

lecture 7/01 : Features (recap). NO CLASS Friday July 4th

- PCA new feat = linear combination (original feat)

- Feature selection using L₁-R_g

- Boosting margins, margin per feature

BINARY Margin(x) = y · classifier score = y · $\frac{F(x)}{\text{score}}$

y ∈ {+1, -1}

if y = +1 : want large score > 0 ⇔ large margin

if y = -1 : want large neg score < 0 ⇔ large margin

if F(x) ≥ 0 ⇔ low margin ⇔ small confidence

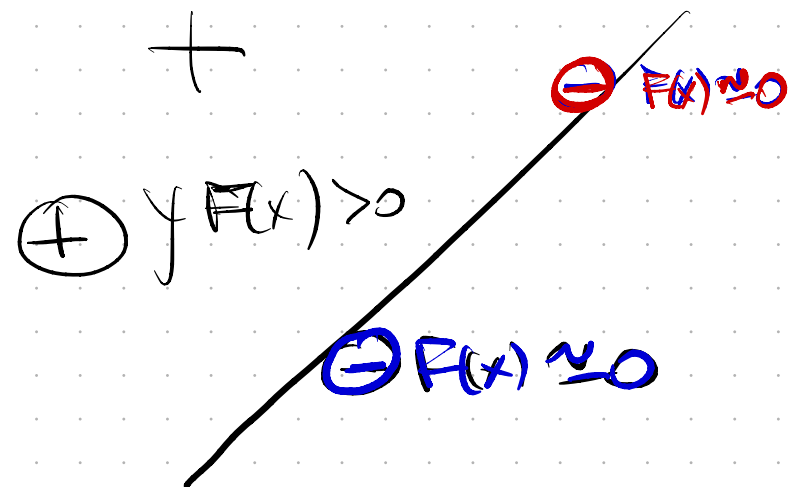
if y F(x) < 0 ⇔ mistake classification

if margin < 0 ⇔ mistake with high confidence

Quality of classifier $\Leftrightarrow$ large (pos) margins

margin better measurement than accuracy

"how far
we predict
from 50-50"
confidence



correct side.

$\ominus y F(x)$ large pos margin

increasing margin $\Leftrightarrow$ increase confidence

GB, Adaboost: increase margins $y F(x)$ during training
even for points classified correctly
already $y F(x) > 0$.

classif
boosting
Alg

margin = $e^{y F(x)}$
maximized

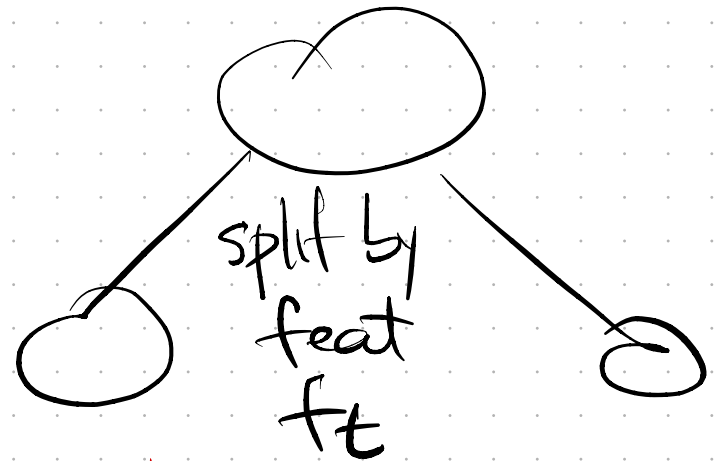
or $\frac{1}{1 + e^{-y F(x)}}$

Margin Analysis GB: $F(x) = \sum_{t=1}^T h_t(x)$

$h_t = \text{dec. stump}$

Adaboost: $F(x) = \sum_{t=1}^T \alpha_t h_t(x)$

Reorder the summation per feature



$$F(x) = \boxed{\sum_{f_t=f_1} \alpha_t h_t(x)} + \boxed{\sum_{f_t=f_2} \alpha_t h_t(x)} + \boxed{\sum_{f_t=f_3} \alpha_t h_t(x)} + \boxed{\sum_{f_t=f_D} \alpha_t h_t(x)}$$

$\text{margin}_{f_1}(x)$ $\text{margin}_{f_2}(x)$ $\text{margin}_{f_3}(x)$ $\text{margin}_{f_D}(x)$

$\Rightarrow$ contribution of each feature to the classifier
Note online, optional PB HW4.

HW4 Feat Selection using L1-reg Regression

→ L1 Reg
→ Logistic Reg
→ Multinom Reg

- use very large L1 penalty train Regression

⇒ many Feat. end up with 0 or close to 0 wgt.

L1 penalty ⇒ sparsity. (acc not that good, irrelevant)

- Select $G = 100$ features by |wgt| largest
200
- Train (real) classifier without L1 penalty.

ex L2-reg Regression

ex. NNNet

ex. Naive Bayes, GDA,

not good: Dec Tree, because D.T. have "natural"
sparsity bc # splits.

PCA recap ① rotates data to align dim with axis

Σ = sigma = covar

Σ = summation

- dim are sorted by $\text{var}[\text{dim}]$

- each dim = projector = linear comb (original feat)

② selects first R dim $\Rightarrow$ losing $\text{var}[\text{dim}]$ information existing on dim $[R+1 : D]$

$$\mu = \sum_{i=1}^N x_i \quad 1 \times D$$

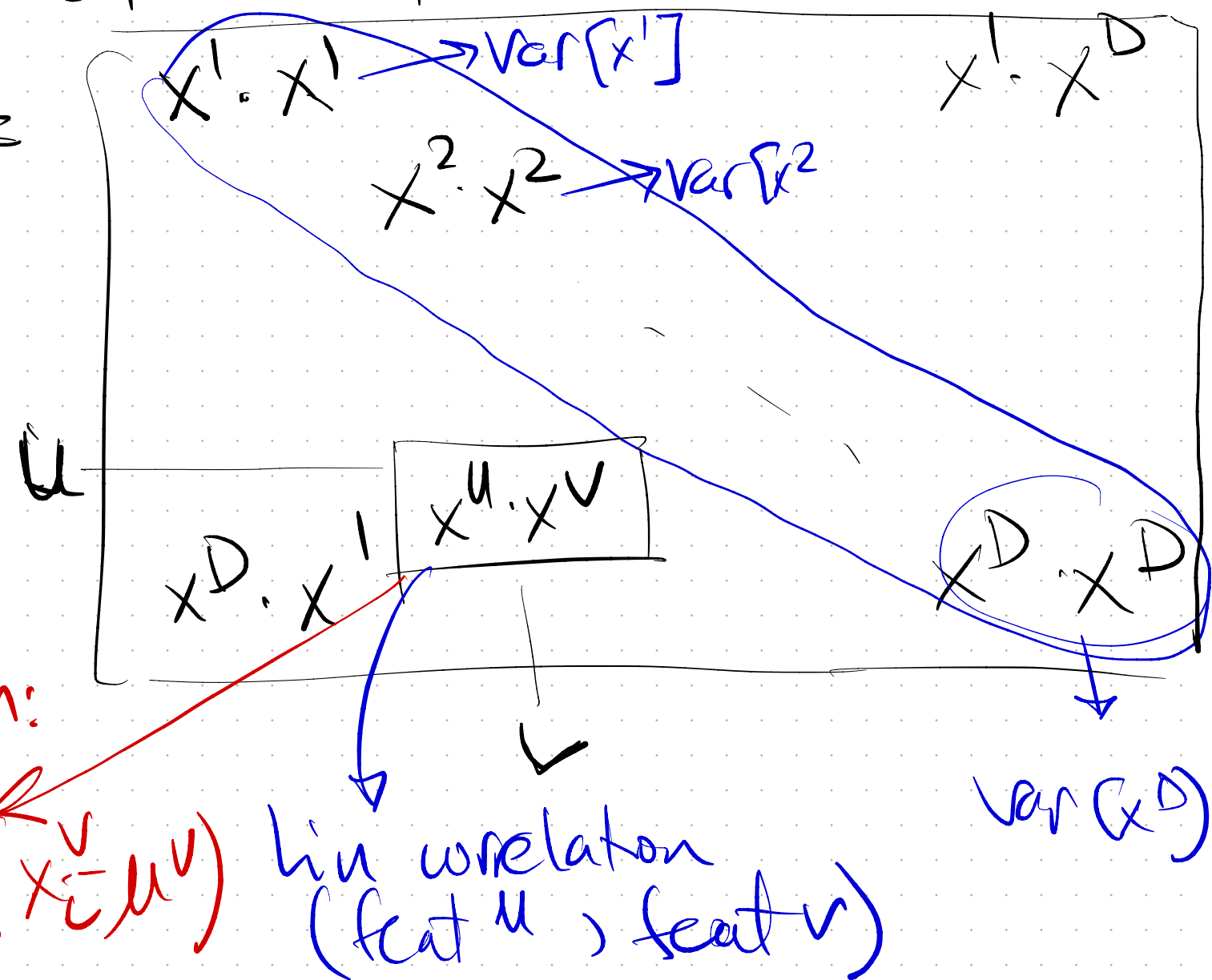
data centered $X = X - \mu$. Each column = feat has 0-mean.

$$\Sigma_{D \times D} = \text{covar}(X) = \frac{1}{N} X^T X \quad D \times N \quad N \times D$$

orthonormal vectors
 $e_1, e_2, \dots, e_D$

• $\|e_i\| = 1$

• $\|e_i e_j\| = 0$ (perpendicular)



$\mu \neq 0$ then:

$$\sum_{i=1}^N (x_i^u - \mu^u)(x_i^v - \mu^v)$$

lin correlation
(feat u , feat v)

Σ = nice math matrix $\Rightarrow$ orthonormal decomposition

- sym
 - pos semidef
 - form $A^T A$
- $(e_1, \lambda_1) (e_2, \lambda_2) \dots (e_D, \lambda_D)$ eigen vector-values
- $e \Sigma = \lambda e$ "e does not change direction mult by Σ "

$$\Sigma_{\text{cov}} = \begin{bmatrix} | & | & & | \\ e_1 & e_2 & \dots & e_D \\ | & | & & | \end{bmatrix} \begin{bmatrix} \lambda_1 & & & 0 \\ & \lambda_2 & & \\ & & \ddots & \\ 0 & & & \lambda_D \end{bmatrix} \begin{bmatrix} - e_1^T \\ - e_2^T \\ \vdots \\ - e_D^T \end{bmatrix}$$

eigen vectors = cols eigen values eigen vect = rows

notes: why $[x \cdot e]$ $\lambda_1 > \lambda_2 > \dots > \lambda_D \geq 0$

① Rotation PCA dim are obtained by projecting x on $e_1, e_2, \dots, e_D$

first dim

$$\begin{array}{ccc} X & \cdot & e_1 \\ N \times D & D \times 1 & N \times 1 \end{array} = X_{\text{new}}^1$$

2nd dim

$$X \cdot e_2 = X_{\text{new}}^2$$

$\Rightarrow X_{\text{new}}^1 = \text{lin combination of orig feat with coef given by } e_1 = [e_1^1 e_2^1 \dots e_D^1]$

② R dim $R \ll D$ (for ex $D=1000$ $R=20$)

$$\underset{N \times D}{X} \cdot \underset{D \times 20}{\begin{bmatrix} | & | & & | \\ e_1 & e_2 & \dots & e_{20} \\ | & | & & | \end{bmatrix}} = \underset{N \times 20}{X_{\text{new}}}$$

Theory: $R=20$ first PCA-dim $\Leftrightarrow$ approx Σ with first 20 (eigenvectors, eigen vals)

$$\underset{\substack{\text{approx} \\ D \times D}}{\Sigma} = \left[\begin{array}{c} \boxed{\begin{array}{c} | & | & & | \\ e_1 & e_2 & \dots & e_R \\ | & | & & | \end{array}}_{D \times R} \quad \boxed{\begin{array}{c} | & | & & | \\ e_{R+1} & \dots & e_D \\ | & | & & | \end{array}}_{R \times (D-R)} \end{array} \right] \left[\begin{array}{c} \boxed{\begin{array}{c} \lambda_1 \quad \lambda_2 \quad \dots \quad \lambda_R \\ \hline \end{array}}_{R \times R} \\ \boxed{\begin{array}{c} \lambda_{R+1} \quad \dots \quad \lambda_D \\ \hline \end{array}}_{(D-R) \times (D-R)} \end{array} \right] \left[\begin{array}{c} \boxed{\begin{array}{c} \hline e_1 \hline \hline e_2 \hline \hline \vdots \hline \hline e_R \hline \end{array}}_{R \times R} \\ \boxed{\begin{array}{c} e_{R+1} \\ \vdots \\ e_D \end{array}}_{(D-R) \times (D-R)} \end{array} \right]_{R \times D}$$

Rank deficient at most R

Information lost. $\Sigma_{\text{approx}} \approx \Sigma$ depends on magnitude of ignored eigenvalues.

- $\lambda_{R+1}, \lambda_{R+2}, \dots, \lambda_D$ close to 0 $\Rightarrow \Sigma_{\text{approx}}$ very good

$R=? \leftarrow$ don't want miss on large eigenvalues.

- if $\lambda_{R+1} = \lambda_{R+2} = \dots = \lambda_D = 0$ (actually 0) $\Rightarrow \Sigma_{\text{approx}} = \Sigma$