

Last updated: Nov 26, 2012

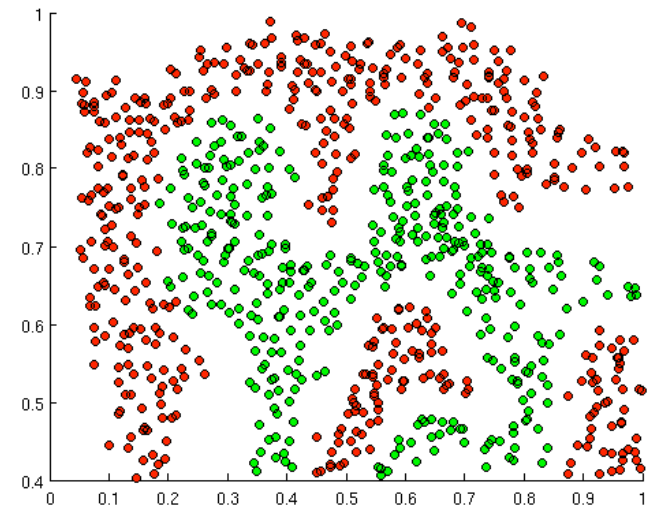
BOOSTING

Limitations of Classifiers

2

Boosting

- All of the classifiers we have studied have limitations: if the problem is ‘interesting’, they don’t tend to perform perfectly.
- Often these limitations stem from the inability of a single classifier to mold a decision surface into exactly the right shape.
- As a consequence, the classifier may perform well for part of the input space, but poorly in another part of the input space.
- How can we overcome this problem?



Combining Classifiers

3

Boosting

- One natural idea is to try to combine multiple classifiers, each of which captures some aspect of the problem.
- Together, the ensemble of classifiers may be able to model complex decision surfaces.

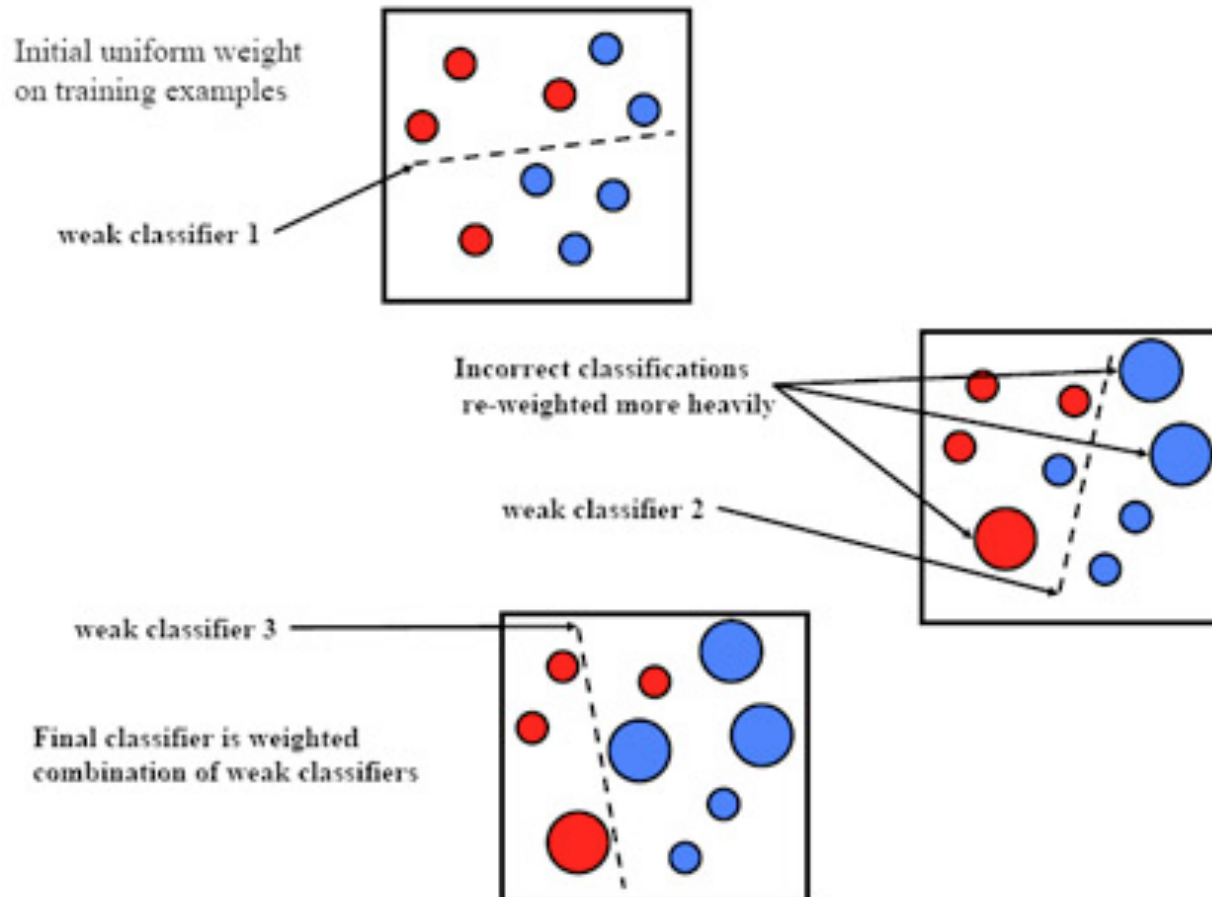
Boosting

- Boosting is one method for combining multiple ‘weak’ classifiers into a single ‘strong’ classifier.
- The approach is greedy. We begin by selecting the best weak classifier (i.e., the one with lowest error on the training dataset).
- We then repeat the following 3 steps until no improvement can be made or an error criterion has been reached:
 1. Reweight the input vectors, up-weighting those that were classified incorrectly by the last classifier selected.
 2. Again select the best weak classifier, on the reweighted training data.
 3. Linearly combine this classifier with the classifiers already selected, with an appropriate weighting.

Boosting: Example

5

Boosting



$$H(x) = \text{sign}(\alpha_1 h_1(x) + \alpha_2 h_2(x) + \alpha_3 h_3(x))$$

Adaboost

Yoav Freund and Robert E. Schapire (1996). Experiments with a new boosting algorithm. *Machine Learning: Proceedings of the Thirteenth International Conference*, 148–156.



Yoav Freund



Robert Schapire

AdaBoost

7

Boosting

- AdaBoost (**Ad**aptive **Boo**sting) is probably the most popular form of boosting.

Let $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N)$ be the training data, with $y_i \in \{-1, 1\}$.

- We seek an optimal classifier of the form

$$F(\mathbf{x}) = \sum_{k=1}^K \alpha_k \phi(\mathbf{x}; \theta_k), \quad \alpha_k > 0 \forall k.$$

Here, each $\phi(\mathbf{x}; \theta_k)$ is a weak classifier that evaluates to -1 or +1, and input vectors $\mathbf{x}$ are classified according to the sign of $F(\mathbf{x})$.

The parameters α_k and θ_k are selected at each stage to minimize

$$\sum_{i=1}^N \exp(-y_i F(\mathbf{x}_i)).$$

AdaBoost: Learning the Parameters

$$F(\mathbf{x}) = \sum_{k=1}^K \alpha_k \phi(\mathbf{x}; \theta_k)$$

The parameters α_k and θ_k are selected at each stage to minimize

$$\sum_{i=1}^N \exp(-y_i F(\mathbf{x}_i)).$$

Note that learning is greedy:

at each stage k only the parameters α_k and θ_k are learned.

Parameters at previous stages $1 \dots k-1$ remain fixed.

At stage k we can thus define the problem as finding the parameters α_k and θ_k that minimize

$$\mathcal{J}(\alpha, \theta) = \sum_{i=1}^N \exp(-y_i (F_{k-1}(\mathbf{x}_i) + \alpha \phi(\mathbf{x}_i; \theta))).$$

AdaBoost: Learning the Parameters

$$\begin{aligned}\mathcal{J}(\alpha, \theta) &= \sum_{i=1}^N \exp\left(-y_i \left(F_{k-1}(\mathbf{x}_i) + \alpha \phi(\mathbf{x}_i; \theta)\right)\right) \\ &= \sum_{i=1}^N w_i^{(k)} \exp\left(-y_i \alpha \phi(\mathbf{x}_i; \theta)\right)\end{aligned}$$

where

$$w_i^{(k)} \triangleq \exp\left(-y_i F_{k-1}(\mathbf{x}_i)\right)$$

AdaBoost: Learning the Parameters

$$\mathcal{J}(\alpha, \theta) = \sum_{i=1}^N w_i^{(k)} \exp(-y_i \alpha \phi(\mathbf{x}_i; \theta))$$

Observe that $\exp(-y_i \alpha \phi(\mathbf{x}_i; \theta))$ evaluates to $e^{-\alpha}$ for input vectors that are correctly classified and $e^{+\alpha}$ for vectors that are incorrectly classified.

Thus the cost function can be re-expressed as

$$\mathcal{J}(\alpha, \theta) = e^{-\alpha} \sum_{i \in I_c(\theta)} w_i^{(k)} + e^{+\alpha} \sum_{i \in I_e(\theta)} w_i^{(k)},$$

where I_c and I_e are the index sets of inputs correctly and incorrectly classified by the k^{th} weak classifier, respectively.

AdaBoost: Learning the Parameters

$$\mathcal{J}(\alpha, \theta) = e^{-\alpha} \sum_{i \in I_c(\theta)} w_i^{(k)} + e^{+\alpha} \sum_{i \in I_e(\theta)} w_i^{(k)},$$

where I_c and I_e are the index sets of inputs correctly and incorrectly classified by the k^{th} weak classifier, respectively.

Thus we seek the parameters that minimize the sum of the weights for the misclassified inputs. Note that this does not depend upon α ! Thus we have:

$$\theta_k = \underset{\theta}{\operatorname{argmin}} \sum_{i \in I_e(\theta)} w_i^{(k)}.$$

AdaBoost: Learning the Parameters

12

Boosting

$$\mathcal{J}(\alpha, \theta) = e^{-\alpha} \sum_{i \in I_c(\theta)} w_i^{(k)} + e^{+\alpha} \sum_{i \in I_e(\theta)} w_i^{(k)},$$

where I_c and I_e are the index sets of inputs correctly and incorrectly classified by the k^{th} weak classifier, respectively.

Once θ_k is chosen, we need to determine α_k .

$$\frac{\partial}{\partial \alpha} \mathcal{J}(\alpha, \theta) = -e^{-\alpha} \sum_{i \in I_c(\theta)} w_i^{(k)} + e^{+\alpha} \sum_{i \in I_e(\theta)} w_i^{(k)} = 0$$

$$\rightarrow e^{2\alpha} = \frac{\sum_{i \in I_c(\theta)} w_i^{(k)}}{\sum_{i \in I_e(\theta)} w_i^{(k)}} \rightarrow \alpha = \frac{1}{2} \log \left(\frac{\sum_{i \in I_c(\theta)} w_i^{(k)}}{\sum_{i \in I_e(\theta)} w_i^{(k)}} \right)$$

AdaBoost: Learning the Parameters

13

Boosting

$$\alpha_k = \frac{1}{2} \log \left(\frac{\sum_{i \in I_c(\theta)} w_i^{(k)}}{\sum_{i \in I_e(\theta)} w_i^{(k)}} \right)$$

We define the **weighted error rate** ε_k as $\varepsilon_k = \frac{\sum_{i \in I_e(\theta)} w_i^{(k)}}{\sum_{i=1}^N w_i^{(k)}}$.

$$\text{Then we have } \alpha_k = \frac{1}{2} \log \left(\frac{1 - \varepsilon_k}{\varepsilon_k} \right)$$

AdaBoost: Learning the Parameters

14

Boosting

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

Note that for the optimal θ , $\varepsilon_k < \frac{1}{2}$.

(Otherwise, we could just flip the sign and have a better weak classifier.)

Thus $\alpha_k > 0$, and is larger for weak classifiers with lower weighted error.

Weights

15

Boosting

$$F(\mathbf{x}) = \sum_{k=1}^K \alpha_k \phi(\mathbf{x}; \theta_k)$$

$$w_i^{(k)} \triangleq \exp(-y_i F_{k-1}(\mathbf{x}_i))$$

$$\rightarrow w_i^{(k+1)} \triangleq \exp(-y_i F_k(\mathbf{x}_i)) = \exp(-y_i F_{k-1}(\mathbf{x}_i)) \exp(-y_i \alpha_k \phi(\mathbf{x}_i, \theta_k))$$

Thus the weight for an input decreases if the current weak classifier classifies it correctly, and increases if it classifies it incorrectly.

The amount of the change depends on the reliability of the weak classifier: If the weak classifier has a relatively low error rate, it changes the weights more aggressively.

Convergence

16

Boosting

- Boosting enjoys a very strong convergence property:
 - ▣ Convergence is not generally monotonic.
 - ▣ However, as the number of iterations $\rightarrow \infty$, error on the training dataset $\rightarrow 0$:

Let the weighted error rate of the k^{th} weak classifier be $\frac{1}{2} - \gamma_k$.

Then it can be shown that after K iterations,

the training error of the strong classifier is at most $\exp\left(-2\sum_{k=1}^K \gamma_k^2\right)$

- Adaboost is found empirically to have good generalization properties:
 - ▣ As the number K of weak classifiers grows, the error on the test set tends to decrease and then level off
 - ▣ Unlike many other methods, Adaboost does not always suffer from overlearning, even for large K .
 - ▣ The error on the test set tends to decrease even after the error on the training set is 0.

Limitations

18

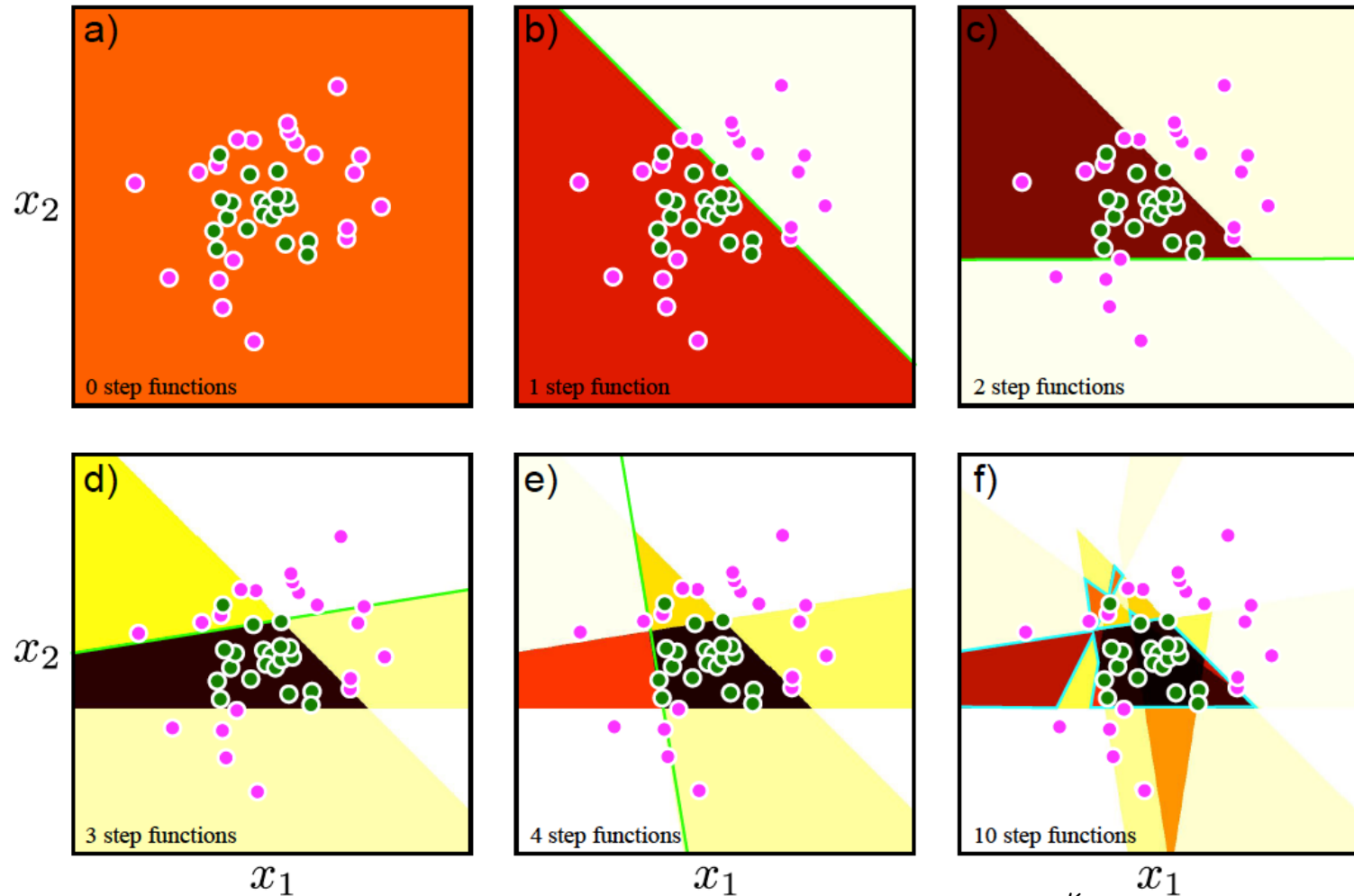
Boosting

- Adaboost is not probabilistic, and thus does not provide a measure of confidence for each classification.
- Can be sensitive to outliers

Example

19

Boosting



NB: There is a biasing term in this example, thus $F(\mathbf{x}) = \alpha_0 + \sum_{k=1}^K \alpha_k \phi(\mathbf{x}; \theta_k)$, $\alpha_k > 0 \forall k$.

□ Example

Boosting in MATLAB

21

Boosting

- MATLAB's Statistics Toolbox provides various boosting algorithms through the function **fitensemble()**.