

Machine Learning in R

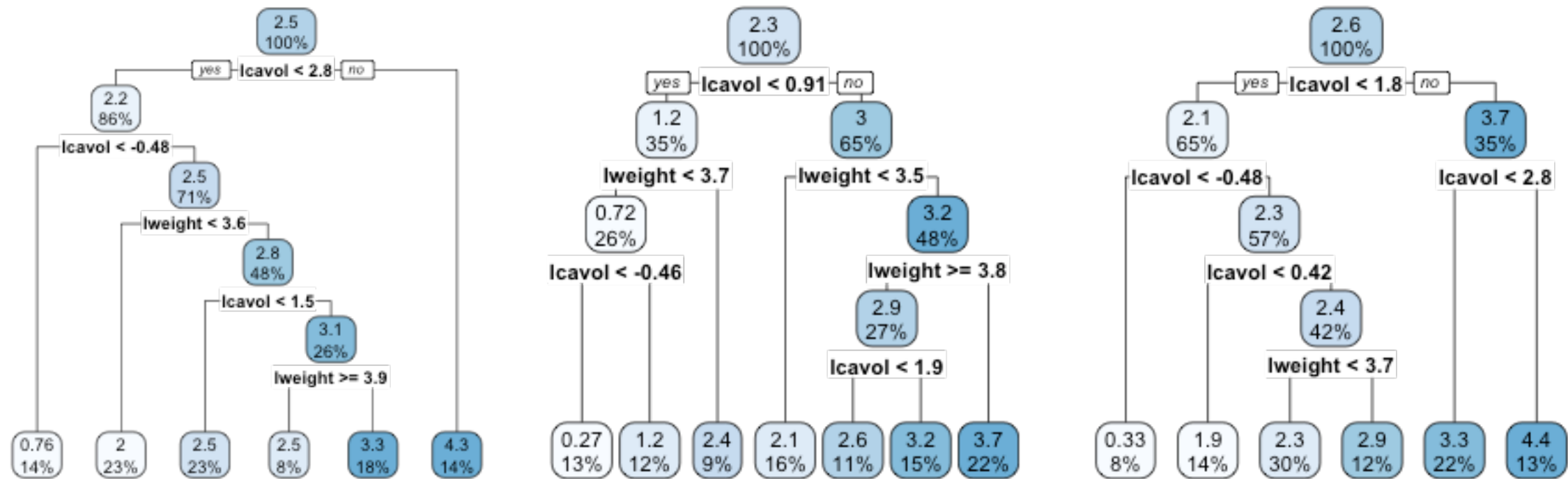
Lecture 6 — Boosting; Things to watch out for when doing supervised learning

Jean Feng — Epidemiology and Biostatistics

Last Lecture

- Decision Trees
- Bagging
- Random Forests
- Variable Importance

Recall: Trees and Forests



- **Classification/regression trees:** This fits a single tree. However, this model is highly sensitive to how the data is sampled and has *high variance*.
- **Random forests:** We can decrease the variance by averaging B trees, where each tree is trained on a bootstrap sample and the predictors are randomly selected for each split.

Today's Lecture

- Boosting
 - AdaBoost
 - Forward Stagewise Additive Modeling
- Gradient boosting
- Issues to think about
 - Confounders
 - Fairness

The hypothesis boosting problem

- Suppose there is an efficient learning algorithm whose performance, on any dataset, is slightly better than random guessing.

“Weak learner”:
an algorithm that does slightly better than random
- Does this imply that there exists an efficient algorithm on any dataset with arbitrarily good prediction accuracy?

“Strong learner”:
an algorithm that achieves arbitrarily good prediction accuracy as the sample size grows
- Can we “boost” a set of weak learners to create a single strong learner?

The answer is yes!

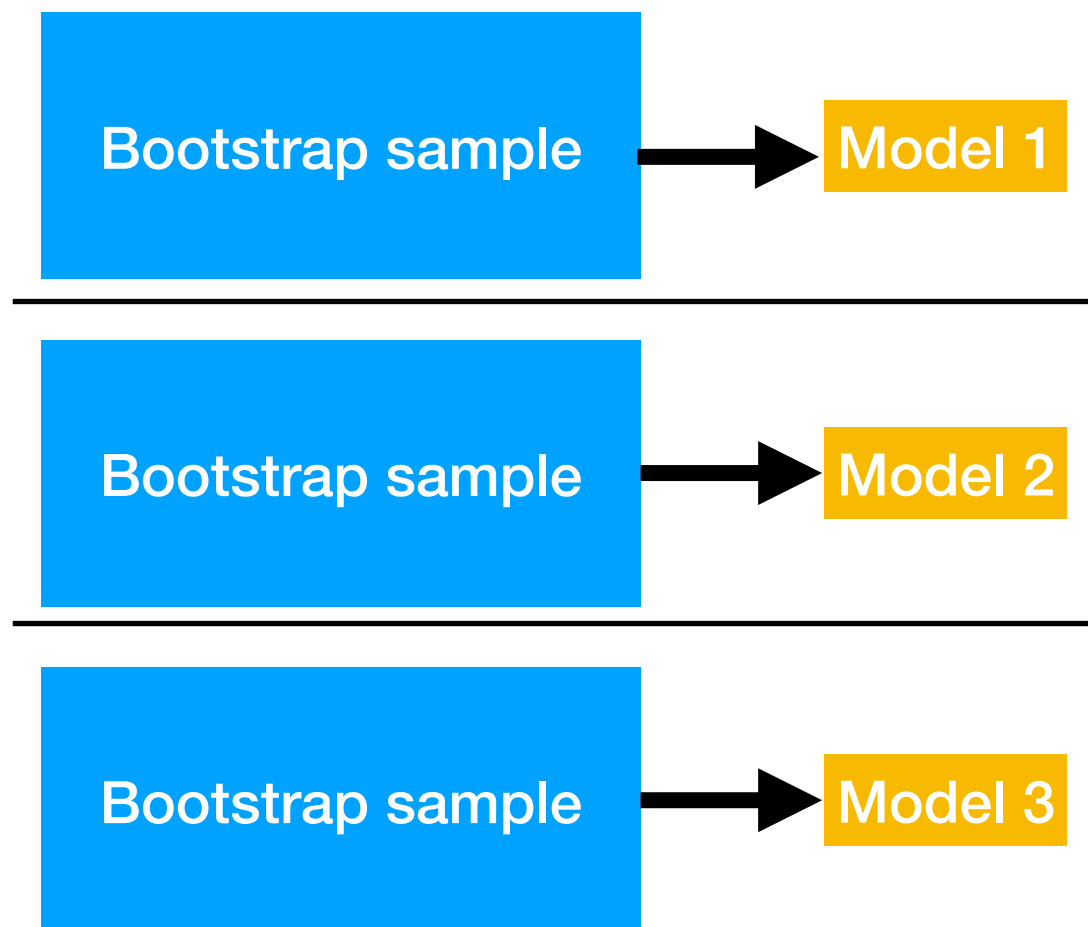
Boosting algorithms

- Boosting is a general-purpose method that builds an ensemble of “weak learners”.
- Like bagging and random forests, boosting averages the predictions from each weak learner to make a prediction.
- Unlike bagging and random forests, boosting trains each weak learner ***sequentially***, which is a fundamentally different approach. Boosting tries to optimize how the weak learners are combined to maximize prediction accuracy.
- Can be used for both regression and classification.

Ensemble methods

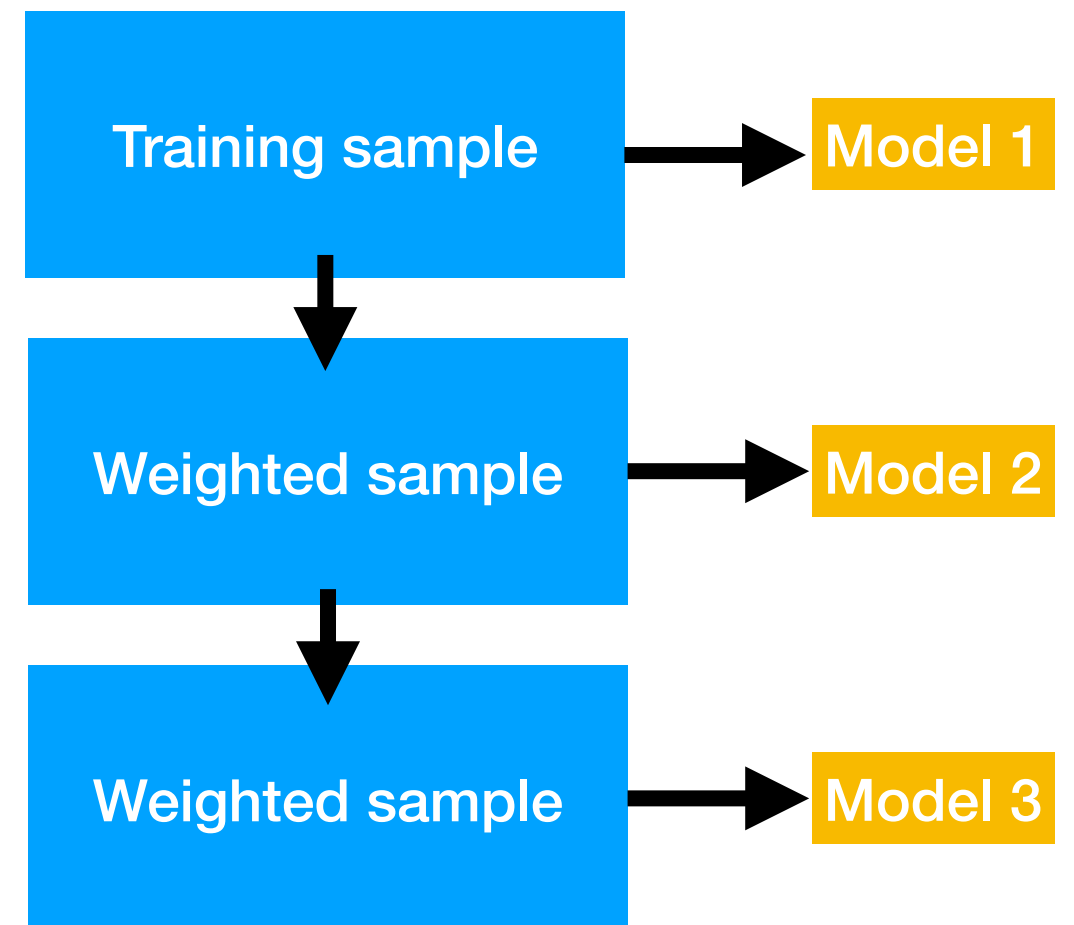
Bagging-type methods:

Models are built independently



Boosting:

Models build off one another

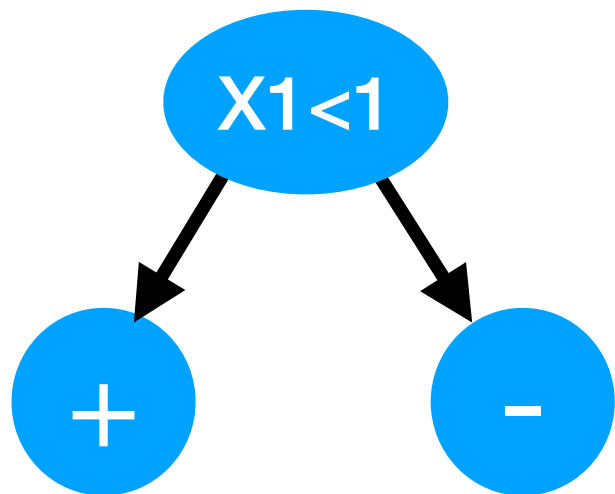


The many variants of Boosting

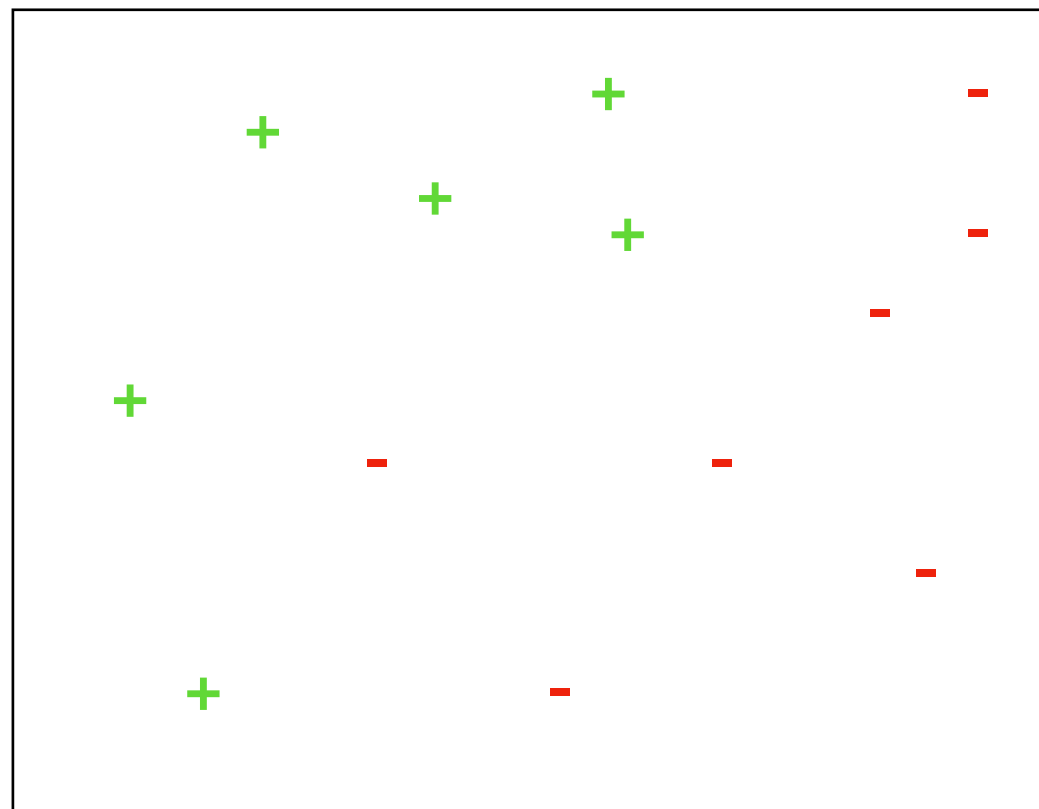
- **Adaptive Boosting (AdaBoost):** The first classification algorithm that adaptively grows a sequence of weak learners
- **Forward Stagewise Additive Modeling:** Generalization of AdaBoost
- **Gradient Boosting:** A gradient-based perspective of Boosting

AdaBoost: the basic idea

- Let's suppose our “weak learners” are decision stumps:

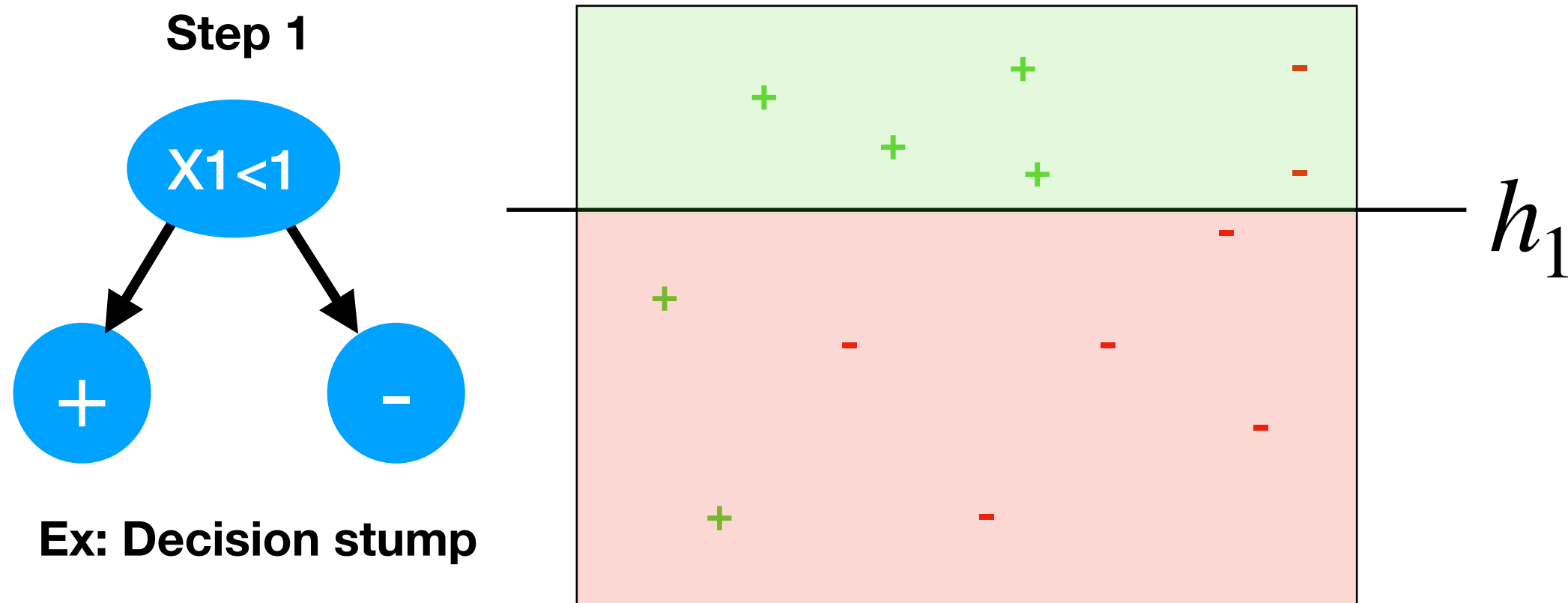


Ex: Decision stump



AdaBoost: the basic idea

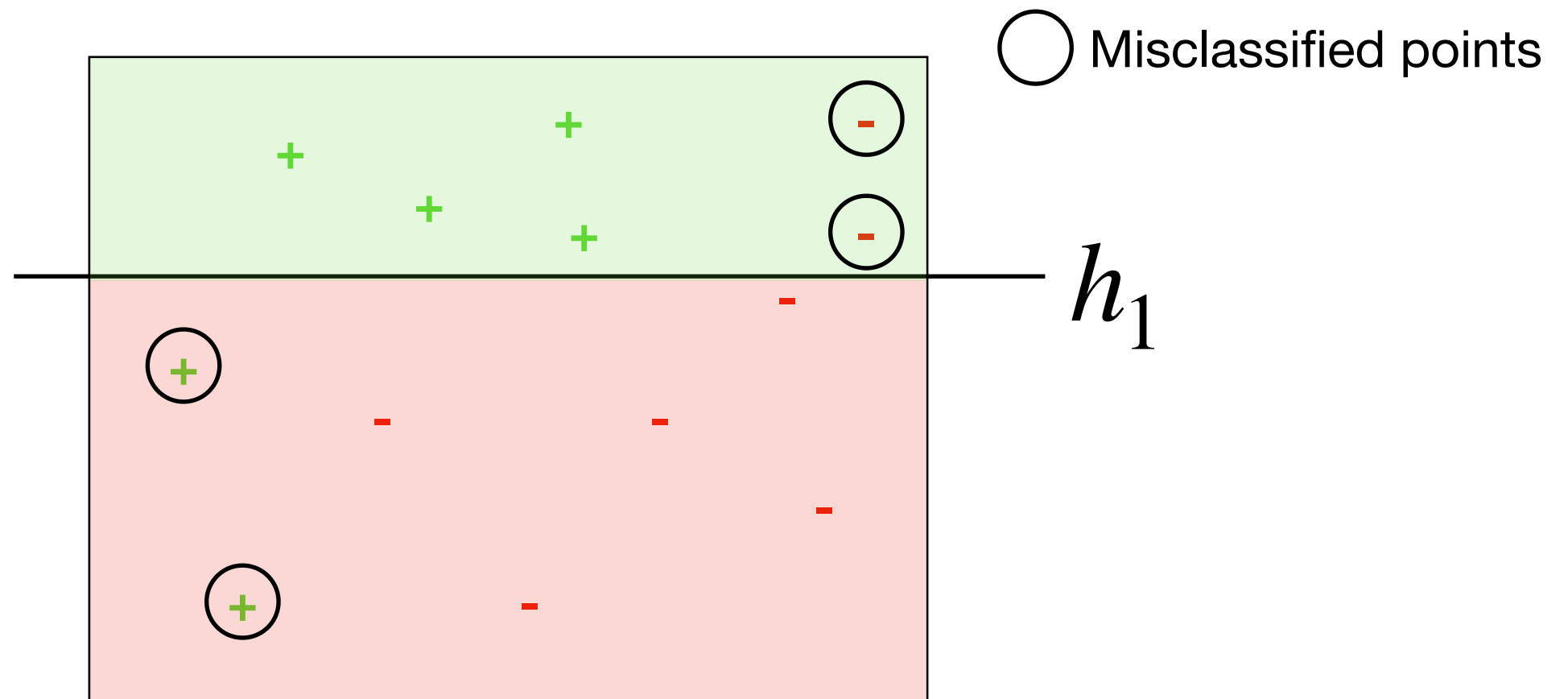
- Let's suppose our “weak learners” are decision stumps:



AdaBoost: the basic idea

- There are a lot of points that are misclassified. Let's try to fix this.

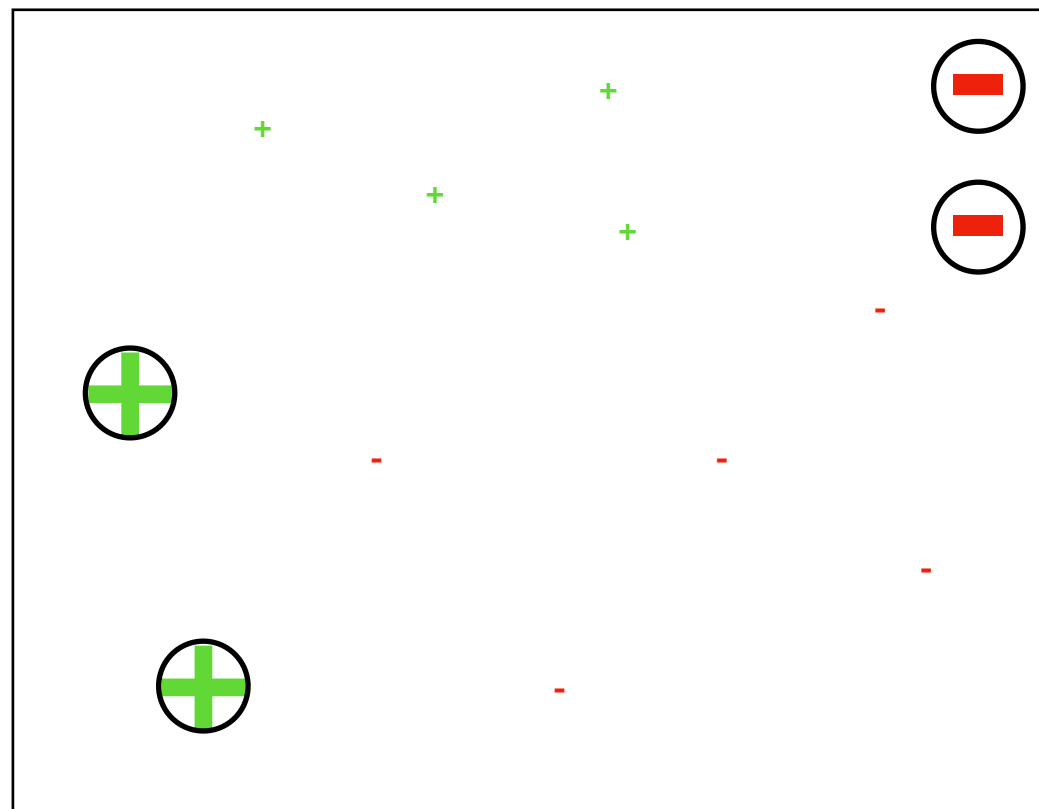
Step 1



AdaBoost: the basic idea

- Let's upweight misclassified points and downweight correctly classified points. Fit a new decision stump on this reweighed dataset.

Step 2

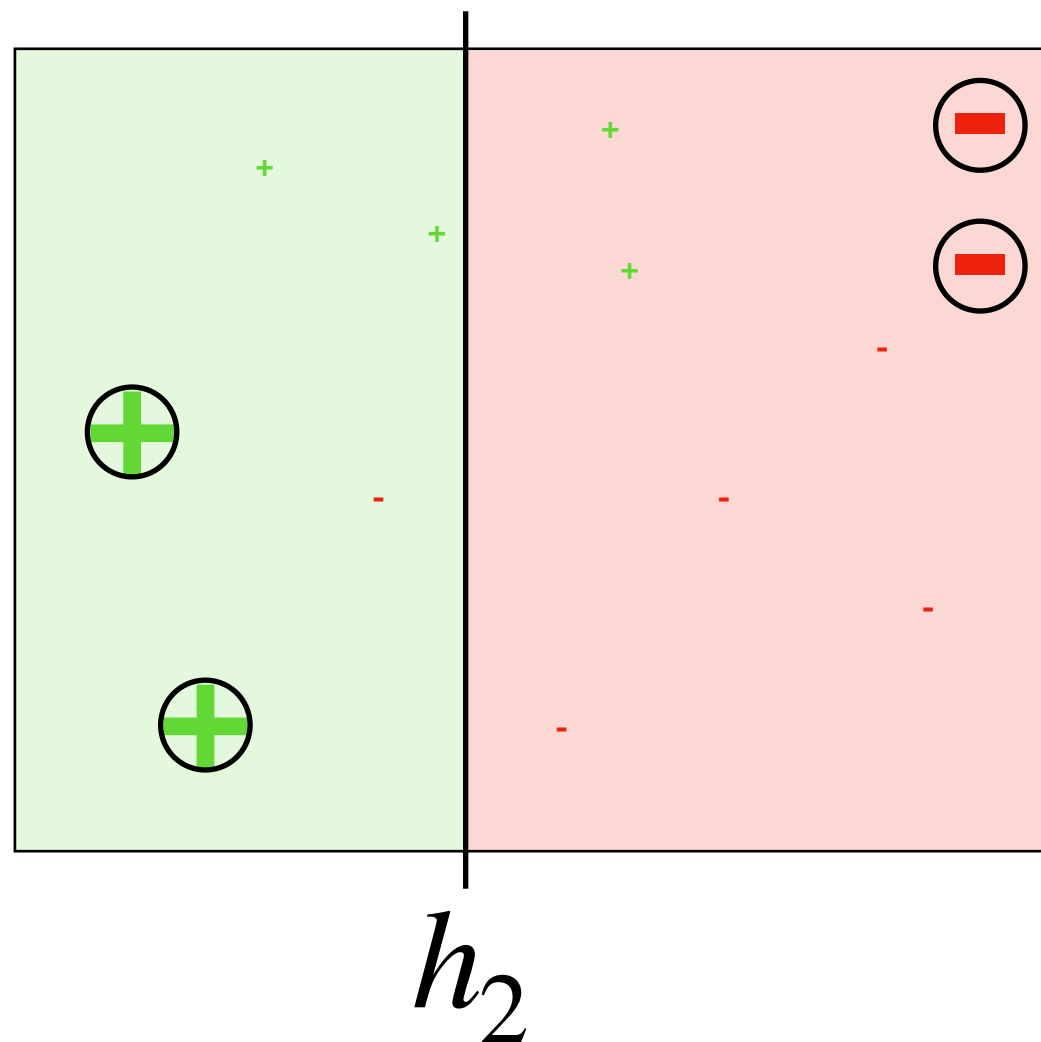


○ Misclassified points

AdaBoost: the basic idea

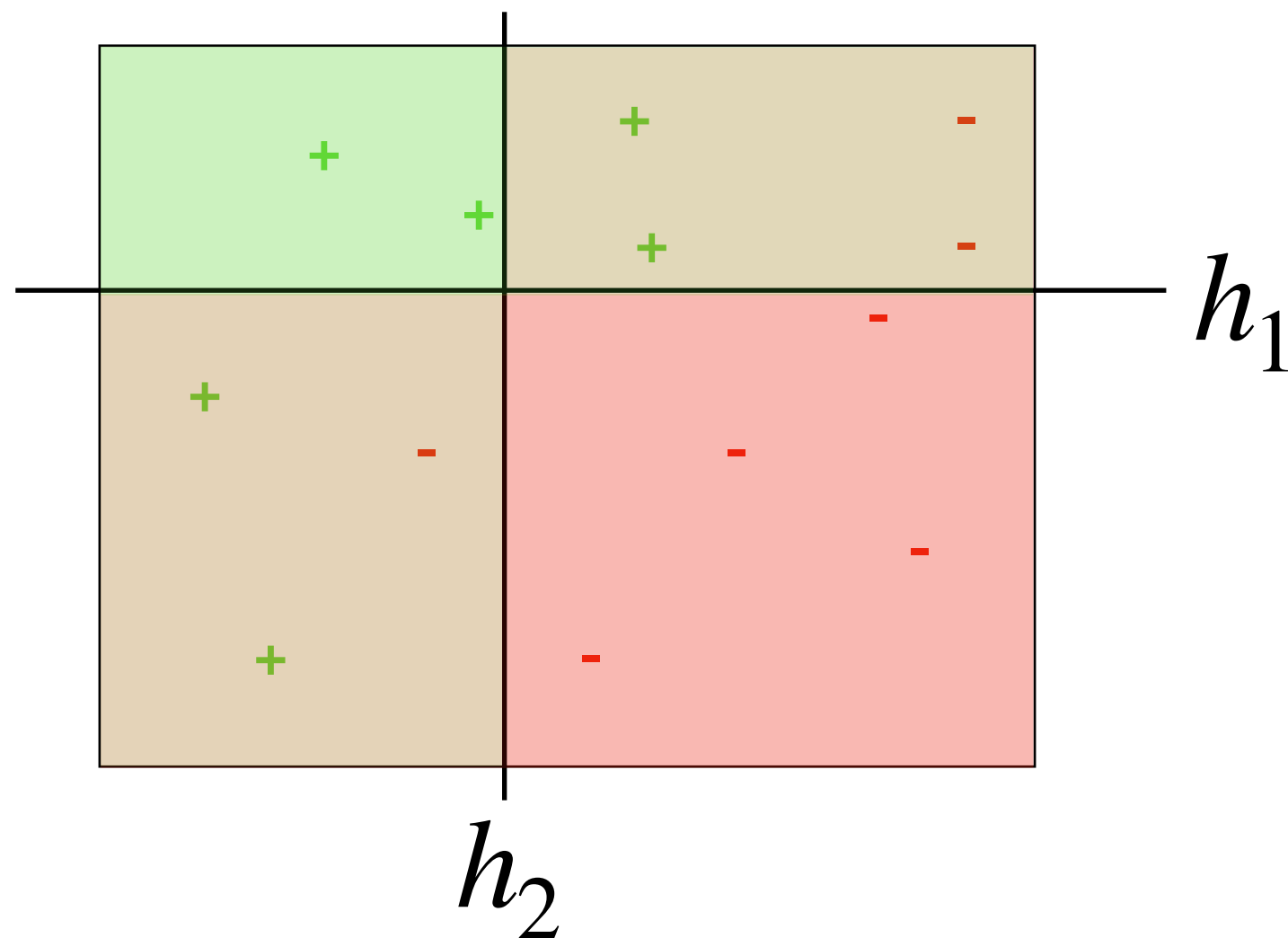
- Based on this reweighed dataset, we fit the following decision stump:

Step 2



AdaBoost: the basic idea

- The combined model looks like:



Moreover, we can take a weighted average of these two models:

$$f(x) = \alpha_1 h_1(x) + \alpha_2 h_2(x)$$

Weights

AdaBoost: general steps

- Start with weighting all observations equally.
- For $b = 1, 2, \dots, B$:
 - Fit a new classifier h_b for weighted dataset.
 - Decide the weight for model h_b , denoted α_b , based on its error on the weighted dataset.
 - Decrease the weight of observations correctly classified by h_b . Increase the weight of observations incorrectly classified by h_b .

- Final model: $f(x) = \text{sign} \left(\sum_{b=1}^B \alpha_b h_b(x) \right)$.

Slowly improving the training error

- Because we are fitting weak learners at each iteration, the training error will decrease ***slowly*** over time.
- This is important, because it gives us very fine control over the bias-variance trade-off.
- A key **tuning parameter** in boosting is the **number of models B** .
 - (Unlike bagging and random forests, the model complexity and the risk of overfitting increases with B .)

Today's Lecture

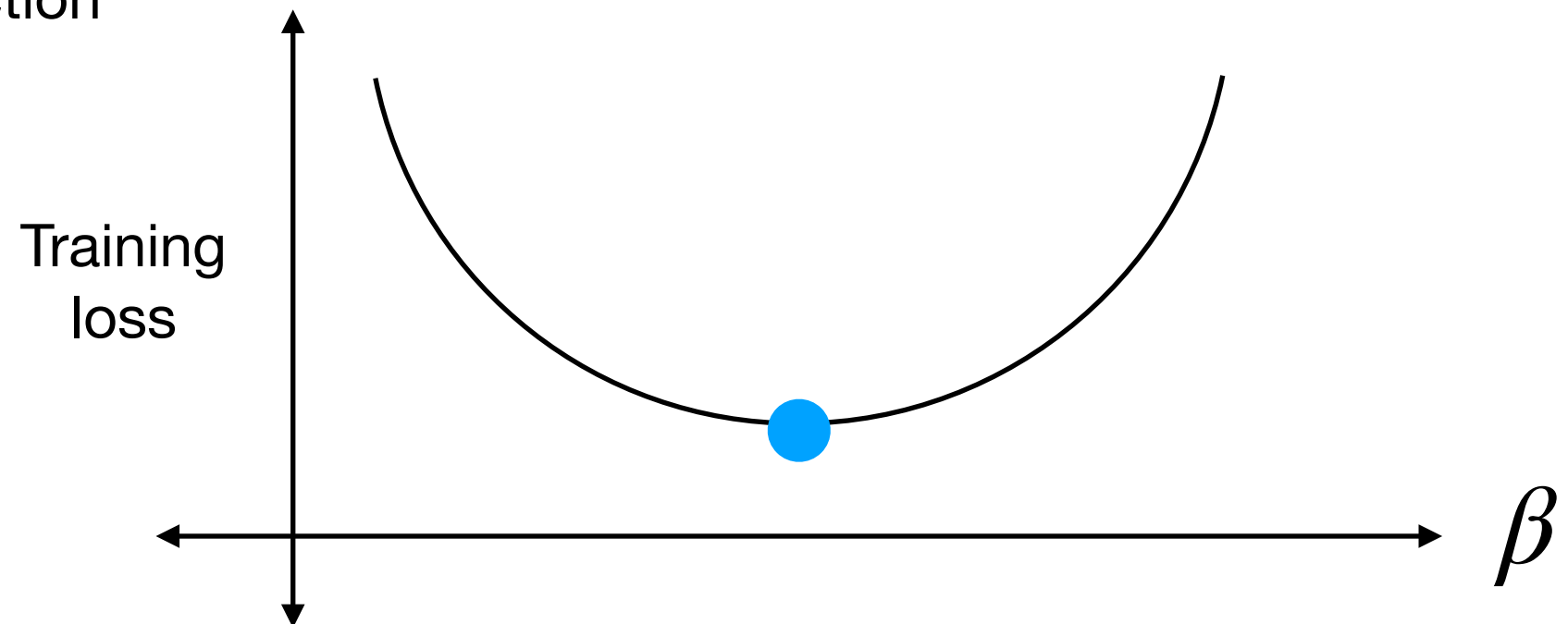
- Boosting
 - AdaBoost
 - Forward Stagewise Additive Modeling
- Gradient boosting
- Issues to think about
 - Confounders
 - Fairness

Boosting as an optimization algorithm

- Recall: Many machine learning algorithms can be defined as the solution to an optimization problem of the form:

$$\min_{\beta \in \mathbb{R}^{p+1}} \frac{1}{n} \sum_{i=1}^n \ell \left(y_i, f_{\beta}(x_i) \right)$$

Loss function



Forward Stagewise Additive Modeling

An additive model has the form $f(x) = \sum_{b=1}^B \alpha_b h_b(x)$.

- For $b = 1, 2, \dots, B$:
 - Find a classifier h_b and weight α_b that minimize the training error

$$\frac{1}{n} \sum_{i=1}^n \ell \left(y_i, \underbrace{\sum_{j=1}^{b-1} \alpha_j h_j(x)}_{\text{Current ensemble}} + \underbrace{\alpha_b h_b(x)}_{\text{New weight \& model}} \right)$$

- Output $f(x) = \sum_{b=1}^B \alpha_b h_b(x)$

Forward Stagewise Additive Modeling

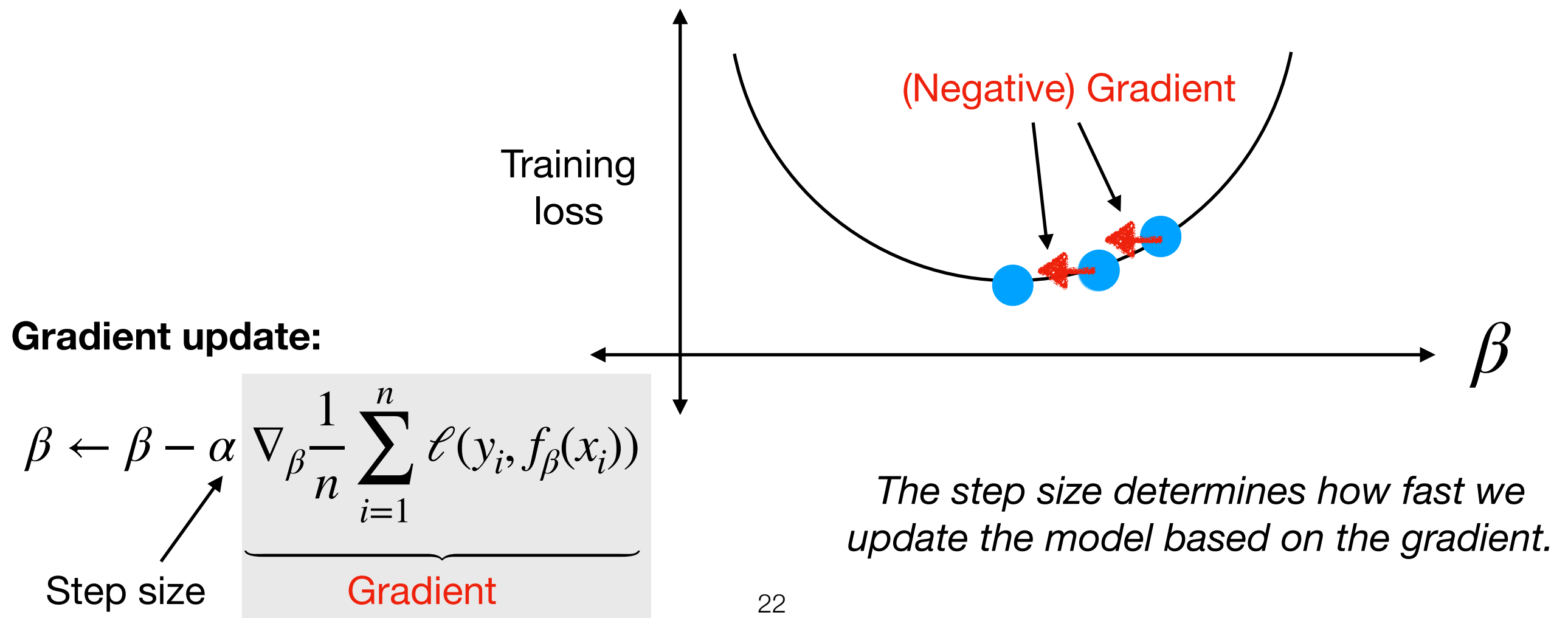
- For special choices of loss functions, this procedure reduces to re-weighting observations at each iteration, as in AdaBoost. This makes it simple to implement, since you simply reweight the data at each step.
- Unfortunately, this doesn't hold in general.
- A more general boosting algorithm uses a **gradient-based approach**.

Today's Lecture

- Boosting
 - AdaBoost
 - Forward Stagewise Additive Modeling
- Gradient boosting
- Issues to think about
 - Confounders
 - Fairness

A gradient-based approach

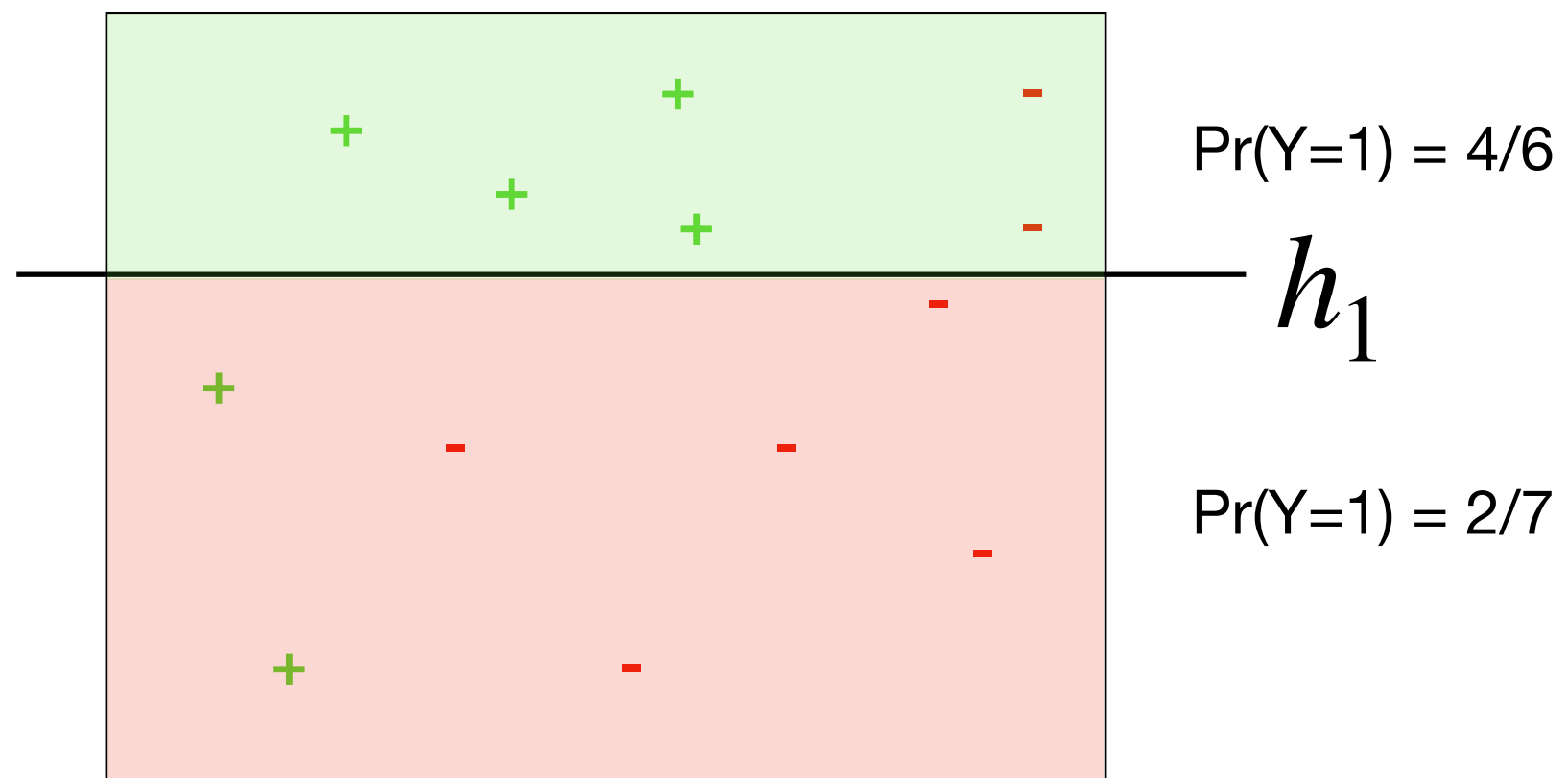
- The **gradient** of a function is the direction that increases the function's value the most. The **negative gradient** is the direction of “**steepest descent**” of the function's value.
- To minimize the training error, many machine learning algorithms compute the gradient and follow the direction of steepest descent.



Gradient-boosting: the basic idea

- Again, let's suppose our “weak learners” are probability decision stumps:

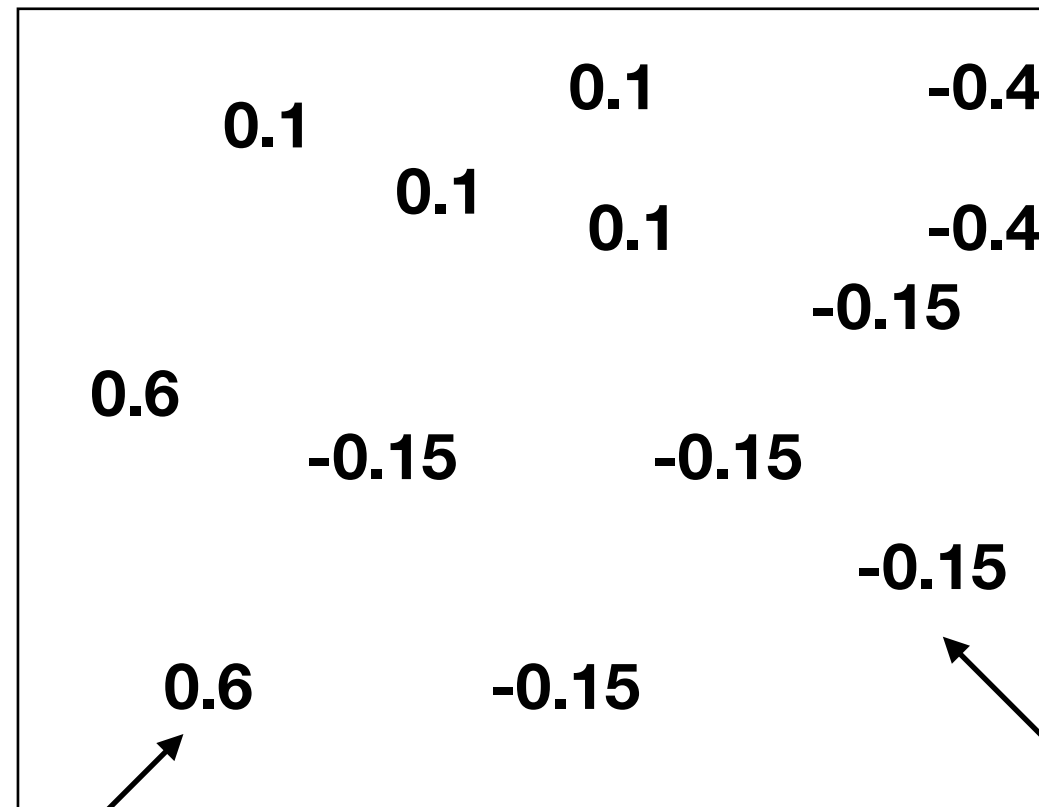
Step 1



Gradient-boosting: the basic idea

Compute the gradient to determine how we should decrease the training error

Step 2



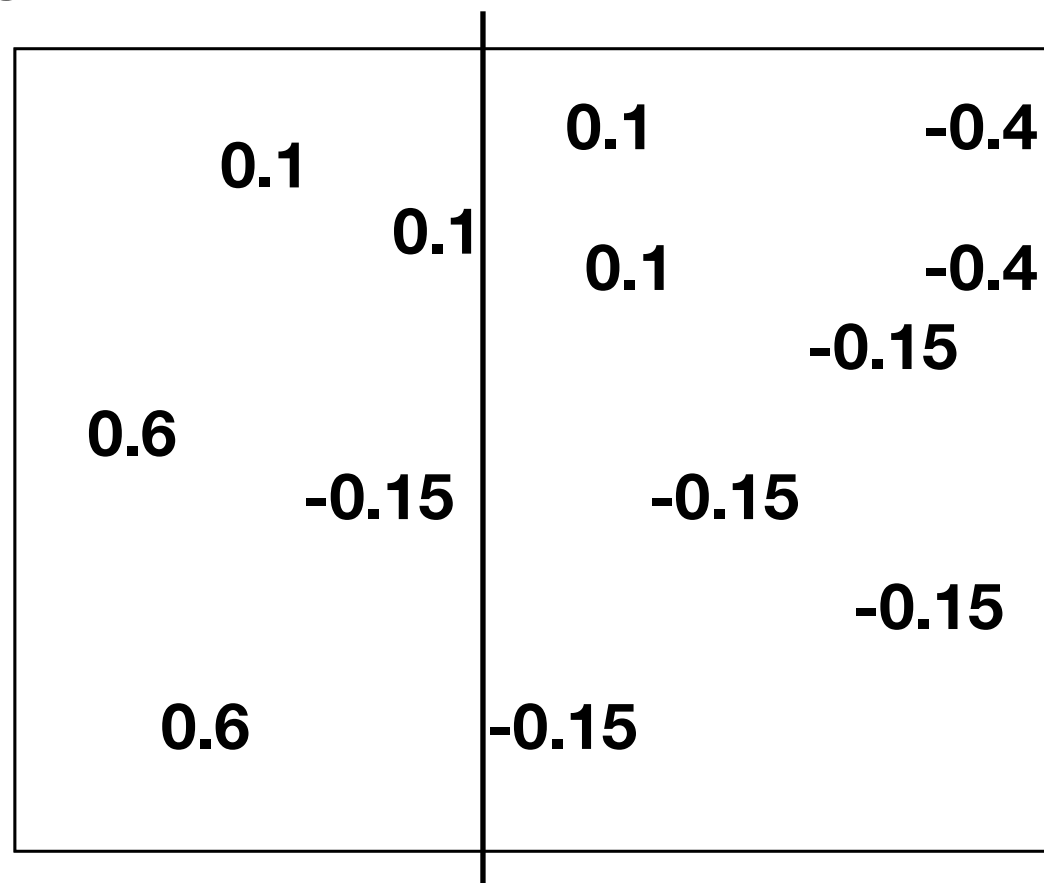
Example: We should increase the probability assigned to class 1 for this label by a lot.

Example: We should decrease the probability assigned to class 1 for this label by a little.

Gradient-boosting: the basic idea

- Our goal is to fit a decision stump that best mimics the values of the gradient.

Step 2



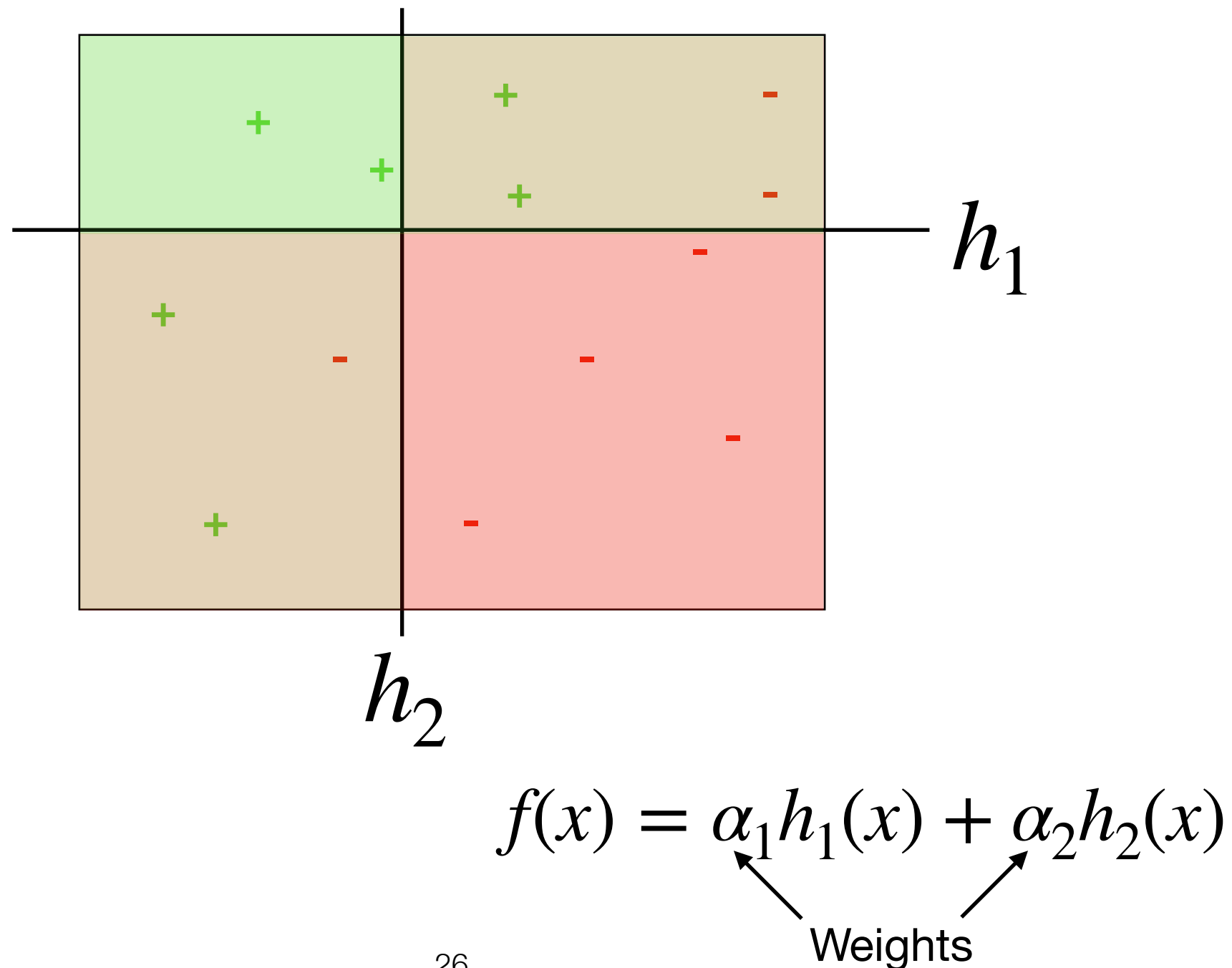
$h_2(x)$ = Average value on the left side
(Increase the probability for class 1)

h_2

$h_2(x)$ = Average value on the right side
(Decrease the probability for class 1)

Gradient-boosting: the basic idea

- The combined model will, again, look like:



Gradient boosting

- At each step, fit a model that best mimics the direction of the gradient.

- For $b = 1, 2, \dots, B$:

- Calculate the gradient of the training loss with respect to the predicted outcomes.

- Find a model h_b that best mimics the gradient.

- Add this new model to the ensemble:

$$f(x) = h_1(x) + \lambda \sum_{j=2}^b \alpha_j h_j(x)$$

Gradient boosted + Trees

- **Gradient boosted trees** are very popular off-the-shelf algorithms because of their high predictive accuracy.
- The tuning parameters you need to think about are:
 - **The number of splits d in each tree:** Controls the complexity of each tree. In practice, $d = 1$ often works well.
 - **The number of trees B :** As B increases, the complexity of our boosted model increases.
 - **The shrinkage parameter λ :** λ controls the rate of learning (i.e. the step size), where a small value means that the model update at each iteration is small.

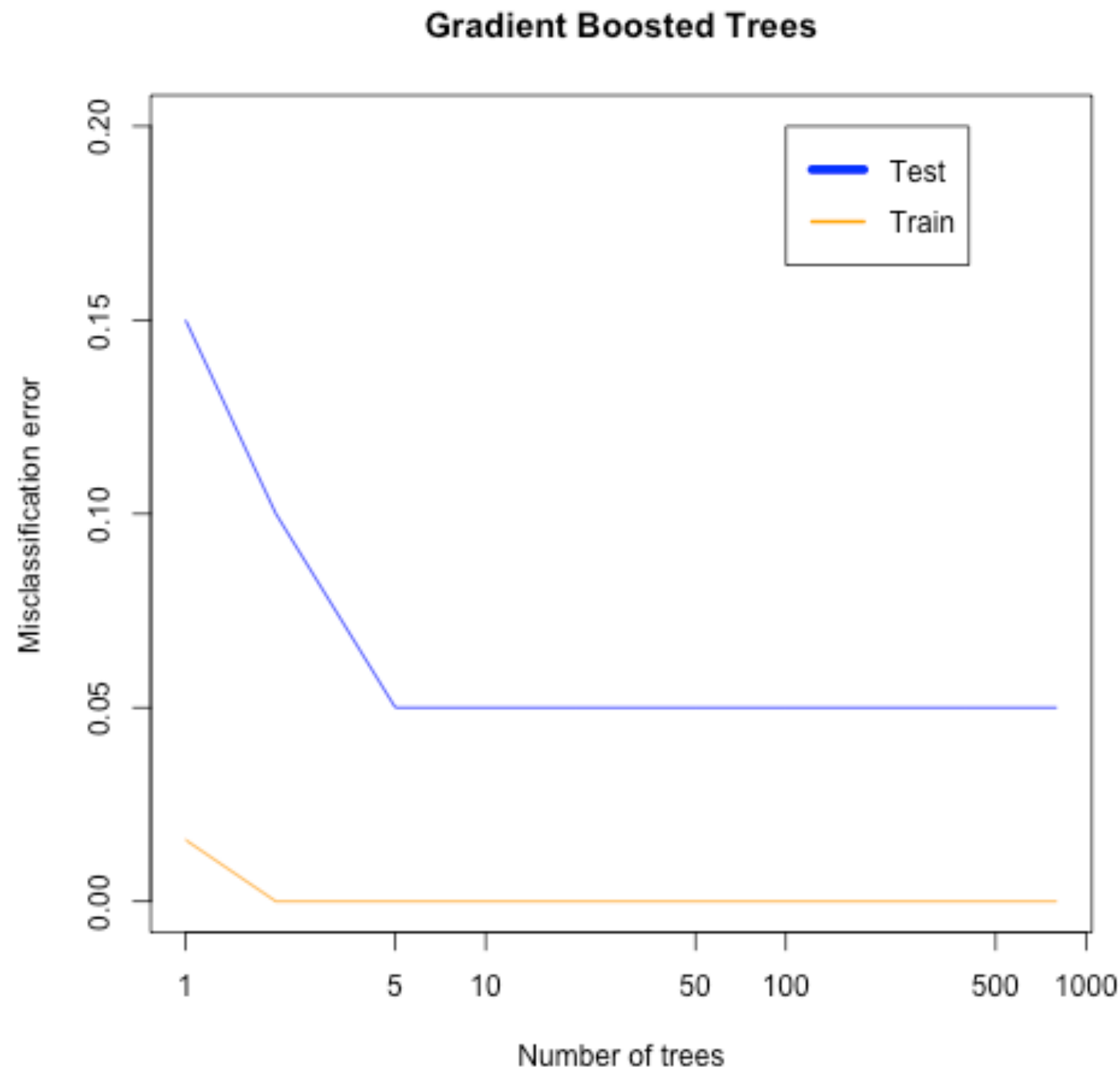
Example: Gene expression data

- Khan gene expression dataset from ISLR:
 - X = Gene expression measurements, $p=2308$
 - Y = Type of small round blue cell tumors, 4 classes
- Training dataset size = 63
- Test dataset size = 20

```
> library(xgboost)
> dtrain <- xgb.DMatrix(Khan$xtrain, label = Khan$ytrain)
> bst <- xgboost(
+   data=dtrain,
+   nrounds = 50,
+   max_depth=2,
+   num_class=num_class,
+   objective="multi:softprob")
```

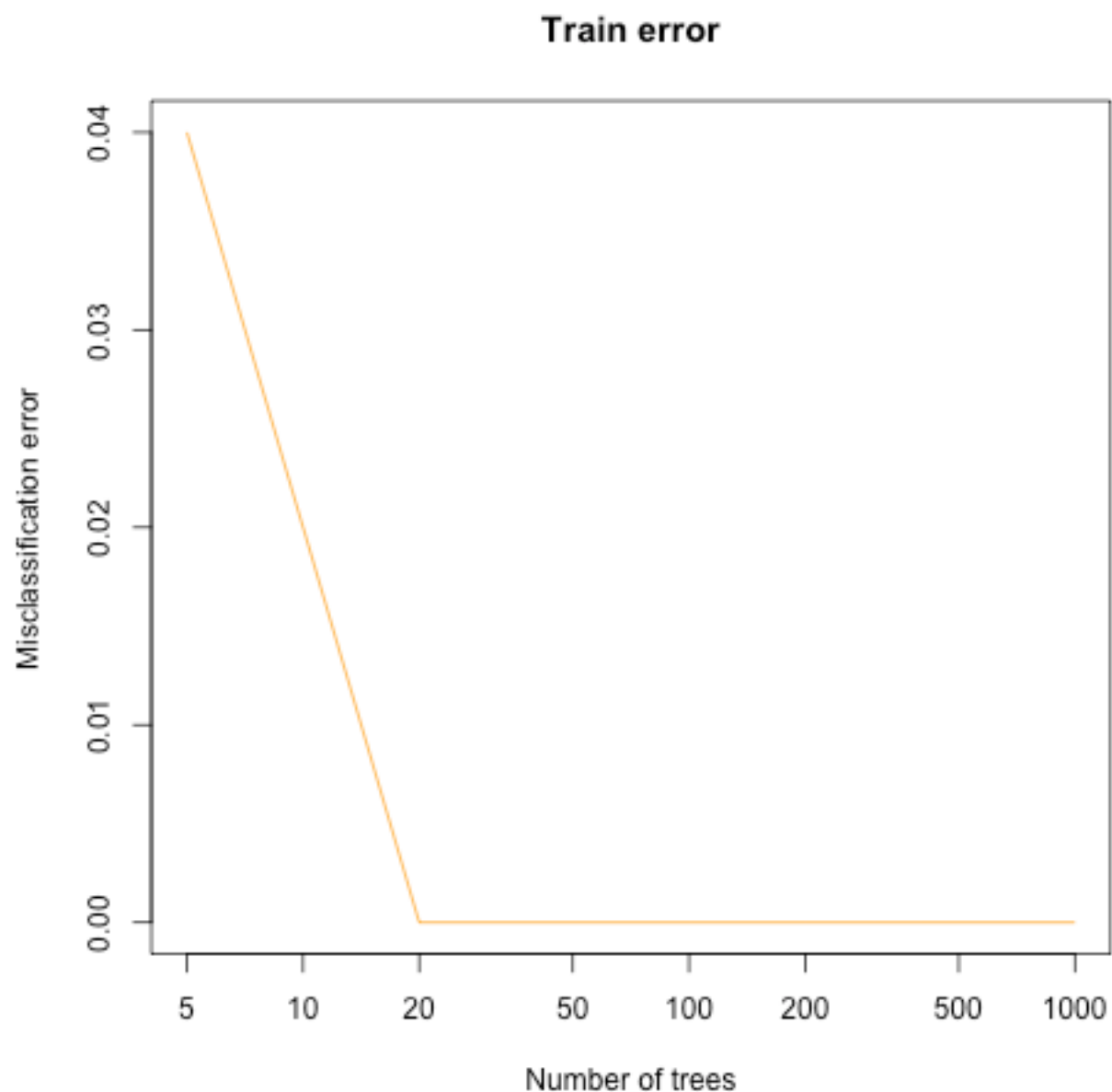
Example: Gene expression data

- Let's vary the number of trees B .



Example: Noisy binary classification problem

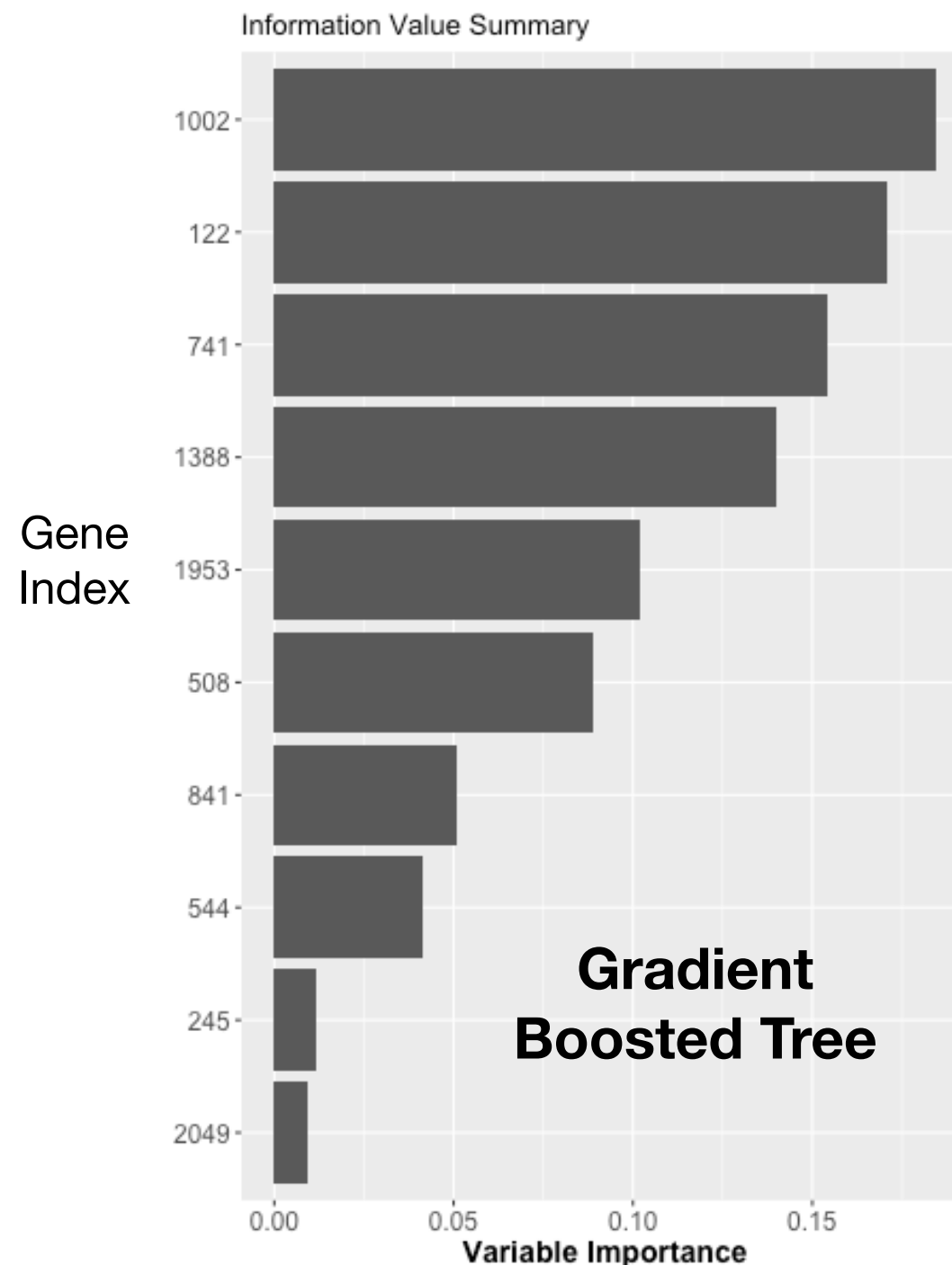
- Simulated data for a noisier binary classification problem.
- We see that boosting overfits in this problem and the number of trees needs to be tuned.



Example: Variable importance

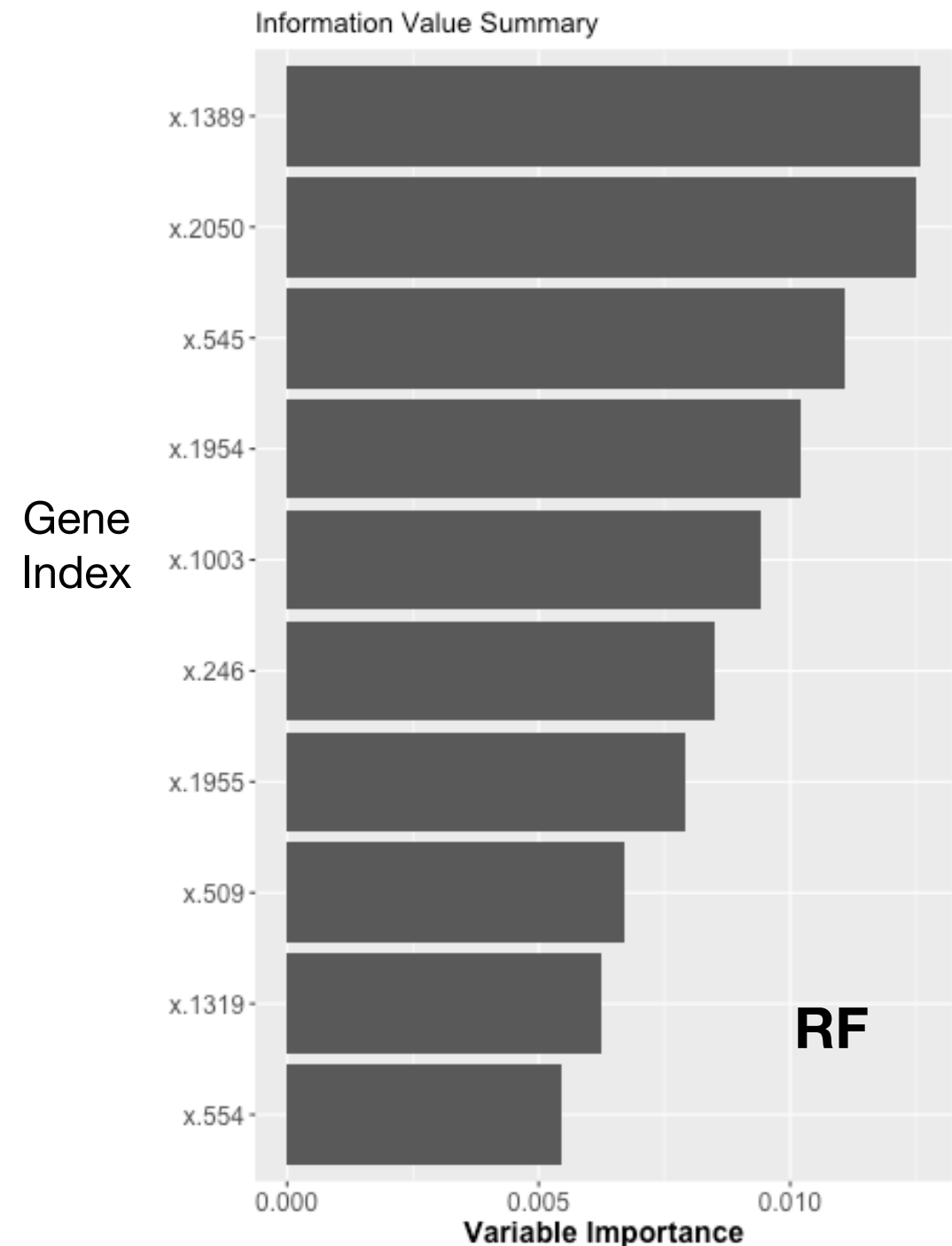
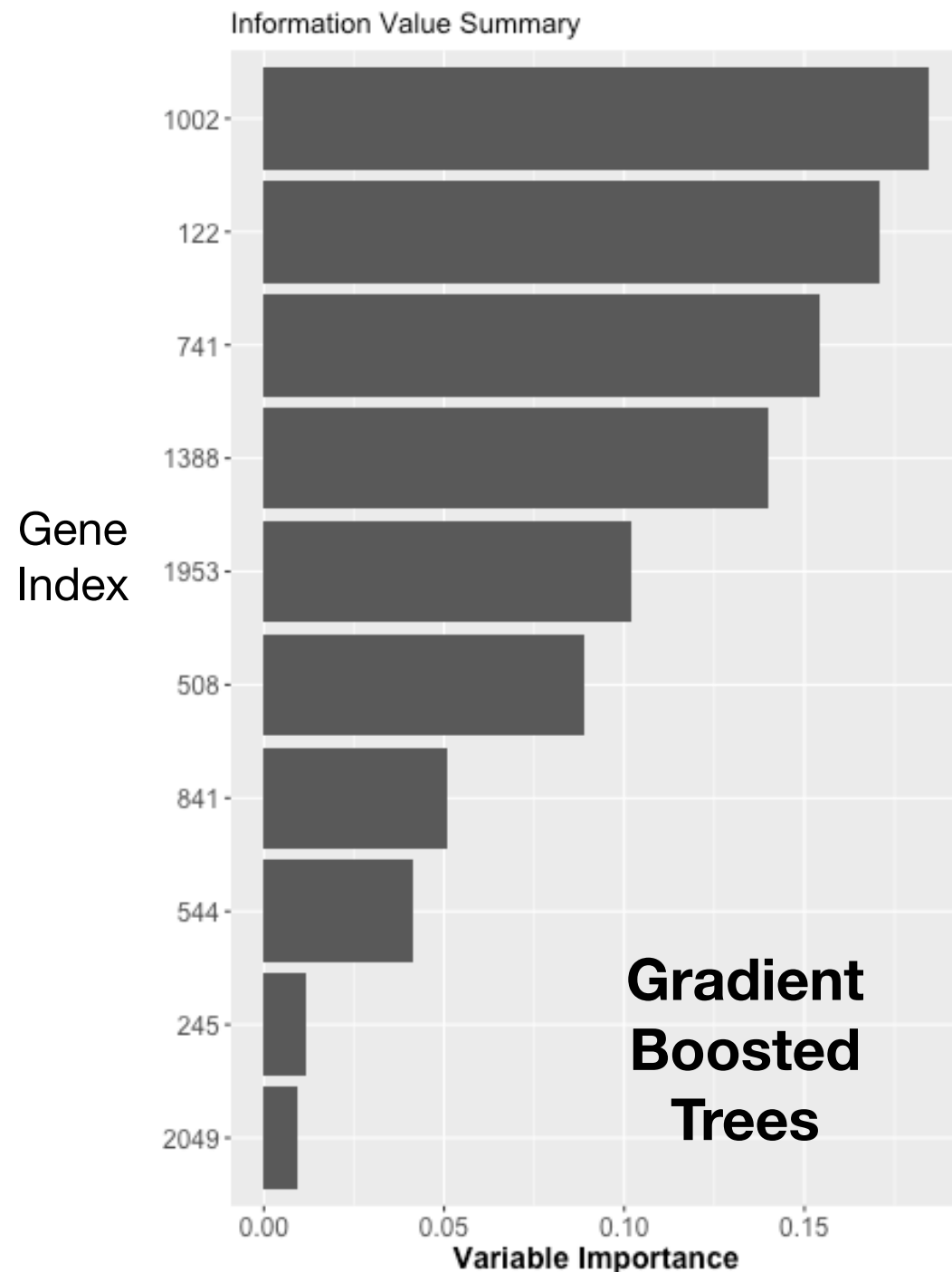
```
> xgb.plot.importance(importance_matrix = head(xgb.importance(model=bst), n = 20))
```

- Let's now try to interpret the gradient boosted tree we fit to the Khan gene expression dataset.
- Here we show the top 10 most important genes according to the variable importance function from xgboost.



Example: Variable importance

- These variable importance estimates are model specific, so be careful when you interpret them!



Variable importance in gradient boosted trees

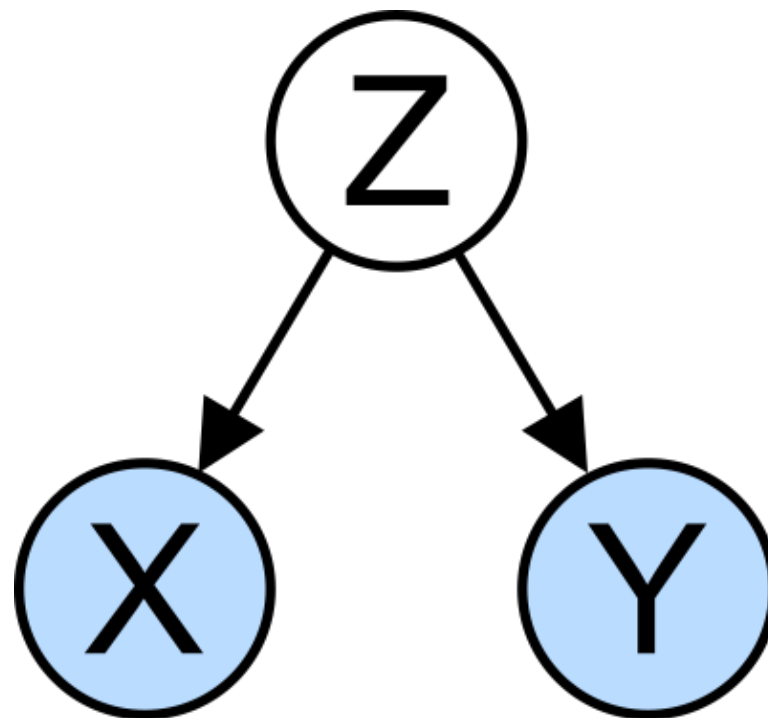
- Consider this example: Suppose there are two highly predictive variables X_1, X_2 that are also highly correlated. Then...
 - Gradient boosted trees will just choose one at random and ignore the other variable.
 - Random forests will choose both of them with similar probabilities.

Today's Lecture

- Boosting
 - AdaBoost
 - Forward Stagewise Additive Modeling
- Gradient boosting
- Issues to think about
 - Confounders
 - Fairness

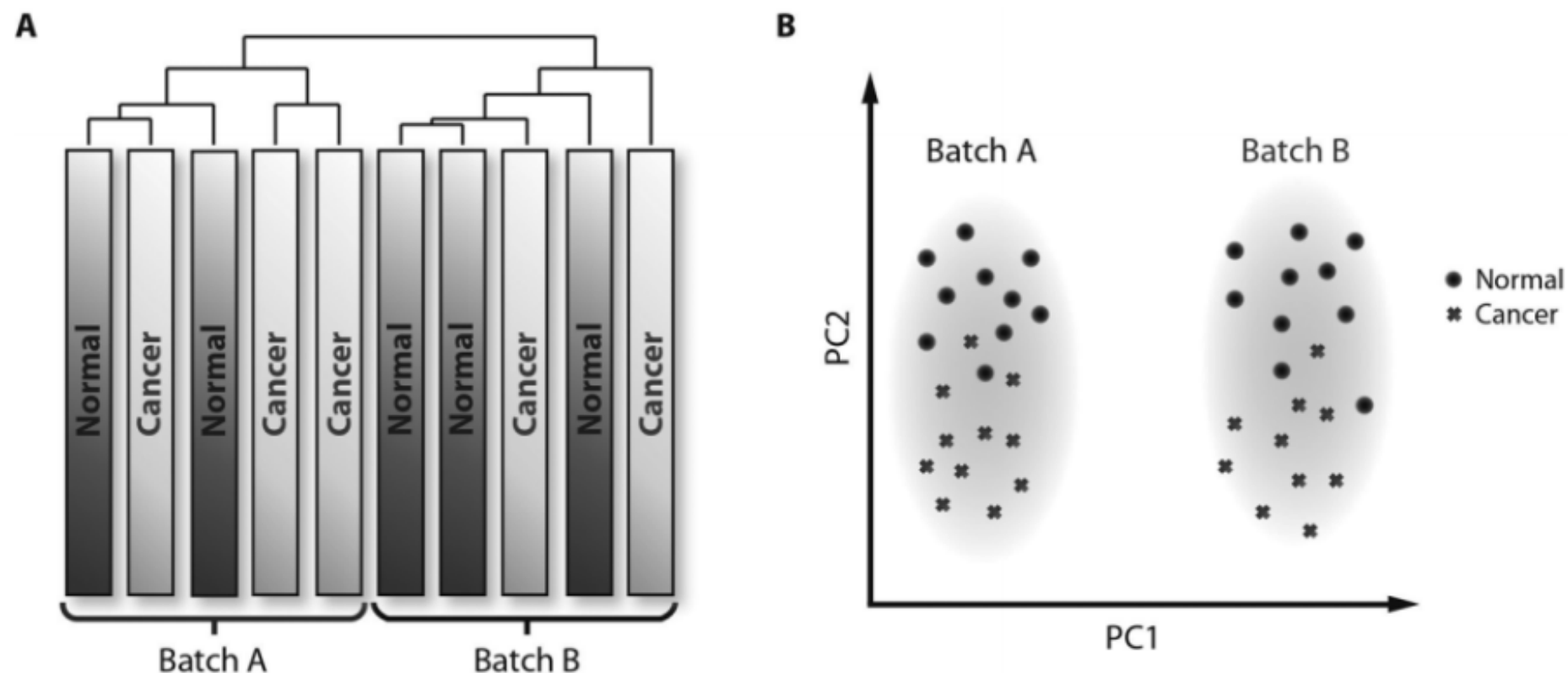
Confounders and spurious associations

- A confounding variable **Z** is associated with both an explanatory variable **X** and the outcome of interest **Y**.
- Machine learning algorithms can rely on confounders to achieve high prediction accuracy for a given dataset. In fact, the prediction accuracies can often be overly optimistic.
- We must carefully consider possible confounding in our datasets. Overlooking these can lead to models that fail to generalize in other settings.



An example from omics

- In omics experiments, the measured variables often have a strong association with non-biological factors (e.g. inter-machine or inter-lab or inter-operator variability, or time of day), which are often referred to as “batch effects.”

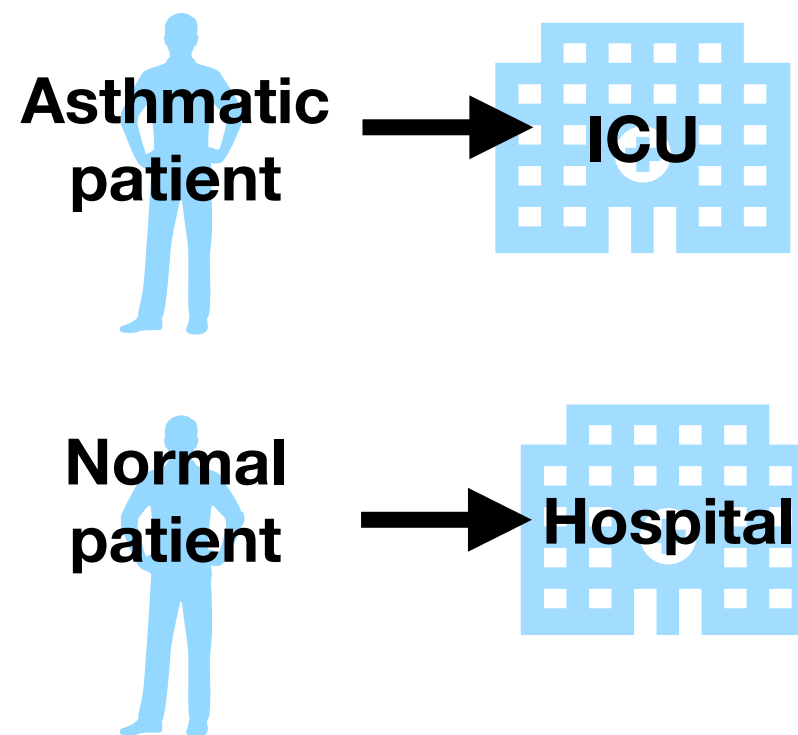


An example from omics

- Petricoin et. al. (2002) reported that their predictive model could use proteomic patterns to discriminate between healthy patients and those with ovarian cancer with 100% sensitivity and 95% specificity!
- However, when independent researchers looked at the publicly released data, they found:
 - inadvertent changes in protocol mid-experiment
 - difference in processing between tumor and normal samples.

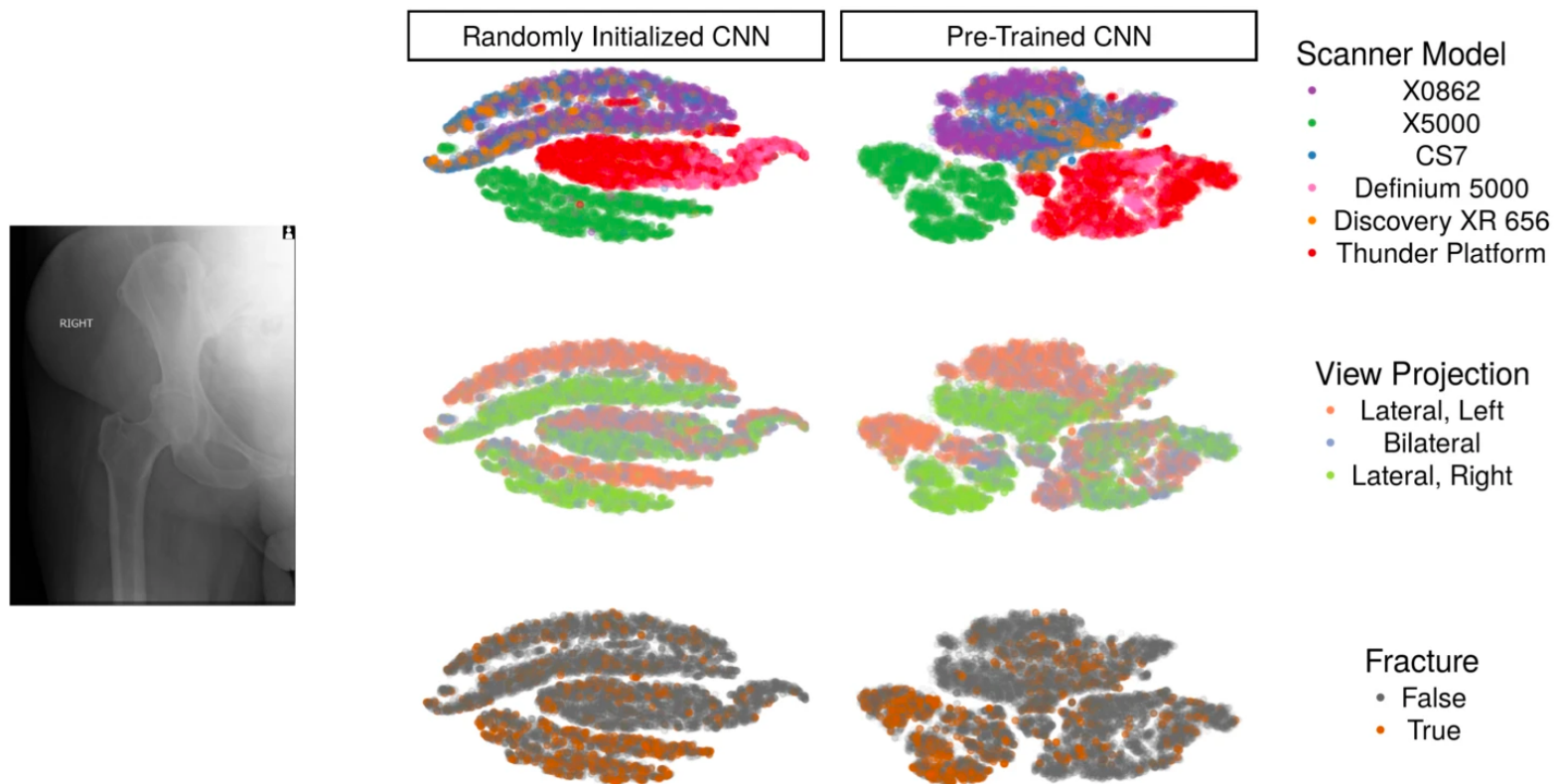
Examples from EHR data

Caruana 2015: ML predicted
that asthmatics are less
likely to die from pneumonia



Examples from imaging

Badgeley 2019: ML made predictions by detecting the scanner model from the radiograph



Confounders and distributional shifts

- Machine learning algorithms simply try to maximize prediction accuracy for the setting from which the data was collected. If future data is going to be collected in exactly the same manner, the prediction accuracy of your ML algorithm will remain stable.
- However, this is often not the case. We often want our model to be usable by other hospitals and labs.
- ML algorithms cannot determine which associations will generalize to other settings and which ones don't (e.g. confounders and spurious associations). In the absence of additional guidance, ML algorithms will learn to depend on whichever associations are the strongest.

Improving data collection

- Good experimental design is the best strategy for reducing spurious associations in the data and avoiding faulty conclusions.
- Randomize sample run times: e.g. don't run cases first and controls second.
- Document how the data was collected (date, machine, cell culture medium, ...)

Improving the analysis

- If you know what confounders commonly exist in your data setting, apply methods that adjust for this known source of confounding.
- Train on data from different institutions, rather than using a single data set. Validate results on independent data sets from a different institution.
- Use an interpretable model and/or understand how the model came to its final conclusion. Check that the model explanations make sense (and fix them if they don't).
- Build in causal reasoning to your prediction model.
- Make sure your data analysis is reproducible. Clearly document your code that contains every step, from data pre-processing to the final analysis.

Cross-validation when combining datasets

- When training models on datasets across different institutions, remember to do site-wise cross-validation!



Today's Lecture

- Boosting
 - AdaBoost
 - Forward Stagewise Additive Modeling
- Gradient boosting
- Issues to think about
 - Confounders
 - Fairness

ML Fairness

RESEARCH ARTICLE

ECONOMICS

Dissecting racial bias in an algorithm used to manage the health of populations

Ziad Obermeyer^{1,2*}, Brian Powers³, Christine Vogeli⁴, Sendhil Mullainathan^{5*†}

“We show that a widely used algorithm ... exhibits significant racial bias: At a given risk score, Black patients are considerably sicker than White patients, as evidenced by signs of uncontrolled illnesses... The bias arises because the algorithm predicts health care costs rather than illness, but unequal access to care means that we spend less money caring for Black patients than for White patients. Thus, despite health care cost appearing to be an effective proxy for health by some measures of predictive accuracy, large racial biases arise. We suggest that the choice of convenient, seemingly effective proxies for ground truth can be an important source of algorithmic bias in many contexts.”

ML fairness

Google Photos Tags Two African-Americans As Gorillas Through Facial Recognition Software

The Best Algorithms Struggle to Recognize Black Faces Equally

US government tests find even top-performing facial recognition systems misidentify blacks at rates five to 10 times higher than they do whites.

- Another well-known example is the Google photo-tagging incident in 2015.
- “Why facial recognition systems perform differently for darker skin tones is unclear. Buolamwini told Congress that many datasets used by companies to test or train facial analysis systems are not properly representative. The easiest place to gather huge collections of faces is from the web, where content skews white, male, and western. Three face-image collections most widely cited in academic studies are 81 percent or more people with lighter skin, according to an IBM review.”

Common causes of unfairness

- **Bias Encoded in Data:** Often the training data that is available encodes human biases.
- **Minimizing Average Error Fits Majority Populations:** The relationships between the variables and outcomes can differ across subpopulations. Unfortunately, most ML algorithms aim to minimize the average error, which may be done at the expense of minority groups.
- **The Need to Explore:** When there is a cost to generate training data for an ML algorithm, do certain subpopulations bear higher burdens?

$$\min_{\beta \in \mathbb{R}^{p+1}} \frac{1}{n} \sum_{i=1}^n \ell \left(y_i, f_{\beta}(x_i) \right)$$

Many ML algorithms aim to minimize the average loss, which down-weights errors made on minority groups.

Chouldechova and Roth. “The Frontiers of Fairness in Machine Learning” 2018.

ML fairness is still a young field

- There are *many* proposals for how to quantify the (un)fairness of an ML algorithms. In general, it is impossible for all of these fairness criteria to hold simultaneously. (Chouldechova 2017)
- There are one-time fairness questions and then there are questions about the dynamics induced by ML algorithms. For example, it is possible for algorithms that were initially fair to become unfair.
- There is also lot of research on how to correct for biases in datasets, which typically requires understanding how the data was generated in the first place.

Chouldechova and Roth. “The Frontiers of Fairness in Machine Learning” 2018.

Today's Lecture

- Boosting
 - AdaBoost
 - Forward Stagewise Additive Modeling
- Gradient boosting
- Issues to think about
 - Confounders
 - Fairness