

Chapter 8: Tree-Based Methods

We will introduce *tree-based* methods for regression and classification.

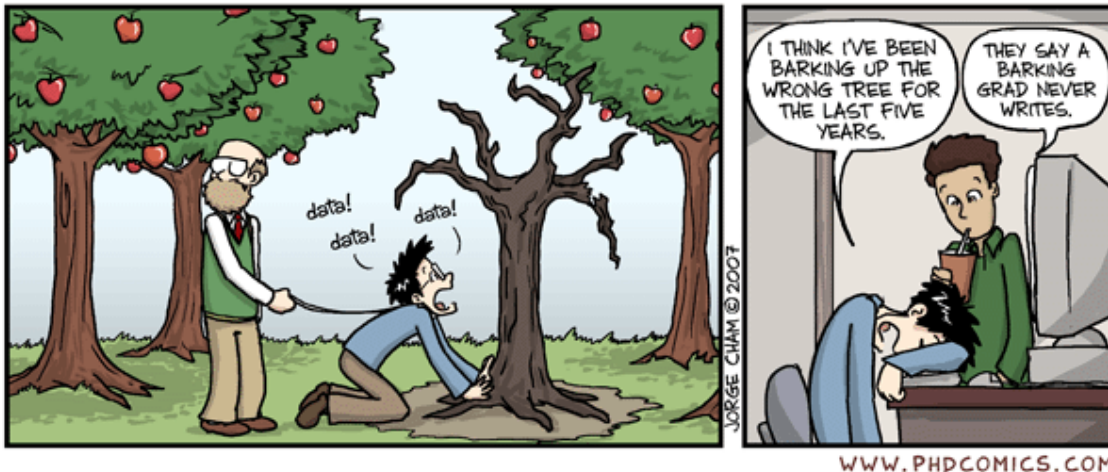
These involve segmenting the predictor space into a number of simple regions.

To make a prediction for an observation, we use the mean or mode of training observations in the region to which it belongs.

The set of splitting rules can be summarized in a tree $\Rightarrow$ “decision trees”.

- simple and useful for interpretation
- not competitive w/ other supervised approaches (e.g. lasso) for prediction.

Combining a large number of trees can often result in dramatic improvements in prediction accuracy at the expense of interpretation.



Credit: <http://phdcomics.com/comics.php?f=852>

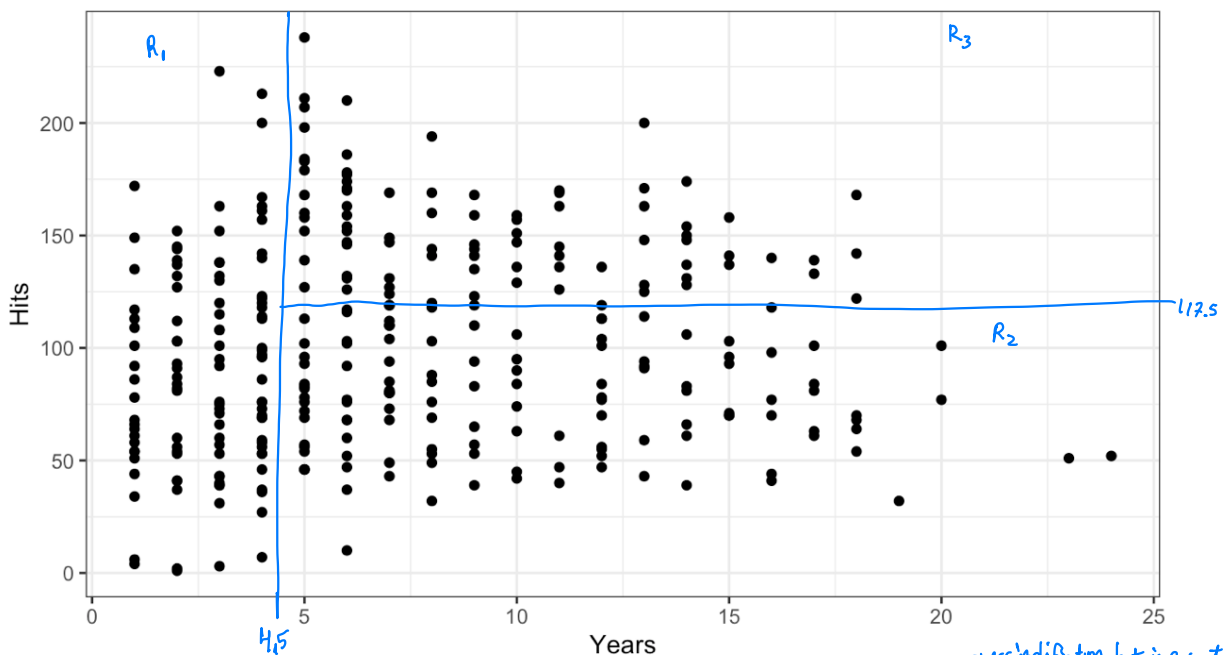
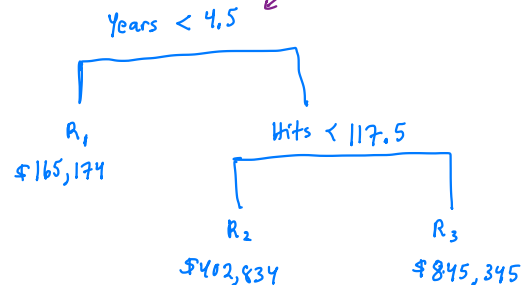
Decision trees can be applied to both regression and classification problems. We will start with regression.

1 Regression Trees

Start with

Example: We want to predict baseball salaries using the **Hitters** data set based on **Years** (the number of years that a player has been in the major leagues) and **Hits** (the number of hits he made the previous year).

We can make a series of splitting rules to create regions and predict salary as the mean in each region.
 according to a tree fit to this data (more later).
 "root" of the tree



The predicted salary for players is given by the mean response value for the players in that box. Overall, the tree segments the players into 3 regions of predictor space.

oversimplification but is easy to interpret and has a nice graphical representation.

terminology:

$R_1, R_2, R_3 =$ terminal nodes or leaves of the tree

points along the tree where predictor space is split = internal nodes

segments of tree that connect nodes = branches

Interpretation

Years is the most important factor in determining salary

→ given that a player has less experience, # hits in previous year play little role in his salary.

→ among players who have been in the league 5+ years, # hits ↑, ↑ salary.

→ quantitative
Y

We now discuss the process of building a regression tree. There are ^Wto steps:

→ set of possible values for $X_1, \dots, X_p$

1. Divide the predictor space

into J distinct and non-overlapping regions $R_1, \dots, R_J$

2. Predict

For every observation that falls in into region R_j we make the same prediction, the mean of the response Y for training values in R_j .

How do we construct the regions $R_1, \dots, R_J$? How to divide the predictor space?
regions could have any shape, but that is too hard (to do, & to interpret)

⇒ divide predictor space into high dimensional rectangles (or "boxes")

The goal is to find boxes $R_1, \dots, R_J$ that minimize the $RSS = \sum_{j=1}^J \sum_{i \in R_j} (y_i - \hat{y}_{R_j})^2$ where

Unfortunately, it is computationally infeasible to consider every possible partition.

$\hat{y}_{R_j}$ = mean response of training data in R_j box.

⇒ take top-down, greedy approach called recursive binary splitting.

The approach is *top-down* because

We start at the top of the tree (all observations belong to a single region) and successively split the predictor space.

The approach is *greedy* because

at each step of the tree building process, the best split is made at that particular step,
↳ not looking ahead to make a split that will lead to a better tree later.

In order to perform recursive binary splitting,

- ① Select the predictor X_j and cutpoint s st. splitting the predictor space into regions $\{X_j < s\}$ and $\{X_j \geq s\}$ leads the greatest reduction in RSS.
 i.e., consider all possible half planes $R_1(j,s) = \{X_j < s\}$ and $R_2(j,s) = \{X_j \geq s\}$. We seek j and s that
 minimize $\sum_{i: x_i \in R_1(j,s)} (y_i - \hat{y}_{R_1})^2 + \sum_{i: x_i \in R_2(j,s)} (y_i - \hat{y}_{R_2})^2 \leftarrow$ finding j and s can be done quickly as long as p not large.
- ② Repeat process looking for best j, s combo, but instead of splitting the entire space, we split $R_1(j,s)$ and $R_2(j,s)$ to minimize RSS.
- ③ Continue until stopping criteria is met (i.e. no region contains more than 5 observations).
- ④ predict using mean of training observations in the region that test obs. fall.

The process described above may produce good predictions on the training set, but is likely to overfit the data.

because resulting tree too complex.

$\Rightarrow$ less regions $R_1, \dots, R_J$

A smaller tree, with less splits might lead to lower variance and better interpretation at the cost of a little bias.

Idea: only split tree if it results in large enough drop in RSS.

Bad idea because seemingly worthless splits early in the tree might be followed by a "good" split.

Better idea

A strategy is to grow a very large tree T_0 and then prune it back to obtain a subtree.

How to prune the tree?

goal: select a subtree that leads to lowest test error rate. $\rightarrow$ could use CV for every subtree but too expensive (large # of possible subtrees).

Solution: "cost complexity pruning"

Consider a sequence of trees indexed by a nonnegative tuning parameter α .

For each value of α , there exists a corresponding subtree T_α st.

$$\sum_{m=1}^M \sum_{x_i \in R_m} (y_i - \hat{y}_{R_m})^2 + \alpha |T| \text{ is as small as possible.}$$

$\nwarrow$
terminal nodes of the tree.

$R_m = m^{\text{th}}$ terminal node region

$\hat{y}_{R_m}$ = predicted response for R_m .

α controls the trade-off between subtree's complexity & fit to training data.

Select α via CV!

Then go back and use full data & chosen α to get subtree.

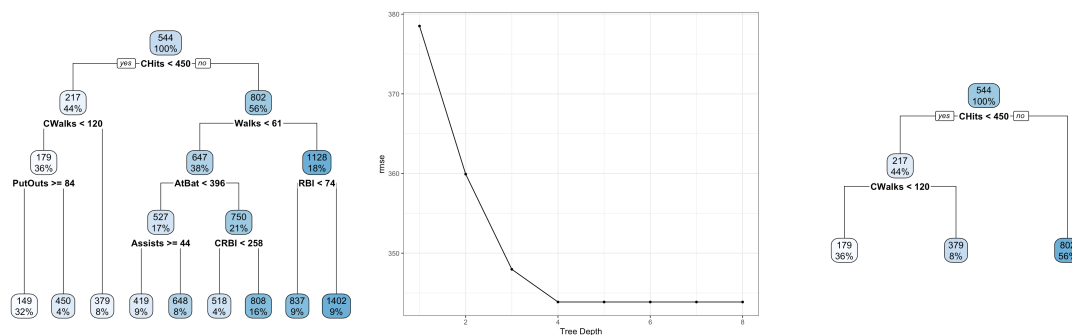
$\rightarrow$ When $\alpha=0$,
 $T=T_0$,

$\alpha \uparrow \Rightarrow$ price to pay for having many terminal nodes $\uparrow \Rightarrow$ smaller tree

Algorithm for building a regression tree:

- ① Use recursive binary splitting to grow a large tree on training data, stopping only when each terminal node has fewer than some # (5?) observations.
- ② Apply cost complexity pruning to the large tree to get a sequence of best trees as a function of α .
- ③ Use k-fold CV to choose α
 - Divide training data into K folds, for each $k=1, \dots, K$
 - (a) repeat ① and ② on all but k^{th} fold
 - (b) Evaluate the MSE on data in k^{th} fold as a function of α .
 - Average results for each value of α to get CV error. Choose α which minimizes CV_k
- ④ Return the subtree from ② that corresponds to α from ③.

Example: Fit regression tree to Hitters use 9 features



2 Classification Trees

A *classification tree* is very similar to a regression tree, except that it is used to predict a categorical response.

For a classification tree, we predict that each observation belongs to the *most commonly occurring class* of training observation in the region to which it belongs.

The task of growing a classification tree is quite similar to the task of growing a regression tree.

It turns out that classification error is not sensitive enough.

When building a classification tree, either the Gini index or the entropy are typically used to evaluate the quality of a particular split.

3 Trees vs. Linear Models

Regression and classification trees have a very different feel from the more classical approaches for regression and classification.

Which method is better?

3.1 Advantages and Disadvantages of Trees

4 Bagging

Decision trees suffer from *high variance*.

Bootstrap aggregation or *bagging* is a general-purpose procedure for reducing the variance of a statistical learning method, particularly useful for trees.

So a natural way to reduce the variance is to take many training sets from the population, build a separate prediction model using each training set, and average the resulting predictions.

Of course, this is not practical because we generally do not have access to multiple training sets.

While bagging can improve predictions for many regression methods, it's particularly useful for decision trees.

These trees are grown deep and not pruned.

How can bagging be extended to a classification problem?

4.1 Out-of-Bag Error

There is a very straightforward way to estimate the test error of a bagged model, without the need to perform cross-validation.

4.2 Interpretation

5 Random Forests

Random forests provide an improvement over bagged trees by a small tweak that decorrelates the trees.

As with bagged trees, we build a number of decision trees on bootstrapped training samples.

In other words, in building a random forest, at each split in the tree, the algorithm is not allowed to consider a majority of the predictors.

The main difference between bagging and random forests is the choice of predictor subset size m .

6 Boosting

Boosting is another approach for improving the prediction results from a decision tree.

While bagging involves creating multiple copies of the original training data set using the bootstrap and fitting a separate decision tree on each copy,

Boosting does not involve bootstrap sampling, instead each tree is fit on a modified version of the original data set.

Boosting has three tuning parameters:

1.

2.

3.