.73 -0.05  0.56  0.13]
 [-0.54  0.35  0.27 -0.08 -0.16  0.69]
 [-0.54  0.35  0.27 -0.08  0.16 -0.69]
 [-0.45 -0.35 -0.27  0.52 -0.56 -0.13]
 [-0.28 -0.71  0.18 -0.62 -0.   -0. ]
 [-0.08 -0.35  0.46  0.57  0.56  0.13]]
```

```
[[4.74 0.   0.   0.   0.   0.   0. ]
 [0.  1.41 0.   0.   0.   0.   0. ]
 [0.   0.  1.41 0.   0.   0.   0. ]
 [0.   0.   0.  0.73 0.   0.   0. ]
 [0.   0.   0.   0.   0.   0.   0. ]
 [0.   0.   0.   0.   0.   0.   0. ]]
```

```
[[-0.23 -0.4 -0.46 -0.4 -0.46 -0.4 -0.23]
 [ 0.5  0.25 -0.25 -0.5 -0.25  0.25  0.5 ]
 [ 0.39 -0.32 -0.19  0.65 -0.19 -0.32  0.39]
 [-0.22  0.42 -0.44  0.42 -0.44  0.42 -0.22]
 [ 0.56 -0.43  0.03  0.  -0.03  0.43 -0.56]
 [-0.42 -0.55 -0.16  0.  0.16  0.55  0.42]
 [-0.12 -0.11  0.69 -0.  -0.69  0.11  0.12]]
```



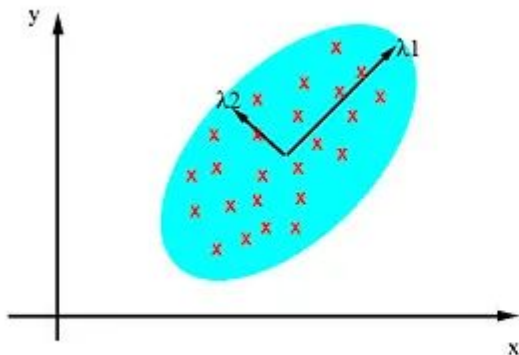
PCA vs LDA (Linear Discriminant Analysis)

Το PCA εστιάζει κυρίως στη μεγαλύτερη απόκλιση μεταξύ όλων των μεταβλητών.

→ Το LDA μας ενδιαφέρει να μεγιστοποιήσουμε τη διαχωριστικότητα μεταξύ όλων των γνωστών κατηγοριών.

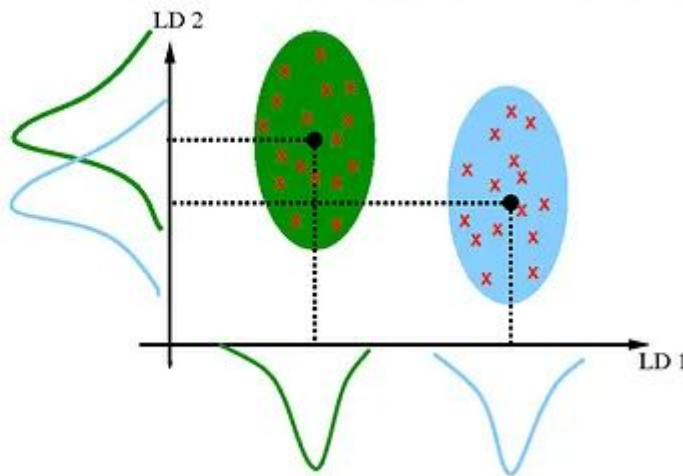
→ Το LDA προβάλλει τα δεδομένα με τρόπο που μεγιστοποιεί τον διαχωρισμό δύο κατηγοριών.

PCA: component axes that maximize the variance



UNSUPERVISED

LDA: maximizing the component axes for class-separation



SUPERVISED

LDA (Linear Discriminant Analysis)

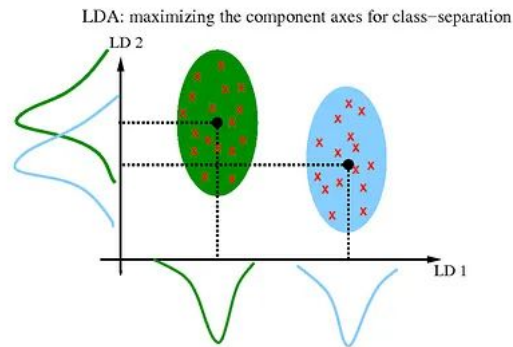
Δύο κριτήρια:

- Μεγιστοποιήστε την απόσταση μεταξύ των μέσων των κατηγοριών
- Ελαχιστοποιήστε τη διασπορά σε κάθε κατηγορία

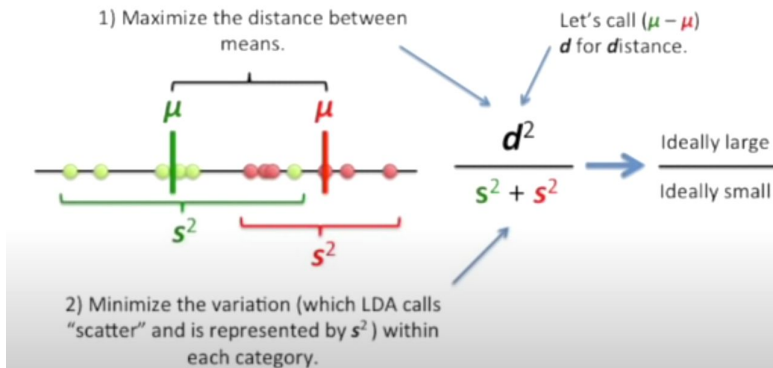
Βήμα 1: Υπολογίζεται ο μέσος όρος και η τυπική απόκλιση κάθε χαρακτηριστικού εντός της κλάσης

Βήμα 2: Υπολογίζονται οι πίνακες διασποράς εντός και μεταξύ των κλάσεων και χρησιμοποιούνται για τον υπολογισμό των ιδιοδιανυσμάτων και των ιδιοτιμών.

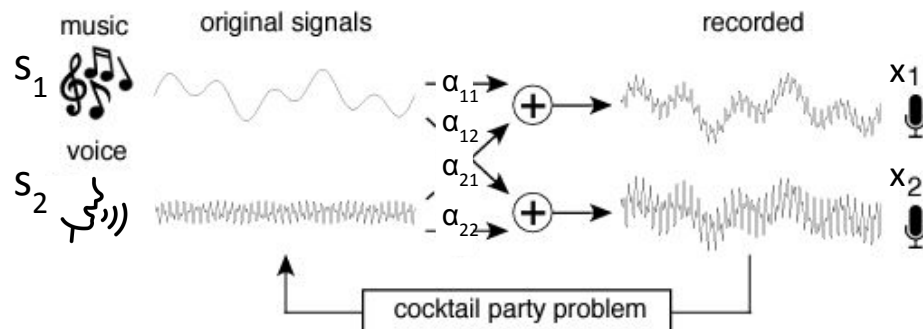
Βήμα 3 Κατασκευάζεται χώρος χαμηλότερων διαστάσεων που μεγιστοποιεί τη διακύμανση μεταξύ κλάσεων και ελαχιστοποιεί τη διακύμανση εντός κλάσης



The new axis is created according to two criteria (considered simultaneously):



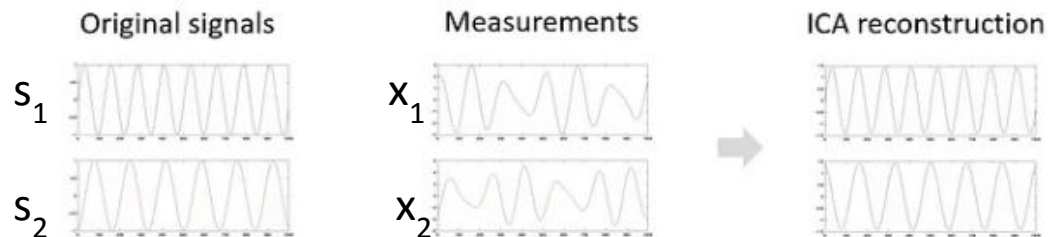
Independent Component Analysis (ICA)



$$x_1(t) = a_{11}s_1 + a_{12}s_2$$

$$x_2(t) = a_{21}s_1 + a_{22}s_2$$

Assume s_1, s_2 : statistically independent



$$\begin{bmatrix} \vec{x}_1 \\ \vec{x}_2 \end{bmatrix} = \mathbf{A} * \begin{bmatrix} \vec{s}_1 \\ \vec{s}_2 \end{bmatrix}$$

$$\begin{bmatrix} \vec{s}_1 \\ \vec{s}_2 \end{bmatrix} = \mathbf{A}^{-1} * \begin{bmatrix} \vec{x}_1 \\ \vec{x}_2 \end{bmatrix}$$

ICA: Στατιστική παρουσίαση

Θεωρήστε ότι δύο ανεξάρτητα στοιχεία (πηγές) περιγράφονται από τις ομοιόμορφες κατανομές:

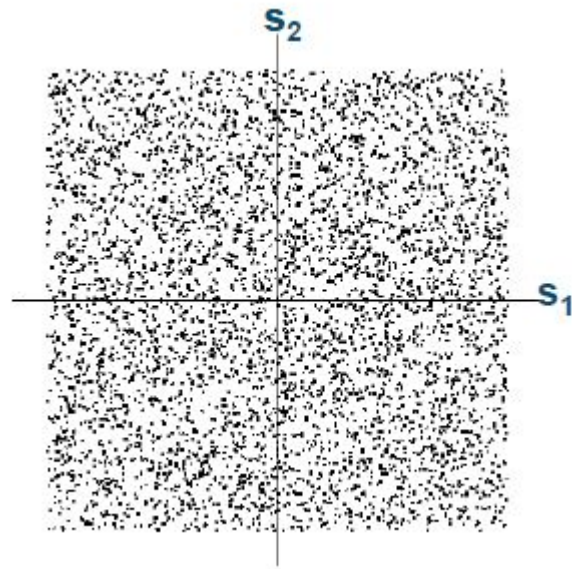
$$p(s_i) = \begin{cases} \frac{1}{2\sqrt{3}} & \text{if } |s_i| \leq \sqrt{3} \\ 0 & \text{otherwise} \end{cases}$$

Αυτή η ομοιόμορφη κατανομή έχει μηδενικό μέσο όρο και η διακύμανση είναι ίση με 1

Ας υποθέσουμε ότι αυτές οι δύο μεμονωμένες πηγές αναμειγνύονται από τον ακόλουθο πίνακα:

$$A_0 = \begin{pmatrix} 2 & 3 \\ 2 & 1 \end{pmatrix}$$

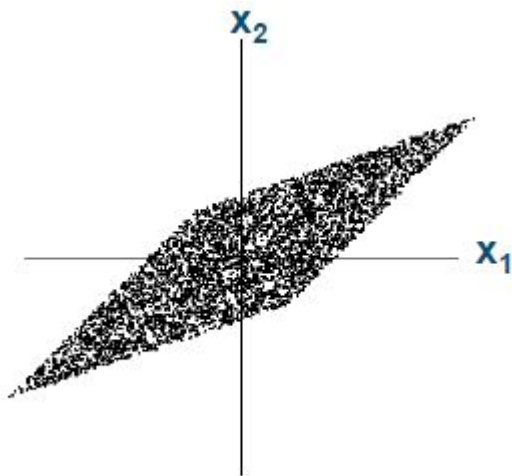
Joint distribution s_i



Οι x μπορούν να υπολογιστούν χρησιμοποιώντας το μοντέλο ICA: $\mathbf{x} = \mathbf{A}\mathbf{s}$

$$\begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{pmatrix} 2 & 3 \\ 2 & 1 \end{pmatrix} \begin{bmatrix} s_1 \\ s_2 \end{bmatrix} \Rightarrow \begin{cases} x_1 = 2s_1 + 3s_2 \\ x_2 = 2s_1 + 1s_2 \end{cases}$$

ICA: Στατιστική παρουσίαση



Joint distribution of mixture x_1, x_2

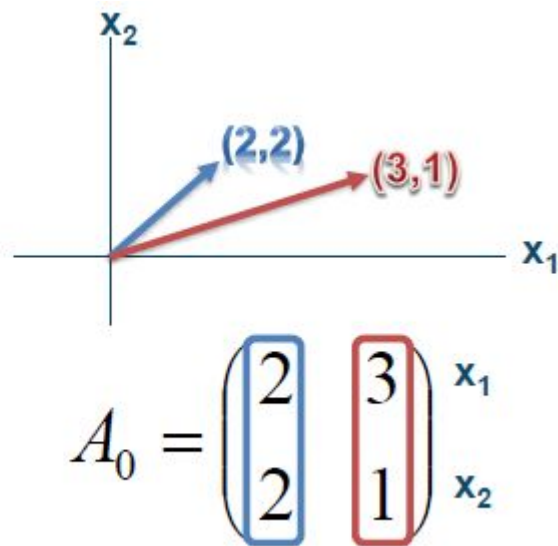
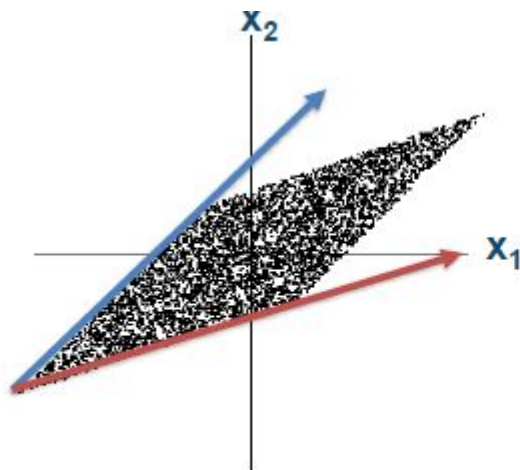
Σημειώστε ότι οι τυχαίες μεταβλητές x_1, x_2 δεν είναι πλέον ανεξάρτητες.

- Ένας εύκολος τρόπος για να το δείτε αυτό είναι να εξετάσετε, εάν είναι δυνατόν να προβλεφθεί η τιμή x_2 από την τιμή της x_1 .
- Παρατηρήστε τον πίνακα A

$$A_0 = \begin{pmatrix} 2 & 3 \\ 2 & 1 \end{pmatrix}$$

ICA: Στατιστική παρουσίαση

Οι ακμές του παραλληλογράμμου είναι οι κατευθύνσεις των στηλών του A_0 .



Αυτό σημαίνει ότι θα μπορούσαμε, κατ' αρχήν, να εκτιμήσουμε το μοντέλο ICA υπολογίζοντας πρώτα την από κοινού κατανομή (joint distribution) των x_1, x_2 , και στη συνέχεια να εντοπίσουμε τις ακμές.

Άρα, το πρόβλημα φαίνεται να έχει λύση.

ICA - Blind Source Separation (BSS)

Στόχος: Προσπαθούμε να εξάγουμε τα ανεξάρτητα συστατικά από τα αναμειγμένα δεδομένα, χωρίς να γνωρίζουμε εκ των προτέρων το πώς έγινε η ανάμειξη των σημάτων.

Σήμα πηγής: Το αυθεντικό, ανεξάρτητο σήμα (π.χ. μια φωνή)

Αναμειγμένα σήμα: Η συνισταμένη δύο ή περισσότερων πηγών που έχουν αναμειχθεί.

Ανάμειξη Σημάτων - ICA Υποθέσεις

Όταν αναμιγνύονται δύο ή περισσότερα σήματα (π.χ. φωνές), παρατηρούνται τρία βασικά αποτελέσματα:

1. **Ανεξαρτησία:** Τα σήματα πηγής είναι στατιστικά ανεξάρτητα, αλλά τα αναμειγμένα σήματα δεν είναι, διότι κάθε ανάμειξη περιέχει και τα δύο σήματα.
2. **Κανονικότητα:** Τα σήματα πηγής έχουν κατανομές που δεν είναι Γκαουσιανές, ενώ τα αναμειγμένα σήματα τείνουν να έχουν Γκαουσιανές κατανομές.
3. **Πολυπλοκότητα:** Η πολυπλοκότητα των αναμειγμένων σημάτων είναι μεγαλύτερη ή ίση με την πολυπλοκότητα των πιο απλών σημάτων πηγής (ο άγνωστος πίνακας ανάμειξης είναι τετράγωνος).

Independent Component Analysis (ICA)

Για να ανακτήσουμε τις αρχικές πηγές από τα αναμεμειγμένα σήματα, χρησιμοποιούμε τις παρακάτω τρεις στρατηγικές:

1. **Ανεξαρτησία:** Εφόσον τα σήματα πηγής είναι ανεξάρτητα, ενώ τα αναμεμειγμένα δεν είναι, $p(x, y) = p(x)p(y)$ το να εξαρτάται
ανεξάρτητα σήματα από τα αναμεμειγμένα σήματα μας επιτρέπει να ανακτήσουμε τις πηγές.
→ Αυτό το κάνει το ICA μέσω της μεγιστοποίησης της στατιστικής ανεξαρτησίας των εξαγόμενων σημάτων.
2. **Κανονικότητα:** Εάν τα σήματα πηγής έχουν μη-Γκαουσιανές κατανομές και τα αναμειγμένα σήματα έχουν Γκαουσιανές κατανομές, τότε η εξαγωγή σημάτων με μη-Γκαουσιανές κατανομές από τα αναμεμειγμένα σήματα μας επιτρέπει να ανακτήσουμε τις πηγές. 
3. **Πολυπλοκότητα:** Εάν τα σήματα πηγής είναι πιο απλά (με χαμηλή πολυπλοκότητα) σε σχέση με τα αναμεμειγμένα σήματα, τότε η εξαγωγή των πιο απλών συστατικών θα μας δώσει τα αρχικά σήματα.

Η στρατηγική είναι να αναζητήσουμε την ιδιότητα (ανεξαρτησία, μη-Γκαουσιανές κατανομές, ή απλότητα) που

Independent Component Analysis (ICA): Παραδοχές

Αριθμός Πηγών: Πρέπει τουλάχιστον ο αριθμός αναμεμειγμένων σημάτων να είναι ίσος με τον αριθμό των πηγών.

Αναμεμειγμένα Σήματα: Αν ο αριθμός των πηγών είναι μεγαλύτερος από τον αριθμό των αναμεμειγμένων σημάτων, δεν μπορεί εύκολα να ανακτηθεί όλος ο αριθμός των πηγών.

Σύγκριση Στρατηγικών για το BSS

- Υπάρχουν διάφορες μέθοδοι για την εφαρμογή του BSS, η κάθε μία με τις δικές της παραδοχές (minimization of mutual information , non-Gaussianity maximization)
- Η μέθοδος που επιλέγεται εξαρτάται από την φύση των δεδομένων και τις υποθέσεις που κάνουμε για τα σήματα και τη διαδικασία ανάμειξης.

ICA - Minimization of mutual information

- Βασίζεται στη θεωρία της πληροφορίας.
- Είναι μέτρο της στατιστικής εξάρτησης μεταξύ συνιστωσών ενός τυχαίου διανύσματος

Διαφορική Εντροπία (H): Η διαφορική εντροπία H ενός τυχαίου διανύσματος y με πυκνότητα πιθανότητας $p(y)$ ορίζεται ως:

$$H(y) = - \int p(y) \log p(y)$$

Αρνητική Εντροπία (Negentropy - J): Μια κανονικοποιημένη έκδοση της εντροπίας, ορίζεται ως αρνητική εντροπία J :

$$J(y) = H(y_{gauss}) - H(y)$$

όπου το y_{gauss} είναι ένα Γκαουσιανό τυχαίο διάνυσμα με τον ίδιο πίνακα συσχέτιση όπως το y .

Η αρνητική εντροπία είναι πάντα θετική, και μηδέν μόνο για Gaussian random vectors

Αμοιβαία Πληροφορία (I): Η αμοιβαία πληροφορία μεταξύ m τυχαίων μεταβλητών y_i , όπου $i=1,...,m$ δίνεται από:

$$I(y_1, y_2, \dots, y_m) = \sum_{i=1}^m H(y_i) - H(y)$$

ICA - Minimization of mutual information

Χρησιμοποιώντας την έννοια της διαφορικής εντροπίας, η αμοιβαία πληροφορία (mutual information) I μεταξύ m τυχαίων μεταβλητών δίνεται από:

$$I(y_1, y_2, \dots, y_m) = \sum_{i=1}^m H(y_i) - H(y)$$

- Η αμοιβαία πληροφορία είναι το φυσικό μέτρο της εξάρτησης μεταξύ τυχαίων μεταβλητών.
- Η τιμή της είναι πάντα μη αρνητική και μηδενική εάν και μόνο εάν οι μεταβλητές εξαρτώνται στατιστικά.

Όταν το αρχικό τυχαίο διάνυσμα $\mathbf{x}$ υφίσταται έναν αντιστρέψιμο γραμμικό μετασχηματισμό $\mathbf{y} = \mathbf{W}\mathbf{x}$ ($\mathbf{x} = \mathbf{A}\mathbf{s}$), η αμοιβαία πληροφορία για το $\mathbf{y}$ ως προς το $\mathbf{x}$ είναι :

$$I(y_1, \dots, y_m) = \sum_i H(y_i) - H(\mathbf{x}) - \log |\det \mathbf{W}|$$

ICA - Minimization of mutual information

Υποθέστε ότι το y_i περιορίζεται να είναι ασυσχέτιστο και με μοναδιαία διακύμανση, τότε θα ισχύει:

$$E\{\mathbf{y}\mathbf{y}^T\} = \mathbf{W}E\{\mathbf{x}\mathbf{x}^T\}\mathbf{W}^T = \mathbf{I}$$

Η εφαρμογή της ορίζουσας σε όλες τις πλευρές της εξίσωσης έχουμε:

$$I(y_1, \dots, y_m) = \sum_i H(y_i) - H(\mathbf{x}) - \log|\det \mathbf{W}| \rightarrow \det I = 1 = \det(\mathbf{W}E\{\mathbf{x}\mathbf{x}^T\}\mathbf{W}^T) = (\det \mathbf{W})(\det E\{\mathbf{x}\mathbf{x}^T\})(\det \mathbf{W}^T)$$

- Το $\det \mathbf{W}$ πρέπει να είναι σταθερό αφού το $\det E\{\mathbf{x}\mathbf{x}^T\}$ δεν εξαρτάται από το $\mathbf{W}$.
- Για y_i μοναδιαίας διακύμανσης, η εντροπία και η αρνητική εντροπία διαφέρουν μόνο κατά σταθερά και πρόσημο.

Επομένως, η θεμελιώδης σχέση ανάμεσα στην εντροπία και την αρνητική εντροπία είναι:

$$I(y_1, \dots, y_n) = C - \sum_i J(y_i)$$

όπου C είναι μια σταθερά που δεν εξαρτάται από το $\mathbf{W}$.

⇒ Βρίσκοντας έτσι έναν αντιστρέψιμο μετασχηματισμό $\mathbf{W}$ που ελαχιστοποιεί την αμοιβαία πληροφορία είναι περίπου ισοδύναμο με την εύρεση κατευθύνσεων στις οποίες η αρνητική εντροπία (μια έννοια που σχετίζεται σε μη γκαουσιανή κατανομή) μεγιστοποιείται.

Maximum Likelihood Estimation

Με τη χρήση της πυκνότητας πιθανότητας ενός γραμμικού μετασχηματισμού υπολογίζεται η πυκνότητα p_x του αναμεμιγμένου δεδομένου $x = As$ όπου $W = A^{-1}$, και f_i δηλώνουν τις πυκνότητες των ανεξάρτητων συστατικών s_i .

$$f_x(x) = |\det W| f_s(s) = |\det W| \prod_{i=1}^n f_i(s_i)$$

Η πυκνότητα p_x μπορεί επίσης να εκφραστεί και του $W = (w_1, w_2 \dots w_n)^T$, δηλαδή,

$$f_x(x) = |\det W| \prod_{i=1}^n f_i(w_i^T x)$$

Υποθέτοντας ότι υπάρχουν T δεδομένα του x , που συμβολίζονται με $x(1), x(2), \dots, x(T)$, και μετά από μερικές πράξεις, η τελική εξίσωση για την πιθανότητα είναι:

$$L = \sum_{t=1}^T \sum_{i=1}^n \log f_i(w_i^T x(t)) + T \log |\det W|$$

•

Πρόβλημα: Οι συναρτήσεις πυκνότητας f_i πρέπει να εκτιμηθούν σωστά, διαφορετικά η Maximum Likelihood Estimation θα δώσει λάθος αποτέλεσμα.

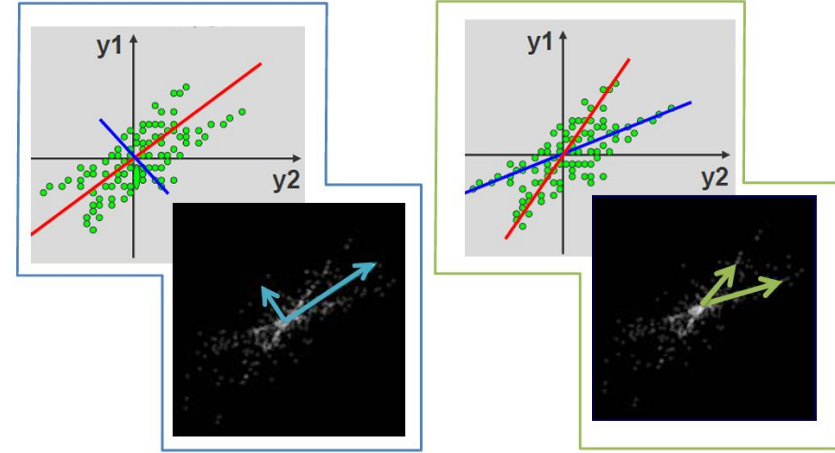
Independent Component Analysis(ICA) vs PCA

Ομοιότητες:

- Εξαγωγή χαρακτηριστικών
- Μείωση διαστάσεων.

Σύγκριση:

- Το PCA βρίσκει κατευθύνσεις μέγιστης διακύμανσης.
- Το ICA βρίσκει κατευθύνσεις μέγιστης ανεξαρτησίας.



PCA finds directions of maximal variance (using second order statistics)

ICA finds directions which maximize independence (using higher order statistics)