

(1)

Ada Boost.

Note Title

5/18/2008

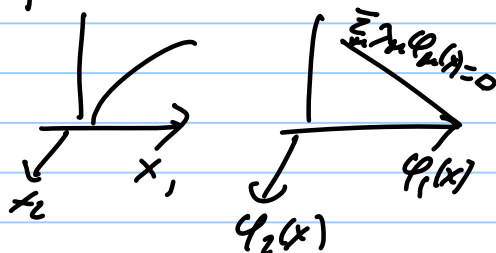
AdaBoost is a method for combining a number of weak classifiers to make a strong classifier.

Input: set of weak classifiers  $\{\varphi_\mu(x) : \mu = 1 \text{ to } M\}$   
 labelled data  $\{(\underline{x}^t, y^t) : t = 1 \text{ to } N\}$   
 $y^t \in \{-1, 1\}$

Output: strong classifier  
 $\text{sign} \left( \sum_{\mu=1}^M \lambda_\mu \varphi_\mu(x) \right)$   
 $\{\lambda_\mu\}$  weights / coefficients.

The strong classifier is a plane in feature space  $\{\varphi_\mu(x)\}$ .

In practice, most of the  $\lambda_\mu = 0$ .



The "selected" weak classifiers are those with  $\lambda_\mu \neq 0$ .

(2) The task of AdaBoost is to select weights  $(\lambda_\mu)$  to make the strong classifier as effective as possible.

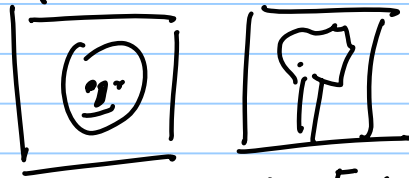
The motivation is that it is often possible to obtain weak classifier for classification tasks - i.e. a weak classifier that is effective 60% of the time. Want to build a strong classifier - effective 99% of the time - that is built by combining weak classifiers.

The combination is by weighted summation  $\sum_{\mu} \lambda_{\mu} P_{\mu}(x)$

Why linear weighted combination?

Because we can use an efficient algorithm - AdaBoost - to estimate them.

### (3) Example: Face Detection (Viola & Jones)



Face

Non-face

threshold

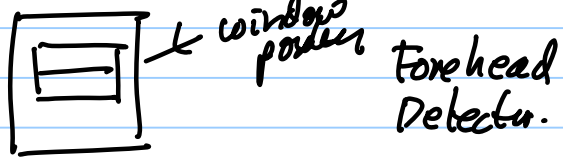
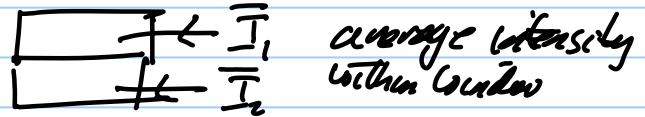
Training examples are a set of images  $x^t$

labelled by  $y^t = 1$  if image contains a face,  $y^t = -1$  if not.

Weak classifier:

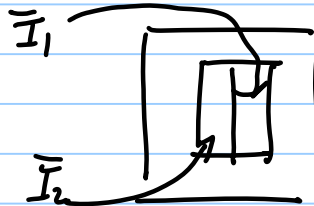
Face: if  $\bar{I}_1(x) - \bar{I}_2(x) > T$

Intensity in forehead is bigger than intensity in eye region.



Forehead Detector.

face: if  $\bar{I}_1 - \bar{I}_2 < \epsilon_1$   
 $\bar{I}_2 - \bar{I}_1 < \epsilon_2$



Symmetry Detector  
 $\rightarrow$  Faces symmetrical

where  $\epsilon_1$  &  $\epsilon_2$  are small constants.

Easy to get weak classifiers of this type — weak classifier is features  $\bar{I}_1(x), \bar{I}_2(x)$  + threshold  $T$  — but each weak classifier is only partially success. AdaBoost gives a way to select weak classifiers and combine them to make a strong classifier.

(Algorithm lab)

(4) AdaBoost: Mathematical Description.

Define  $Z[\lambda_1, \dots, \lambda_M] = \sum_{t=1}^N e^{-y^t \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t)}$

This is an upper bound of the error rate of the strong classifier  $S(x) = \text{sign}(\sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x))$ .

Error Rate:  $E[\lambda_1, \dots, \lambda_M] = \sum_{t=1}^N \{1 - I(S(x^t), y^t)\}$

where  $I(S(x^t), y^t) = 1$ , if  $S(x^t) = y^t$  (correct answer)  
 $= 0$ , otherwise.

Claim:  $E[\lambda_1, \dots, \lambda_M] \leq Z[\lambda_1, \dots, \lambda_M]$

compare each term in the summation  $\sum_{t=1}^N$

case (i) If  $S(x^t) = y^t$ , then  $\text{sign}(\sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t)) = \text{sign } y^t$

so  $y^t \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t) = A > 0$  (Definition A)

Error Term  $\langle 1 - I(S(x^t), y^t) \rangle = 0$

Z Term  $e^{-y^t \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t)} = e^{-A} > 0$  ✓

case (ii) If  $S(x^t) \neq y^t$ , then  $y^t \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t) = -B < 0$  ( $B > 0$ )

Error Term  $\langle 1 - I(S(x^t), y^t) \rangle = 1$

Z term  $e^{-y^t (\sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t))} = e^B > 1$  ✓

so Z term is bigger than error term in both cases.

Goal: AdaBoost minimize

(5)  $\sum [\lambda_1 \dots \lambda_N]$ . This will guarantee that the error rate is small (but not necessarily the minimum error rate).

Strategy to minimize  $\sum [\lambda_1 \dots \lambda_N]$ .

Initialize by  $\lambda_1 = \dots = \lambda_N = 0$ .  
(i.e.  $H_0(x) = 0 \rightarrow$  no weak classifier selected).

Minimize  $\sum [\lambda_1 \dots \lambda_N]$  by coordinate descent.

At time step  $l$ .

For each  $i$  <sup>state</sup>  $\lambda_1^l, \dots, \lambda_N^l$ .  
minimize  $\sum$  w.r.t.  $\lambda_i$   
with  $\lambda_j^l$  fixed for  $j \neq i$ .

Solve  $\frac{\partial \sum}{\partial \lambda_i} = 0$  to solve for  $\hat{\lambda}_i$   
for each  $i$ .

compute  $\sum [\lambda_1^l, \dots, \lambda_{i-1}^l, \hat{\lambda}_i, \lambda_{i+1}^l, \dots, \lambda_N^l]$  for each  $i$ .

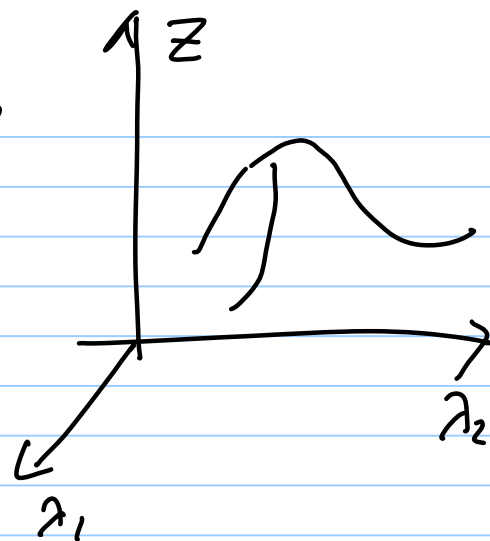
select  $\hat{i} = \text{Argmin}_i \sum [\lambda_1^l, \dots, \lambda_{i-1}^l, \lambda_i, \lambda_{i+1}^l, \dots, \lambda_N^l]$

set  $\lambda_j^{l+1} = \lambda_j^l$ ,  $j \neq \hat{i}$ ,  $\lambda_{\hat{i}}^{l+1} = \hat{\lambda}_{\hat{i}}$ .

(6) Intuition: at each time  
step 1.

calculate how much you  
can decrease  $Z$  by changing  
only one of the  $\{\lambda_i\}$ .

select the  $\lambda_i^*$  which  
maximally decreases  $Z$ .



Each step of this algorithm is guaranteed  
to decrease  $Z$ .

So algorithm will converge to a minimum  
of  $Z$ . (But  $Z$  is convex in  $\lambda$ , so the algorithm  
converges to global minimum).

Why  $Z$ ? Why this algorithm?

Practical — we can compute  $\frac{\partial Z}{\partial \lambda_i}$  and the  
minimum of  $Z$  easily.

(7)

AdaBoost Algorithm.Data  $\{ (x^t, y^t) : t=1, \dots, N \}$ Set of weak classifiers  $\{ \phi_\mu(x), \mu=1, \dots, M \}$ .

For each weak classifier, divide the training data into two classes

$$W_\mu^+ = \{ t : y^t \phi_\mu(x^t) = 1 \} \quad \phi_\mu \text{ correct}$$

$$W_\mu^- = \{ t : y^t \phi_\mu(x^t) = -1 \} \quad \phi_\mu \text{ wrong.}$$

$$W_\mu^+ \cup W_\mu^- = \{ 1, \dots, N \}$$

At each time step  $l$ , define a set of "weights" for the training examples.

$$D_t^l = \frac{e^{-y^t \sum_{\mu=1}^M \lambda_\mu^l \phi_\mu(x^t)}}{\sum_t e^{-y^t \sum_{\mu=1}^M \lambda_\mu^l \phi_\mu(x^t)}} \quad \sum_t D_t^l = 1$$

$$D_t^l \geq 0.$$

Gives bigger weights to data incorrectly classified by current "strong classifier".

i.e.  $D_t^l$  is large if  $y^t \sum_{\mu=1}^M \lambda_\mu^l \phi_\mu(x^t) < 0$

implies  $\text{sign}(\sum_{\mu=1}^M \lambda_\mu^l \phi_\mu(x^t)) \neq y^t$ .

$D_t^l$  is small if  $y^t \sum_{\mu=1}^M \lambda_\mu^l \phi_\mu(x^t) > 0$ .

implies  $\text{sign}(\sum_{\mu=1}^M \lambda_\mu^l \phi_\mu(x^t)) = y^t$

(2)

## Adaboost Algorithm.

Initialize  $\lambda_1 = \lambda_2 = \dots = \lambda_N = 0$

At time step  $\lambda_1^l, \lambda_2^l \dots \lambda_N^l$

For each  $i$ , calculate  $\Delta_i^l = \frac{1}{2} \log \left( \frac{\sum_{t \in w_i^+} D_t^l}{\sum_{t \in w_i^-} D_t^l} \right)$

(change in  $\Delta_i^l$  due to solving  $\frac{\partial Z}{\partial \lambda_i} = 0$ , see later)

calculate  $\sqrt{\sum_{t \in w_i^+} D_t^l} \sqrt{\sum_{t \in w_i^-} D_t^l}$

(change in  $Z$  due to setting  $\lambda_i$  to  $\hat{\lambda}_i$ , see later)

select  $\hat{i} = \underset{i}{\text{ARG MIN}} \sqrt{\sum_{t \in w_i^+} D_t^l} \sqrt{\sum_{t \in w_i^-} D_t^l}$

set  $\lambda_j^{l+1} = \lambda_j^l, j \neq i$

$\lambda_i^{l+1} = \lambda_i^l + \Delta_i^l.$

repeat, until convergence.



(9) Intuition for the  $\sum_{t \in w_i^+} D_t^l$  and  $\sum_{t \in w_i^-} D_t^l$  terms.

Initially,  $D_t^{l=0} = \frac{1}{N}$  the data is equally weighted.

$\sum_{t \in w_i^+} D_t^{l=0}$  is the proportion of data that is correctly classified by  $\phi_i(x)$

$\sum_{t \in w_i^-} D_t^{l=0}$  is proportion incorrectly classified

. i.e.  $\sum_{t \in w_i^-} D_t^{l=0}$  is the normalized

error rate if we just use classifier  $\phi_i(x)$   
- normalized by  $1/N$  //

For  $l \neq 0$ ,  $\sum_{t \in w_i^+} D_t^l$  is data correctly classified by  $\phi_i(x)$  taking into account the previously selected classifier (those for which  $\lambda_j^l \neq 0$ ).

$\phi_i(x)$  is a useless classifier if  $\sum_{t \in w_i^+} D_t^l = \sum_{t \in w_i^-} D_t^l$   
i.e. weighted error =  $1/2$ .

corresponds to  $\Delta_i^l = 0$  (i.e. no change in weight)  
 $\lambda_i$

(10)

Term  $\frac{1}{\sqrt{\sum_{t \in w_i^+} D_t^l}} \frac{1}{\sqrt{\sum_{t \in w_i^-} D_t^l}}$  is a non-linear function of the weighted error rate of  $\phi_i(x)$ .

It can be rewritten as

$$\frac{1}{\left( \sum_{t \in w_i^+} D_t^l \right) \left( 1 - \sum_{t \in w_i^+} D_t^l \right)}$$

because  $\sum_{t \in w_i^+} D_t^l + \sum_{t \in w_i^-} D_t^l = 1$

It's smallest values are if

$\sum_{t \in w_i^+} D_t^l = 0$  ,  $\phi_i(x)$  has optimal weighted classified

$\sum_{t \in w_i^+} D_t^l = 1$  ,  $\phi_i(x)$  worst possible classifier  
 $\rightarrow$  implies  $-\phi_i(x)$  best possible classifier

its largest values are when

$\sum_{t \in w_i^+} D_t^l = \frac{1}{2}$  , i.e. when weighted error is  $\frac{1}{2}$  ,  $\phi_i(x)$  useless

(11)

When does AdaBoost converge?

It stops when all weak classifiers are useless - i.e. when  $\sum_{t \in w_i^+} D_t^l = 1/2$ , for all  $i$ .

In this case  $\sqrt{\sum_{t \in w_i^+} D_t^l} = \sqrt{\sum_{t \in w_i^-} D_t^l}$  takes its biggest (i.e. worst) value of  $1/2$ , for all  $i$ .

The weight update  $\Delta_i^l = \frac{1}{2} \log \left\{ \frac{\sum_{t \in w_i^+} D_t^l}{\sum_{t \in w_i^-} D_t^l} \right\}$  is 0 for all  $Q_i(x)$  (since  $\log 1 = 0$ ).

In general, at time step  $l$  select the classifier  $Q_i(x)$  with smallest "weighted error rate"

$$\sqrt{\sum_{t \in w_i^+} D_t^l} = \sqrt{\sum_{t \in w_i^-} D_t^l} \quad \Delta_i^l = \frac{1}{2} \log \left\{ \frac{\sum_{t \in w_i^+} D_t^l}{\sum_{t \in w_i^-} D_t^l} \right\}$$

$$\text{Update } \lambda_i^l \rightarrow \lambda_i^l + \Delta_i^l$$

the smaller the weighted error rate - i.e. smaller  $\sum_{t \in w_i^+} D_t^l$  or  $\sum_{t \in w_i^-} D_t^l$  - then the bigger the change  $\Delta_i^l$ .

(12) How does AdaBoost algorithm relate to AdaBoost mathematics?

AdaBoost mathematics requires:

(i) efficient solution of  $\frac{\partial Z}{\partial \lambda_i} = 0$ .

(ii) efficient computation of  $Z$ .

(i)  $\frac{\partial Z}{\partial \lambda_i}$

$$\frac{\partial Z}{\partial \lambda_i} = \sum_{t=1}^n \{-y^t \phi_i(x^t)\} e^{-y^t \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t)}$$

sol.  $\lambda_i \rightarrow \lambda_i + \Delta_i$  solve for  $\Delta_i$ .

$$\frac{\partial Z}{\partial \lambda_i} = 0 \Rightarrow \sum_{t=1}^n \{y^t \phi_i(x^t)\} e^{-y^t \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t)} e^{-y^t \Delta_i \phi_i(x^t)} = 0$$

$$\sum_{t=1}^n \{y^t \phi_i(x^t)\} D_t e^{-y^t \Delta_i \phi_i(x^t)} = 0.$$

Divide

$$\sum_{t=1}^n = \sum_{t \in w_i^+} + \sum_{t \in w_i^-}$$

$$\sum_{t \in w_i^+} D_t e^{-\Delta_i} - \sum_{t \in w_i^-} D_t e^{\Delta_i} = 0.$$

$$e^{2\Delta_i} = \left( \sum_{t \in w_i^+} D_t \right) / \left( \sum_{t \in w_i^-} D_t \right)$$

$$\Delta_i = \frac{1}{2} \log \left( \sum_{t \in w_i^+} D_t / \sum_{t \in w_i^-} D_t \right) //$$

$$\left\{ \begin{array}{l} \text{Recall } D_t = \frac{e^{-y^t \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t)}}{\sum_{t=1}^n e^{-y^t \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t)}} \end{array} \right.$$

(13) (2) Computation of  $Z$ .

$$Z[\lambda_1, \dots, \lambda_i, \lambda_i + \Delta_i, \dots, \lambda_n]$$

$$= \sum_{t=1}^N e^{-y^t \left( \sum_{\mu=1}^M \lambda_{\mu} \phi_{\mu}(x^t) + \Delta_i \phi_i(x^t) \right)}$$

$$= K \sum_{t=1}^N D_t e^{-y^t \Delta_i \phi_i(x^t)}$$

where  $K = \sum_{t=1}^N D_t e^{-y^t \Delta_i \phi_i(x^t)}$   
is independent of  $i$ .

$$Z[\lambda_1, \dots, \lambda_i, \lambda_i + \Delta_i, \dots, \lambda_n]$$

$$= K \left\{ \sum_{t \in w_i^+} D_t e^{-\Delta_i} + \sum_{t \in w_i^-} D_t e^{\Delta_i} \right\}$$

$$= 2K \sqrt{\sum_{t \in w_i^+} D_t} \sqrt{\sum_{t \in w_i^-} D_t}$$

using  $e^{\Delta_i}$  from previous page.

Hence, coordinate descent reduces to  
computing  $\Delta_i = \frac{1}{2} \log \left( \frac{\sum_{t \in w_i^+} D_t}{\sum_{t \in w_i^-} D_t} \right) \rightarrow$  how much to  
change  $\lambda_i$

solving  $\hat{\lambda}_i = \arg \min_i \sum_{t \in w_i^+} D_t e^{-\Delta_i} + \sum_{t \in w_i^-} D_t e^{\Delta_i}$   
to find best  $\lambda_i$  to change.

## (14) Error to Avoid in AdaBoost.

Once a weak classifier  $\phi_i(x)$  has been selected, it can be selected again.

This should be obvious from the coordinate descent formulation - if you decide to update  $\lambda_i$  at time step  $l$ , then you can also update  $\lambda_i$  at a later time step.

---

## Probabilistic Interpretation.

It can be shown (Friedman, Hastie, Tibshirani) that AdaBoost relates to logistic regression.

$$P(y | x) = \frac{e^{y \sum_{\mu} \lambda_{\mu} \phi_{\mu}(x)}}{e^{\sum_{\mu} \lambda_{\mu} \phi_{\mu}(x)} + e^{-\sum_{\mu} \lambda_{\mu} \phi_{\mu}(x)}}$$

This result is asymptotic - only true in the limit as the number of samples  $n$  becomes infinitely large.

Note: standard sigmoid regression means that you specify a small number of features  $\phi_{\mu}(x)$  that are not necessarily binary-valued.

(15)

The main advantage of AdaBoost is that you can specify a large set of weak classifiers and the algorithm decides which weak classifier to use - by assigning them non-zero  $\lambda_i$ .

Standard logistic regression only uses a small set of features (like weak classifiers).

SVM uses the kernel trick  $K(x, x') = \phi(x) \cdot \phi(x')$  to simplify the dependence on  $\phi(x)$ , but doesn't say how to select  $K(\cdot, \cdot)$  or  $\phi(\cdot)$ .

Multilayer perceptron can be interpreted as selecting weak classifiers  $\rightarrow$  but in a non-optimal manner.

Recent work, suggests that AdaBoost can be improved by making it more similar to logistic regression.

AdaBoost can be extended to multiclass.