

Lecture 6/25

- HW3 demo

- Next: Chapter 4

- boosting / ensemble

- Multilabel data, in particular ECOC

- Feature selection

- Feature aggregation / PCA

- Features from similarity / TSNE

- HW4

- chapter 5: SVM + kernels

- chapter 6: advanced NN

- CNN
 - RecNN vs Transformers

"random Forests"

- Active learning.

$u_t(x)$ = score produced
by t^{th} weak learner
(not probs, not class
prediction)

Boosting (weak learner) = additive ensemble

datapoint
 $x = (x^1, x^2, \dots, x^D)$
 $t = 1:T$
 x_i

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

final score

$$h_t(x)$$

weak scores

$\frac{h(x)}{h(x)}$
 weak learner

acc $\approx 65\%$
 better than
 random

but poor

default: shallow
 decision tree

dec stump: 1-split DT

• Regression mode label y = quantity

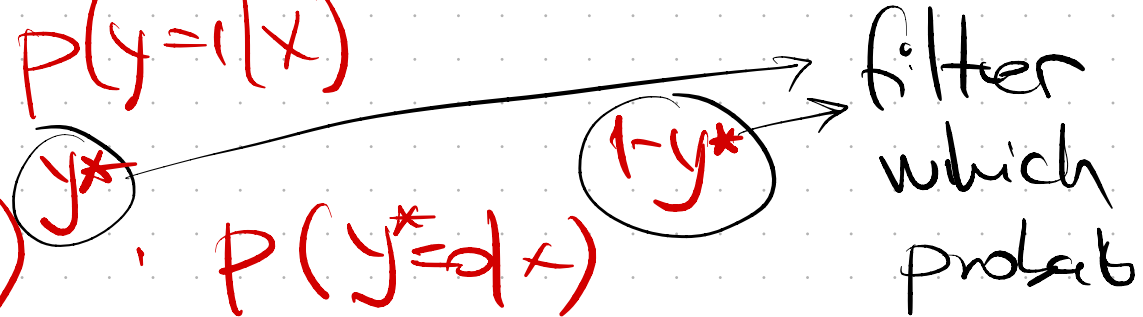
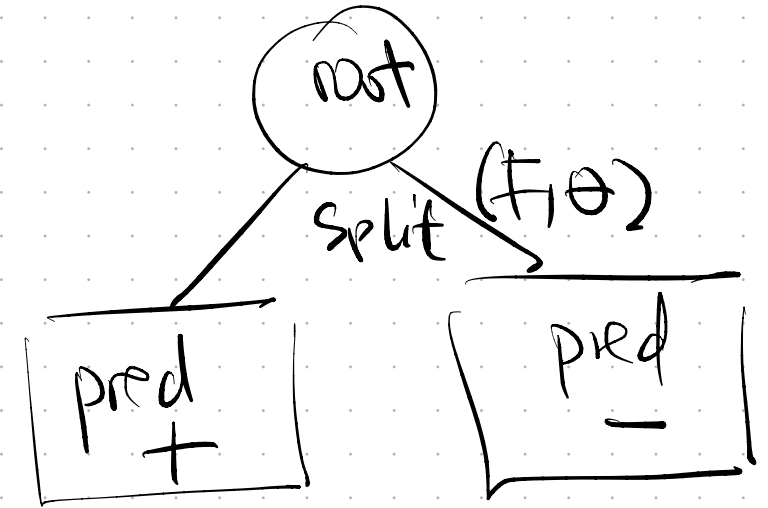
want $(F(x) - y)^2 \approx 0$

all i $\sum_{i=1}^N (F(x_i) - y_i)^2 \approx 0$

• Classification mode $y \in \{-1, +1\}$ $y^* \in \{0, +1\}$

$F(x)$ score $\xrightarrow{\text{map}}$ $\begin{cases} \text{pr}(y=+1|x) \\ \text{pr}(y=-1|x) = 1 - \text{pr}(y=+1|x) \end{cases}$

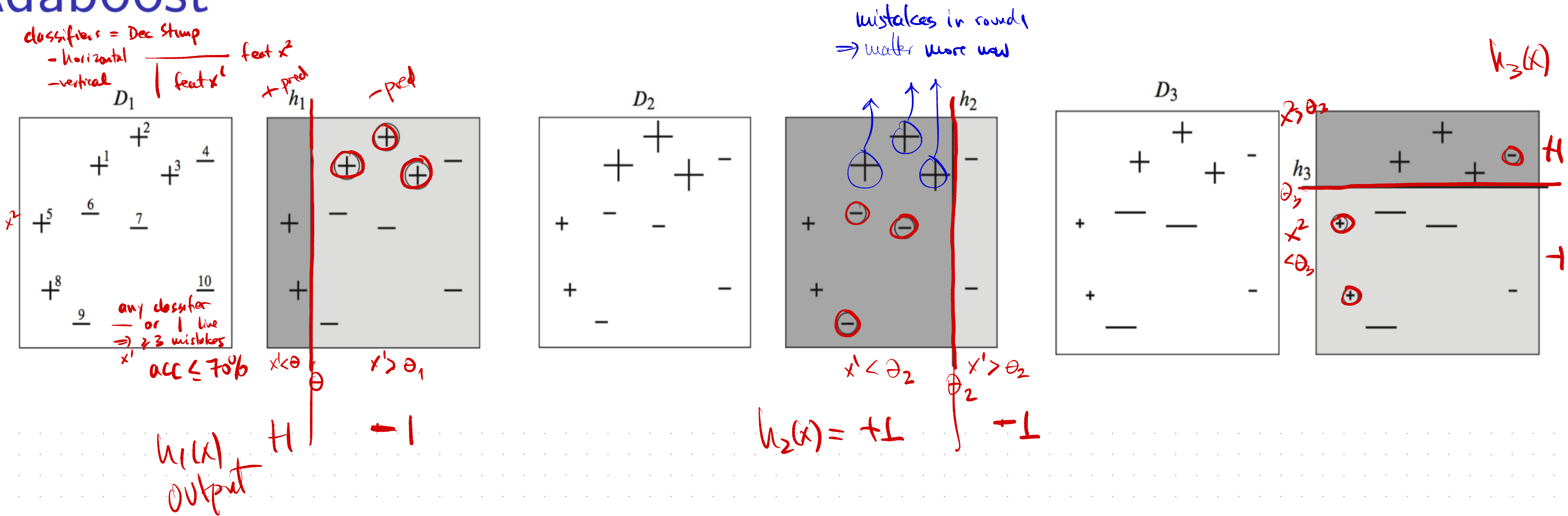
want max likelihood $\prod_{i=1}^N p(y^*_i|x) \cdot p(y^*_i=1|x)$



intuition cca 1992: $F(x)$ cannot be much better than $h(x)$
 VERY WRONG weak learners.

Gradient Boosting = Gradient Descent + Boosting

Adaboost

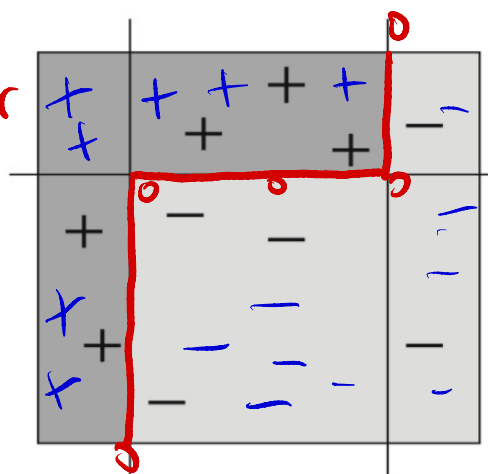


$$F(x) = \sum_t \rho_t h_t(x)$$

(fixed coefficients ρ_t for $h_t(x)$)

$$F = \text{sign} \left(\begin{array}{c} ? \\ 0.42 \\ ? \end{array} + \begin{array}{c} ? \\ 0.65 \\ ? \end{array} + \begin{array}{c} ? \\ 0.92 \\ ? \end{array} \right)$$

$F(x) = \text{add scores}$
 $\text{sign}(F(x)) = \text{classifier}$
 100% acc =



2 ways to implement boosting

1st (historically 1996): ADABOOST

→ look at current weak learner $h_t(x)$

→ reweight data $D_{t+1}(x_i) =$

= importance of datapoint x_i next round

→ train $h_{t+1}() =$ weak learner

with error weighted by $D_{t+1}()$ distribution

$$e_{t+1} = \sum_{i=1}^N \mathbb{1}[h_{t+1}(x_i) \neq y_i] \cdot D_{t+1}(x_i)$$

acc $\begin{cases} 1 & \text{mistake} \\ 0 & \text{correct} \end{cases}$

weight
for round $t+1$

add $h_{t+1}()$ to $F(x)$

$$F(x) = F(x) + h_{t+1}(x)$$

$$F(x) = h_1(x) + h_2(x) + \dots + h_{t+1}(x)$$

Key: After $t = 1 \dots T$ rounds of Adaboost

weight on x $D_{t+1}(x) \approx \# \text{ rounds } t \text{ where } x \text{ mistake}$

$$\approx \exp \left(\sum_{t=1}^T \alpha_t [h_t(x) \cdot -y] \right)$$

$-1: h_t(x) = y$ correct
 $+1: h_t(x) \neq y$ incorrect

• $D_{t+1}(x)$ small
 $\Rightarrow$ most classif $h_t(x) = y$

$D_{t+1}(x) = \text{large} \Rightarrow$ next classifier $h_{t+1}(x)$ focus on datapoint x

• $D_{t+1}(x) = \text{large} \Rightarrow$ most classif $h_t(x) \neq y$

Th proper weights, $h()$ training etc (see Adaboost details)
training can finish in one of those 2 ways.

- either train error ≈ 0
 $F(x)$ ensemble $\sum_{i=1}^N \frac{1}{[F(x_i) \neq y_i]} \rightarrow 0$

• or stuck: $D_{T+1}()$ dist of point importance for next round is s.t. the best $h_{T+1}(x)$ classifier is random

$$E_{T+1} = \sum_{i=1}^N \frac{1}{[h_{T+1}(x_i) \neq y_i]} \cdot D_{T+1}(x_i) \approx 0.5$$

as long as $h_{T+1}(x)$ makes some progress (better than 0.5)
training error ($F(x)$) has to go to zero.

2nd way
much better

$$F(x) = h_1(x) + h_2(x) + \dots + h_T(x)$$

Gradient
Boosting.

want $F(x) \approx y$ (regression mode)
"house prices"

next weak
learner

$$h_{T+1}(x) = F_{\text{new}}(x) - F_{\text{current}}(x)$$

$$\text{error } \frac{1}{2}(F(x) - y)^2 = J(x)$$

diff
w.r.t $F(x)$

$$= F(x) - y = \frac{\partial J}{\partial F(x)}$$

? Grad descent mode for $F(x)$ as variable ?? Simple $\lambda=1$

$$F_{\text{new}}(x) = F(x) - \lambda \cdot \frac{\partial J}{\partial F(x)} = F(x) + y - F(x)$$

trivial? want $F_{\text{new}}(x) = y$ (wrong thinking)

GD update : $F_{\text{new}}(x) = F(x) + \underbrace{h_{T+1}(x)}_{\approx y - F(x)}$

want the next weak learner

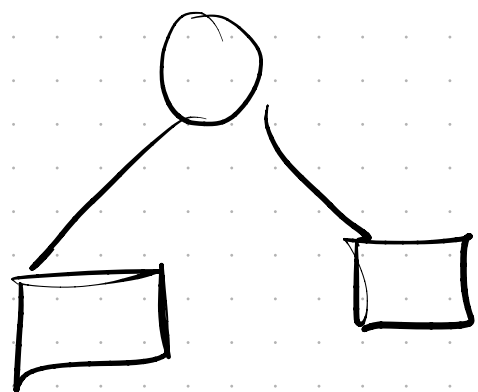
$$h_{T+1}(x) \approx y - F(x)$$

- if $h_{T+1}(x) \approx y - F(x)$ well $\Rightarrow$

next w. learner *residual label*

$$F_{\text{new}}(x) = F(x) + h_{T+1}(x) \approx y \quad (\text{good})$$

- practice: train $h_{T+1}()$ to match $h_{T+1}(x) \approx y - F(x)$
weak learner
"Dec Stump"
new label



- req updating label after every round
 $y_{\text{new}} = y_{\text{orig}} - F(x)$

- does not require weights like Adaboost, modify train procedure

Train $h_{T+1}(x)$: $\text{DecTreeAlg}(x, y - F(x))$

Classification $y \in \{-1, 1\}$ $y^* = \frac{1+y}{2} \in \{1, 0\}$ categories

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

Score

$$p(x) = p(y^*=1|x) = \frac{e^{F(x)}}{e^{F(x)} + e^{-F(x)}}$$

$$p(\bar{x}) = p(y^*=0|x) = \frac{e^{-F(x)}}{e^{F(x)} + e^{-F(x)}} = 1 - p(x)$$

want MAX

$$\text{log likelihood} = \log \left(\prod_{i=1}^N p(x)^{y^*} (1-p(x))^{1-y^*} \right)$$

filter $y^* \rightarrow 0$
 $y^* \rightarrow 1$ which probab.