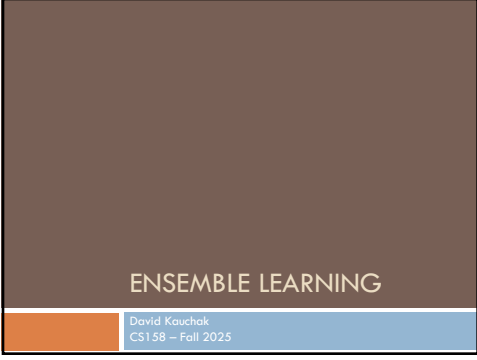
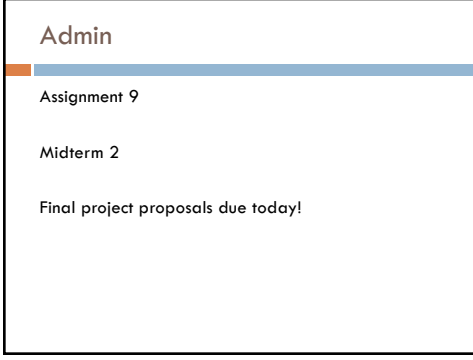




ENSEMBLE LEARNING

David Kauchak
CS158 – Fall 2025

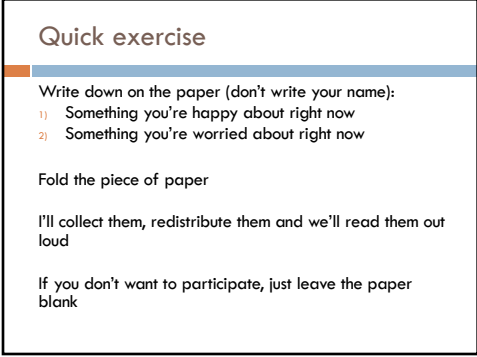
1



Admin

- Assignment 9
- Midterm 2
- Final project proposals due today!

2



Quick exercise

Write down on the paper (don't write your name):

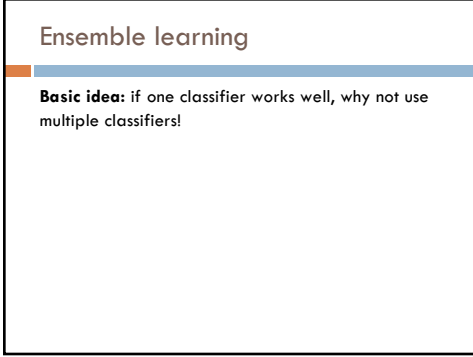
- 1) Something you're happy about right now
- 2) Something you're worried about right now

Fold the piece of paper

I'll collect them, redistribute them and we'll read them out loud

If you don't want to participate, just leave the paper blank

3



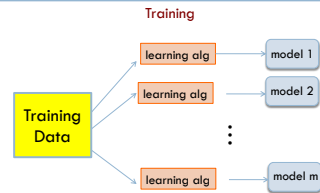
Ensemble learning

Basic idea: if one classifier works well, why not use multiple classifiers!

4

Ensemble learning

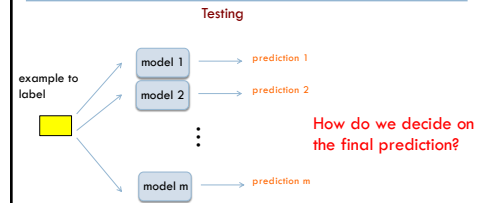
Basic idea: if one classifier works well, why not use multiple classifiers!



5

Ensemble learning

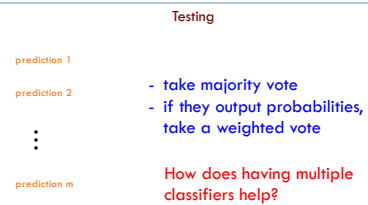
Basic idea: if one classifier works well, why not use multiple classifiers!



6

Ensemble learning

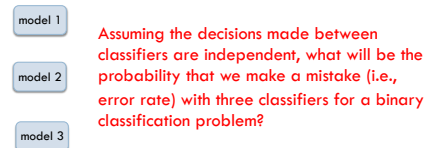
Basic idea: if one classifier works well, why not use multiple classifiers!



7

Benefits of ensemble learning

Assume each classifier makes a mistake with some probability (e.g. 0.4, that is a 40% error rate)



8

Benefits of ensemble learning

Assume each classifier makes a mistake with some probability (e.g. 0.4, that is a 40% error rate)

model 1	model 2	model 3	prob
C	C	C	.6*.6*.6=0.216
C	C	I	.6*.6*.4=0.144
C	I	C	.6*.4*.6=0.144
C	I	I	.6*.4*.4=0.096
I	C	C	.4*.6*.6=0.144
I	C	I	.4*.6*.4=0.096
I	I	C	.4*.4*.6=0.096
I	I	I	.4*.4*.4=0.064

9

Benefits of ensemble learning

Assume each classifier makes a mistake with some probability (e.g. 0.4, that is a 40% error rate)

model 1	model 2	model 3	prob
C	C	C	.6*.6*.6=0.216
C	C	I	.6*.6*.4=0.144
C	I	C	.6*.4*.6=0.144
C	I	I	.6*.4*.4=0.096
I	C	C	.4*.6*.6=0.144
I	C	I	.4*.6*.4=0.096
I	I	C	.4*.4*.6=0.096
I	I	I	.4*.4*.4=0.064

0.096+
0.096+
0.096+
0.064 =
35% error!

10

Benefits of ensemble learning

3 classifiers in general, for r = probability of mistake for individual classifier:

$$p(\text{error}) = 3r^2(1-r) + r^3 \quad \text{binomial distribution}$$

r	p(error)
0.4	0.35
0.3	0.22
0.2	0.10
0.1	0.028
0.05	0.0073

11

Benefits of ensemble learning

5 classifiers in general, for r = probability of mistake for individual classifier:

$$p(\text{error}) = 10r^3(1-r)^2 + 5r^4(1-r) + r^5$$

r	p(error) 3 classifiers	p(error) 5 classifiers
0.4	0.35	0.32
0.3	0.22	0.16
0.2	0.10	0.06
0.1	0.028	0.0086
0.05	0.0073	0.0012

12

Benefits of ensemble learning

m classifiers in general, for r = probability of mistake for individual classifier:

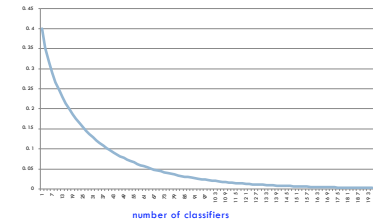
$$p(\text{error}) = \sum_{i=(m+1)/2}^m \binom{m}{i} r^i (1-r)^{m-i}$$

(cumulative probability distribution for the binomial distribution)

13

Given enough classifiers...

$$p(\text{error}) = \sum_{i=(m+1)/2}^m \binom{m}{i} r^i (1-r)^{m-i} \quad r = 0.4$$



14

What's the catch?

Assume each classifier makes a mistake with some probability (e.g. 0.4, that is a 40% error rate)

model 1

model 2

model 3

Assuming the decisions made between classifiers are independent, what will be the probability that we make a mistake (i.e. error rate) with three classifiers for a binary classification problem?

15

What's the catch?

Assume each classifier makes a mistake with some probability (e.g. 0.4, that is a 40% error rate)

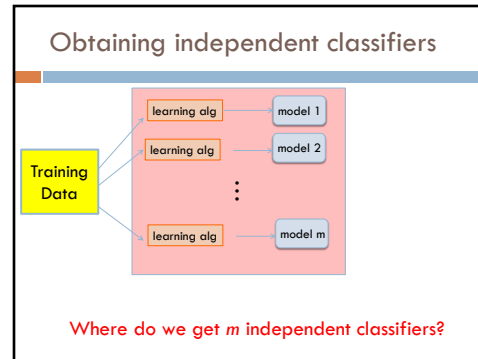
model 1

model 2

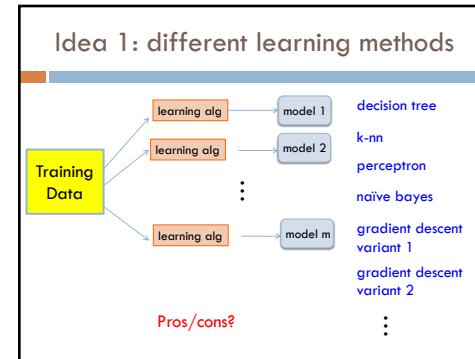
model 3

Assuming the decisions made between classifiers are independent, what will be the probability that we make a mistake (i.e. error rate) with three classifiers for a binary classification problem?

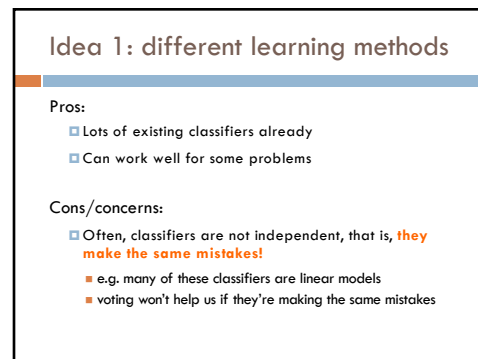
16



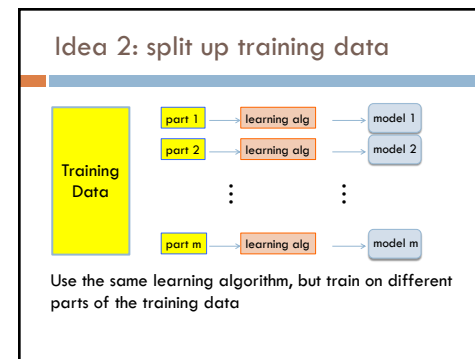
17



18



19



20

Idea 2: split up training data

Pros:

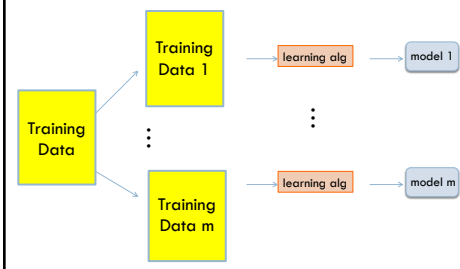
- ▣ Learning from different data, so can't overfit to same examples
- ▣ Easy to implement
- ▣ fast

Cons/concerns:

- ▣ Each classifier is only training on a small amount of data
- ▣ Not clear why this would do any better than training on full data and using good regularization

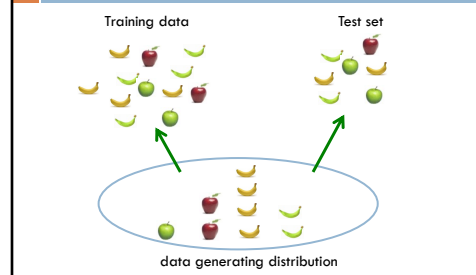
21

Idea 3: bagging



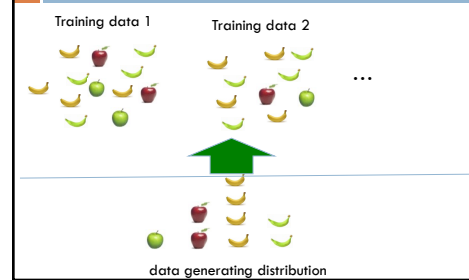
22

data generating distribution

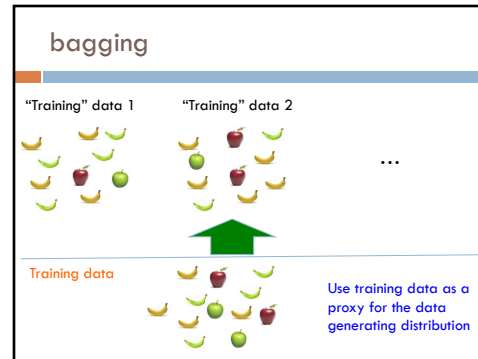


23

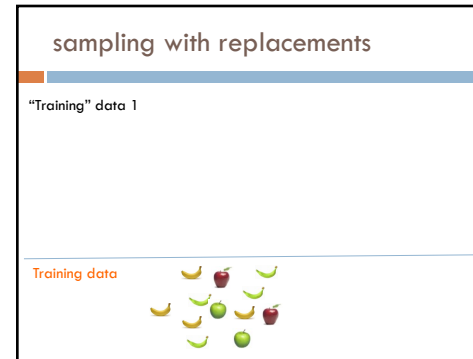
Ideal situation



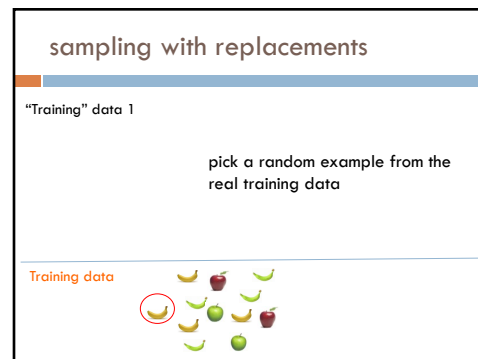
24



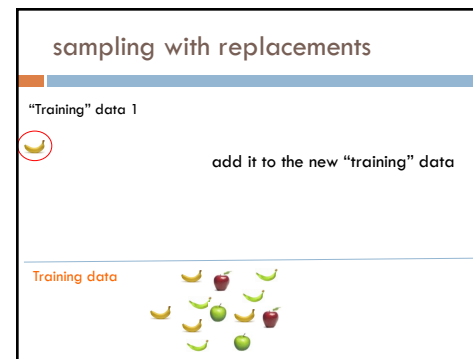
25



26



27



28

sampling with replacements

"Training" data 1

put it back (i.e. leave it) in the original training data

Training data

29

sampling with replacements

"Training" data 1

pick another random example

Training data

30

sampling with replacements

"Training" data 1

pick another random example

Training data

31

sampling with replacements

"Training" data 1

keep going until you've created a new "training" data set

Training data

32

bagging

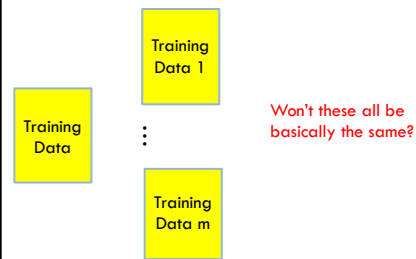
create m "new" training data sets by sampling with replacement from the original training data set (called m "bootstrap" samples)

train a classifier on each of these data sets

to classify, take the majority vote from the m classifiers

33

bagging concerns



34

bagging concerns

For a data set of size n , what is the probability that a given example will **NOT** be select in a "new" training set sampled from the original?

Training data



35

bagging concerns

What is the probability it isn't chosen the first time?

$$1 - 1/n$$

Training data



36

bagging concerns

What is the probability it isn't chosen the **any** of the n times?

$$(1 - 1/n)^n$$

Each draw is independent and has the same probability

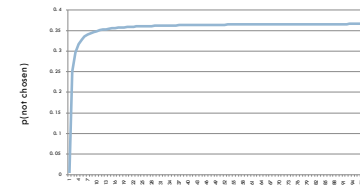
Training data



37

probability of overlap

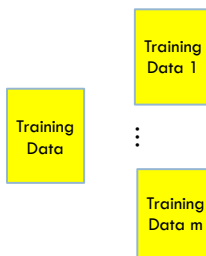
$$(1 - 1/n)^n$$



Converges very quickly to $1 - 1/e \approx 37\%$

38

bagging overlap



Won't these all be basically the same?

On average, a randomly sampled data set will only contain 63% of the examples in the original

39

When does bagging work

Let's say 10% of our examples are noisy (i.e. don't provide good information)

For each of the "new" data set, what proportion of noisy examples will they have?

- They'll still have ~10% of the examples as noisy
- However, these examples will only represent about two-thirds of the original noisy examples

For some classifiers that have trouble with noisy classifiers, this can help

40

When does bagging work

Bagging tends to reduce the **variance** of the classifier

By voting, the classifiers are more robust to noisy examples

Bagging is most useful for classifiers that are:

- Unstable: small changes in the training set produce very different models
- Prone to overfitting

Often has similar effect to regularization

41

Idea 4: boosting

training data

"training" data 2

"training" data 3

Data	Label	Weight	Data	Label	Weight	Data	Label	Weight
	0	0.2		0	0.1		0	0.05
	0	0.2		0	0.1		0	0.2
	1	0.2		1	0.4		1	0.2
	1	0.2		1	0.1		1	0.05
	0	0.2		0	0.3		0	0.5

42

"Strong" learner

Given

- a reasonable amount of training data
- a target error rate ϵ
- a failure probability p

A **strong learning algorithm** will produce a classifier with error rate $< \epsilon$ with probability $1-p$



43

"Weak" learner

Given

- a reasonable amount of training data
- a failure probability p

A **weak learning algorithm** will produce a classifier with error rate < 0.5 with probability $1-p$

Weak learners are much easier to create!



44

weak learners for boosting

Data	Label	Weight
	0	0.2
	0	0.2
	1	0.2
	1	0.2
	0	0.2

 weak learning algorithm  weak classifier

Which of our algorithms can handle weights?

Need a weak learning algorithm that can handle **weighted** examples

45

boosting: basic algorithm

Training:

start with equal example weights

for some number of iterations:

- learn a weak classifier and save
- change the example weights

Classify:

- get prediction from all learned weak classifiers
- weighted vote based on how well the weak classifier did when it was trained (i.e. in relation to training error)

46

boosting basics

Start with equal weighted examples

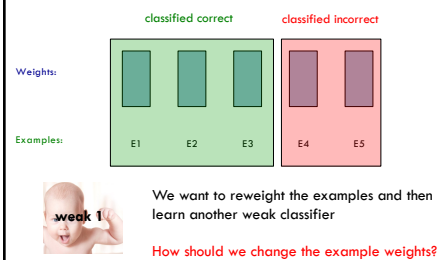
Weights:					
					
Examples:	E1	E2	E3	E4	E5

Learn a weak classifier:

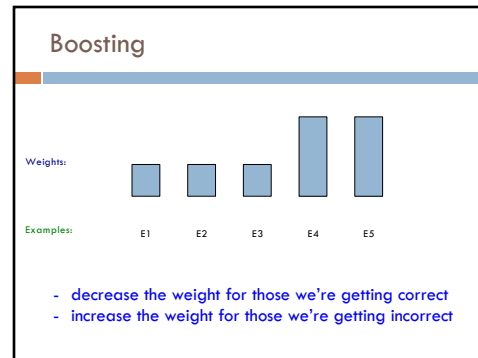


47

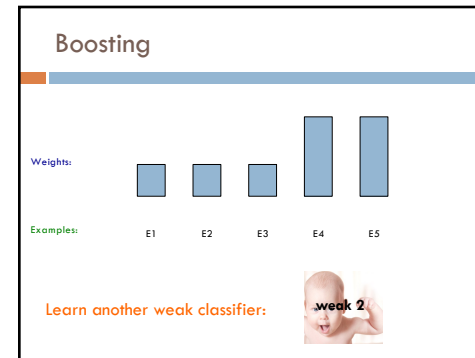
Boosting



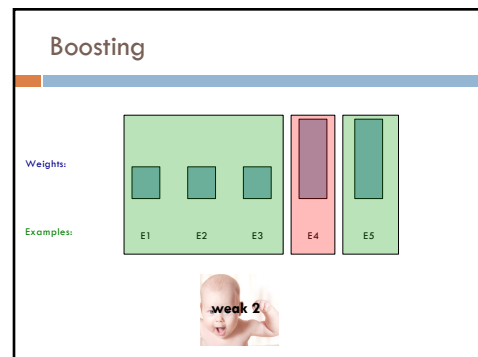
48



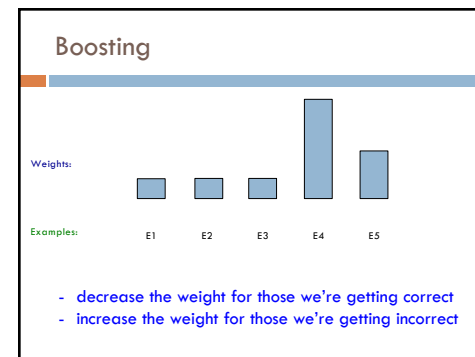
49



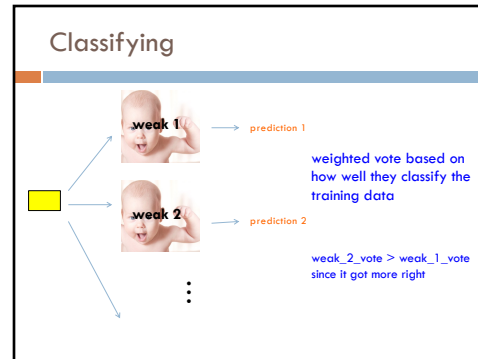
50



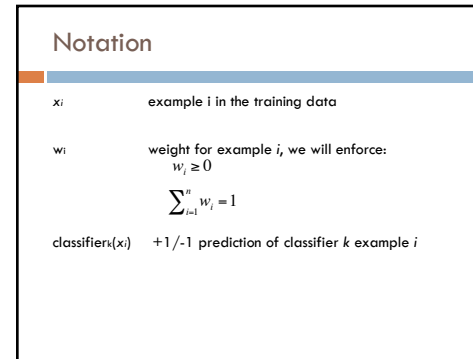
51



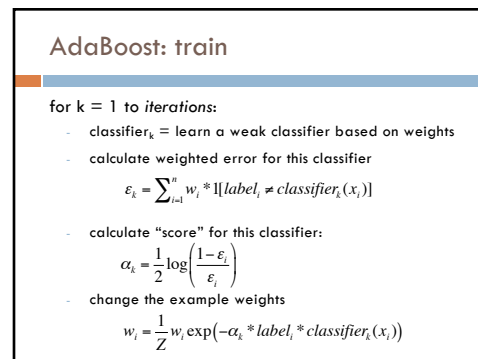
52



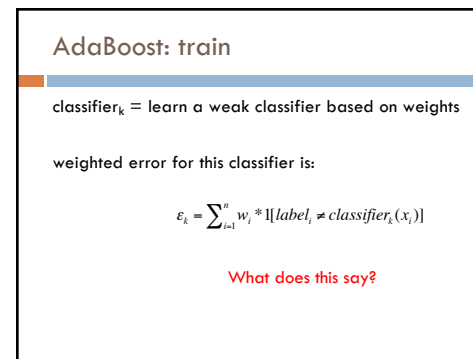
53



54



55



56

AdaBoost: train

classifier_k = learn a weak classifier based on weights

weighted error for this classifier is:

$$\varepsilon_k = \sum_{i=1}^n w_i * \underbrace{1[\text{label}_i \neq \text{prediction}]}_{\text{did we get the example wrong}}$$

What is the range of possible values?

weighted sum of the errors/mistakes

57

AdaBoost: train

classifier_k = learn a weak classifier based on weights

weighted error for this classifier is:

$$\varepsilon_k = \sum_{i=1}^n w_i * \underbrace{1[\text{label}_i \neq \text{prediction}]}_{\text{did we get the example wrong}}$$

Between 0 (if we get all examples right) and 1 (if we get them all wrong)

weighted sum of the errors/mistakes

58

AdaBoost: train

classifier_k = learn a weak classifier based on weights

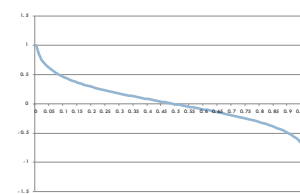
"score" or weight for this classifier is:

$$\alpha_k = \frac{1}{2} \log \left(\frac{1 - \varepsilon_i}{\varepsilon_i} \right)$$

What does this look like (specifically for errors between 0 and 1)?

59

AdaBoost: train



- ranges from $+\infty$ to $-\infty$
- for most reasonable values: ranges from 1 to -1
- errors of 50% = 0
- error < 50% = positive error > 50% = negative

60

AdaBoost: classify

$$\text{classify}(x) = \text{sign} \left(\sum_{k=1}^{\text{iterations}} \alpha_k * \text{classifier}_k(x) \right)$$

What does this do?

61

AdaBoost: classify

$$\text{classify}(x) = \text{sign} \left(\sum_{k=1}^{\text{iterations}} \alpha_k * \text{classifier}_k(x) \right)$$

The weighted vote of the learned classifiers
weighted by α (remember α generally varies
from 1 to -1 training error)

What happens if a classifier has error >50%

62

AdaBoost: classify

$$\text{classify}(x) = \text{sign} \left(\sum_{k=1}^{\text{iterations}} \alpha_k * \text{classifier}_k(x) \right)$$

The weighted vote of the learned classifiers
weighted by α (remember α generally varies
from 1 to -1 training error)

We vote the opposite!

63

AdaBoost: train, updating the weights

update the example weights

$$w_i = \frac{1}{Z} w_i \exp(-\text{label}_i * \alpha_k * \text{classifier}_k(x_i))$$

Remember, we want to enforce:

$$w_i \geq 0$$

$$\sum_{i=1}^n w_i = 1$$

Z is called the **normalizing constant**. It is used
to make sure that the weights sum to 1

What should it be?

64

AdaBoost: train

update the example weights

$$w_i = \frac{1}{Z} w_i \exp(-\text{label}_i * \alpha_k * \text{classifier}_k(x_i))$$

Remember, we want to enforce:

$$w_i \geq 0$$

$$\sum_{i=1}^n w_i = 1$$

normalizing constant (i.e. the sum of the “new” w_i):

$$Z = \sum_{i=1}^n w_i \exp(-\alpha_k * \text{label}_i * \text{classifier}_k(x_i))$$

65

AdaBoost: train

update the example weights

$$w_i = \frac{1}{Z} w_i \exp(-\text{label}_i * \alpha_k * \text{classifier}_k(x_i))$$

What does this do?

66

AdaBoost: train

update the example weights

$$w_i = \frac{1}{Z} w_i \exp(-\text{label}_i * \alpha_k * \text{classifier}_k(x_i))$$

correct → positive
incorrect → negative

correct ?
incorrect

67

AdaBoost: train

update the example weights

$$w_i = \frac{1}{Z} w_i \exp(-\text{label}_i * \alpha_k * \text{classifier}_k(x_i))$$

correct → positive
incorrect → negative

correct → small value
incorrect → large value

Note: only change weights based on current classifier (not all previous classifiers)

68

AdaBoost justification

update the example weights

$$w_i = \frac{1}{Z} w_i \exp(-\text{label}_i * \alpha_k * \text{classifier}_k(x_i))$$

Does this look like anything we've seen before?

70

AdaBoost justification

update the example weights

$$w_i = \frac{1}{Z} w_i \exp(-\text{label}_i * \alpha_k * \text{classifier}_k(x_i))$$

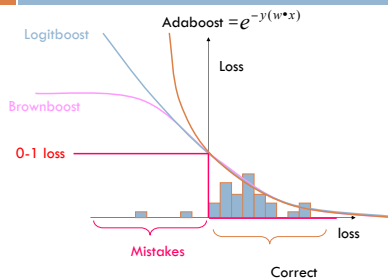
Exponential loss!

$$l(y, y') = \exp(-yy')$$

AdaBoost turns out to be another approach for minimizing the exponential loss!

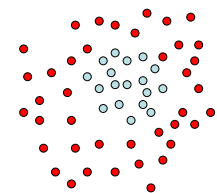
71

Other boosting variants



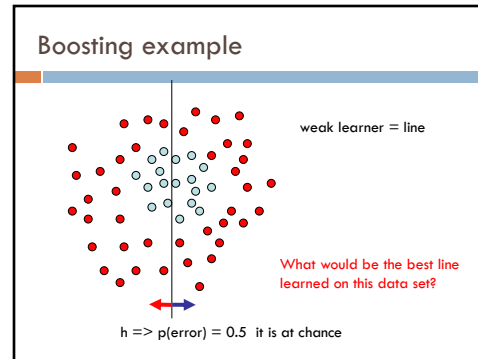
72

Boosting example

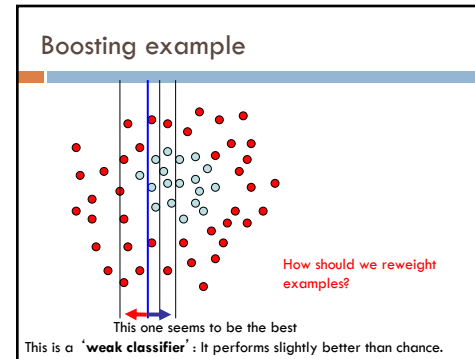


Start with equal weighted data set

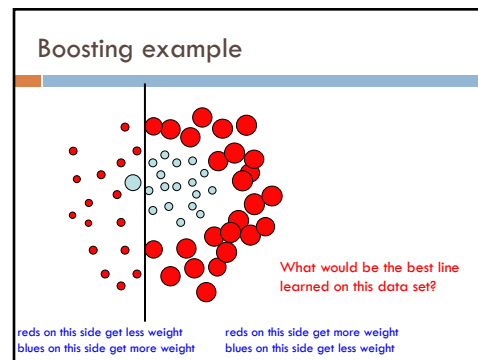
73



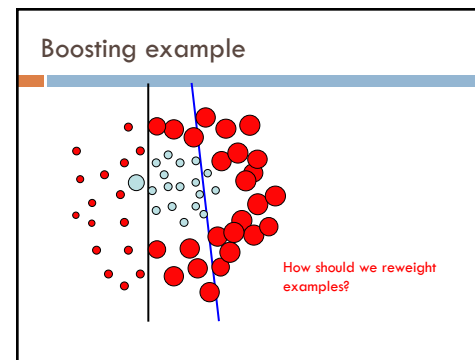
74



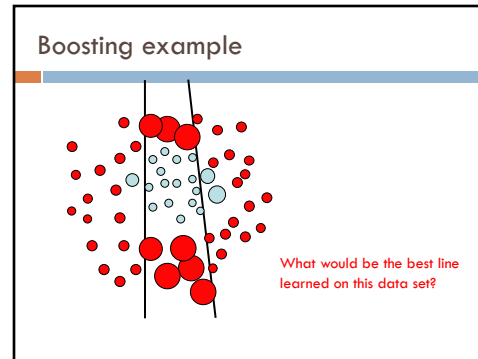
75



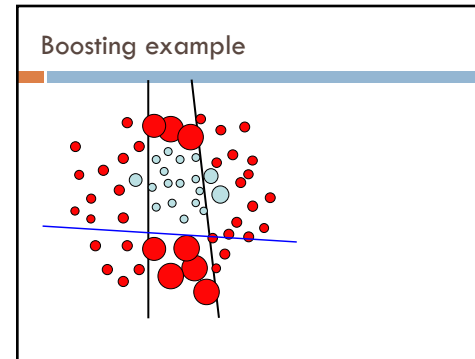
76



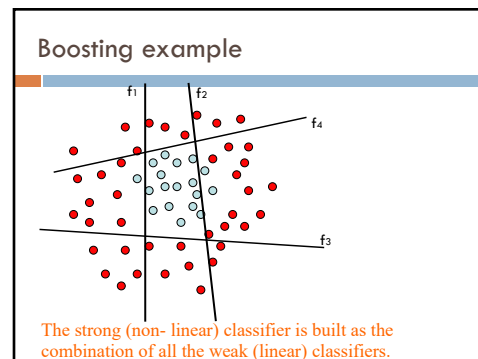
77



78



79



80

AdaBoost: train

for $k = 1$ to *iterations*:

- $\text{classifier}_k = \text{learn a weak classifier based on weights}$
- weighted error for this classifier is:
- "score" or weight for this classifier is:
- change the example weights

What can we use as a classifier?

81

AdaBoost: train

for $k = 1$ to iterations:

- $\text{classifier}_k = \text{learn a weak classifier based on weights}$
- weighted error for this classifier is:
- "score" or weight for this classifier is:
- change the example weights
- Anything that can train on weighted examples
- For most applications, must be fast!

Why?

82

AdaBoost: train

for $k = 1$ to iterations:

- $\text{classifier}_k = \text{learn a weak classifier based on weights}$
- weighted error for this classifier is:
- "score" or weight for this classifier is:
- change the example weights
- Anything that can train on weighted examples
- For most applications, must be fast!
- Each iteration we have to train a new classifier

83

Boosted decision stumps

One of the most common classifiers to use is a decision tree:

- can use a shallow (2-3 level tree)
- even more common is a 1-level tree
- called a **decision stump** 😊
- asks a question about a single feature

What does the decision boundary look like for a decision stump?

84

Boosted decision stumps

One of the most common classifiers to use is a decision tree:

- can use a shallow (2-3 level tree)
- even more common is a 1-level tree
- called a **decision stump** 😊
- asks a question about a single feature

What does the decision boundary look like for boosted decision stumps?

85

Boosted decision stumps

One of the most common classifiers to use is a decision tree:

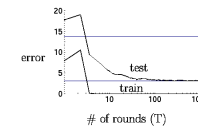
- can use a shallow (2-3 level tree)
- even more common is a 1-level tree
 - called a **decision stump** 😊
 - asks a question about a single feature
- **Linear classifier!**
- Each stump defines the weight for that dimension
 - If you learn multiple stumps for that dimension then it's the weighted average

86

Boosting in practice

Very successful on a wide range of problems

One of the keys is that boosting tends not to overfit, even for a large number of iterations



Using <10,000 training examples can fit >2,000,000 parameters!

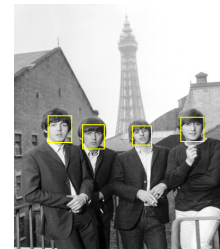
87

Adaboost application example: face detection



88

Adaboost application example: face detection



89

Rapid Object Detection using a Boosted Cascade of Simple Features

Paul Viola
viola@merl.com
Mitsubishi Electric Research Labs
201 Broadway, 8th FL
Cambridge, MA 02139

Michael Jones
mjones@crl.dec.com
Compaq CRL
One Cambridge Center
Cambridge, MA 02142

Rapid object detection using a boosted cascade of simple features
P. Viola, M. Jones - ... on computer vision and pattern recognition ... 2001 - ieeeexplore.ieee.org
... robust and extremely **rapid object detection**. This framework is demonstrated on, and in part motivated by, the task of face **detection**. Toward this end we have constructed a frontal face ...
This set consists of 130 images with 507 labeled frontal faces. A ...
Cited by 8422 Related articles All 129 versions Cite Save More

90

Rapid object detection using a boosted cascade of simple features

P. Viola, M. Jones - ... on computer vision and pattern recognition ... 2001 - ieeeexplore.ieee.org
... robust and extremely **rapid object detection**. This framework is demonstrated on, and in part motivated by, the task of face **detection**. Toward this end we have constructed a frontal face ...
☆ Save 95 Cite Cited by 25943 Related articles All 88 versions

To give you some context of importance:



The anatomy of a large-scale hypertextual web search engine

S. Brin, L. Page - Computer networks and ISDN systems, 1998 - Elsevier
... of a **large scale** search engine which makes heavy use of the **structure** present in **hypertext**.
... The prototype with a full text and **hyperlink** database of at least 24 million pages is available ...
☆ Save 97 Cite Cited by 2553 Related articles All 189 versions Web of Science: 7372

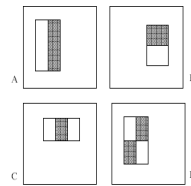
Or: Modeling word burstiness using the Dirichlet distribution

R.E. Mullen, D. Kuzubak, C. Elkan
Proceedings of the 22nd International conference on Machine learning, 2005 - aclm.org
Multinomial distributions are often used to model text documents. However, they do not capture well the phenomenon that words in a document tend to appear in bursts: if a word appears once, it is more likely to appear again. In this paper, we propose the Dirichlet compound multinomial model (DCM) as an alternative to the multinomial. The DCM model has one additional degree of freedom, which allows it to capture burstiness. We show experimentally that the DCM is substantially better than the multinomial at modeling text.
SHOW ABSTRACT

☆ Save 82 Cite Cited by 458 Related articles All 22 versions

91

"weak" learners



4 Types of "Rectangle filters"
(Similar to Haar wavelets
Papageorgiou, et al.)

Based on 24x24 grid:
160,000 features to choose from

$$g(x) = \text{sum(WhiteArea)} - \text{sum(BlackArea)}$$

92

"weak" learners

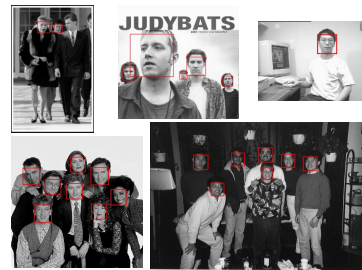


$$F(x) = \alpha_1 f_1(x) + \alpha_2 f_2(x) + \dots$$

$$f_i(x) = \begin{cases} 1 & \text{if } g_i(x) > \theta_i \\ -1 & \text{otherwise} \end{cases}$$

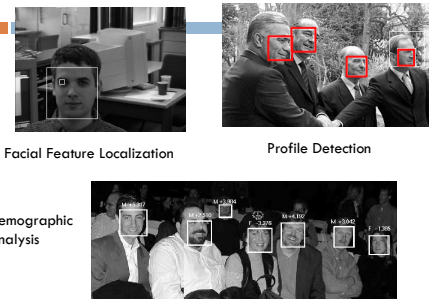
93

Example output



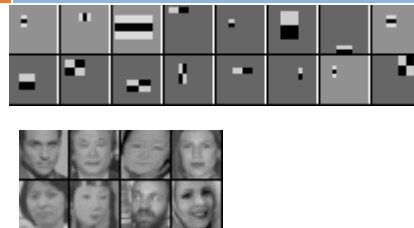
94

Solving other “Face” Tasks



95

“weak” classifiers learned



96

Bagging vs Boosting

Journal of Artificial Intelligence Research 11 (1999) 149-158

Submitted 1/99, published 8/99

Popular Ensemble Methods: An Empirical Study

David Oplis
Department of Computer Science
University of Montana
Missoula, MT 59812 USA

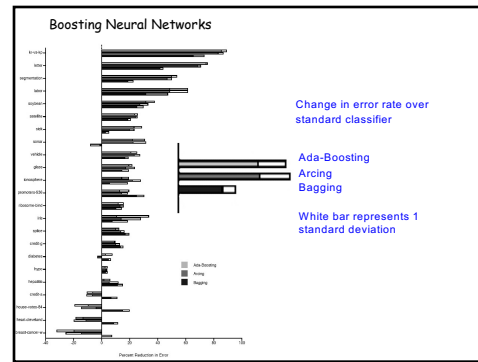
OPITZ@CS.UMT.EDU

Richard MacLlin
Computer Science Department
University of Minnesota
Duluth, MN 55812 USA

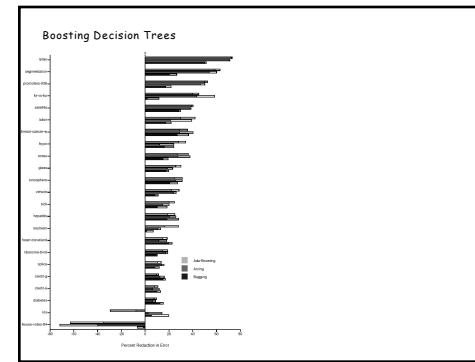
RMACLIN@D.UMN.EDU

<http://arxiv.org/pdf/1106.0257.pdf>

97



98



99