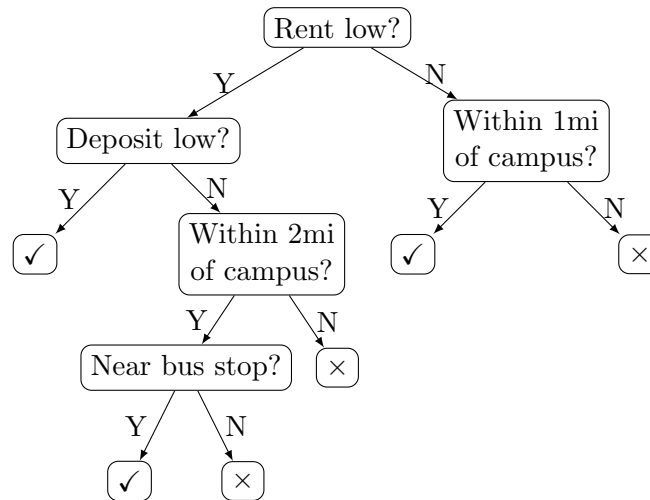


07-280 Decision Trees

Nihar B. Shah

Consider the following setting. You want to rent an apartment. Among the many options available, you want to decide which apartments to check out in person. How do you decide?



- Internal nodes split on some attribute
- Leaves have output labels (✓ or ×)

This is called a **decision tree**.

Major Benefit: Human interpretable.

We created this using a rule-based approach. In many situations, you don't have access to such rules or the setting is just too complex.

ML approach: Given examples, how do you construct a decision tree?

What do the examples comprise?

In this problem, each example is 1 prospective apartment:

- “ x ” is answers to “Rent low?”, “Deposit low?”, “1 mi?”, “2 mi?”, “bus stop?”
- “ y ” is answer to whether to go check the apt or not.

Formally:

$N = \# \text{ prospective apts}$

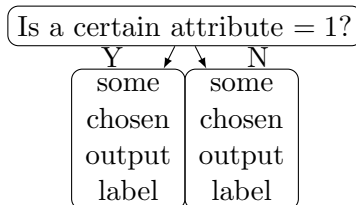
$\mathcal{X} = \{\text{Yes, No}\}^5 = \{0, 1\}^5$

$\mathcal{Y} = \{Y, N\} = \{0, 1\}$

Once you construct a decision tree, given any new apt $x^{(\text{new})} \in \mathcal{X}$, you run it through the tree to get an output as to whether you should check it out or not.

More generally, consider $\mathcal{X} = \{0, 1\}^d$ (d attributes), and $\mathcal{Y} = \{0, 1\}$.

Key building block: Decision Stump Suppose you can split just once...



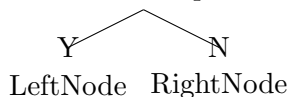
Constructing the Decision Tree

Algorithm 1 Decision Tree Construction

```

1: function DT(dataset  $S$ , set of attributes  $A \subseteq \{1, 2, \dots, d\}$ )
2:   if  $A = \emptyset$  or all examples in  $S$  have same label then
3:     return a leaf with value = majority vote of labels in  $S$ 
4:   end if
5:   Let  $b = \text{BESTATTRIBUTE}(S, A) \in A$ 
6:   Let LeftNode = DT( $\{(x, y) \in S \mid x_b = 1\}, A \setminus \{b\}$ )
7:   Let RightNode = DT( $\{(x, y) \in S \mid x_b = 0\}, A \setminus \{b\}$ )
8:   return tree      “Is  $b^{\text{th}}$  attribute equal to 1?”

```



```

9: end function

```

This function is called with $S = \{(x_i, y_i)\}_{i=1}^N$ and $A = \{1, 2, \dots, d\}$.

Consider one kind of BestAttribute selection: Assume you can split only once (i.e., create a decision stump). Then what choice of attribute to split (and appropriate choice of labels for the two resulting leaves) will lead to lowest training error?

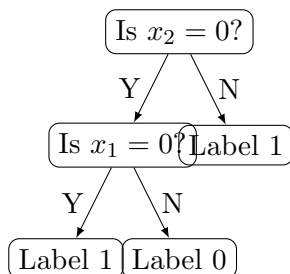
Example (Poll): Training data

y	x_1	x_2	x_3
0	1	0	0
0	1	0	1
0	1	0	0
1	0	0	1
1	1	1	0
1	1	1	1
1	1	1	0
1	1	1	1

Answer: Attribute 2.

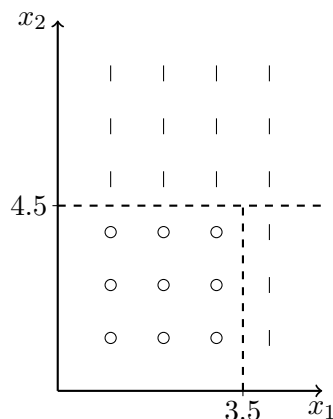
The decision tree construction algorithm initially picks the attribute which would result in the lowest training error if allowed only one split. (Note that this is a greedy approach towards ERM, where greediness is suboptimal but helps in computational efficiency.) It then passes all training samples which have that attribute equal to 1 to the left node, and remaining samples to the right node. It now proceeds recursively until a stopping criterion is met. At this point, we are at a leaf and the algorithm outputs a label.

In the aforementioned example, the tree is:



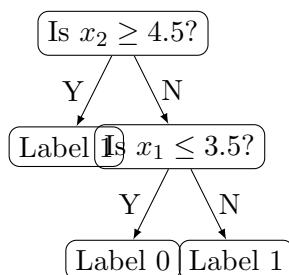
Various Extensions

- **Attributes having more than $\{0, 1\}$ values**



E.g., $\mathcal{X} = \mathbb{R}^2$, $\mathcal{Y} = \{0, 1\}$

Tree we get:



“Decision boundaries” the set of lines that separate output ‘0’ from output ‘1’ regions.

- Need not remove chosen attribute when calling function again in recursion
- **Different stopping criteria** (e.g., stop if the split doesn’t reduce training error by much)

- **BestAttribute()**:
 - Minimize training error (earlier)
 - Minimize Gini impurity (won't discuss here)
 - Maximize mutual information (below)

Information Theory

We will discuss a few concepts on quantifying uncertainty and information that are widely used in machine learning and in many other disciplines.

Entropy

First consider a “20 questions” like game. Suppose I choose one of the students in the class uniformly at random. There are 128 students in the class. You can ask me Yes/No questions and have to guess who I chose. How many questions would you need on average?

$$\text{Answer} = 7 = \log_2 128 = \log_2 \frac{1}{1/128}$$

Notice that the probability with which I chose any given student was $1/128$. Thus the term $\log_2 \frac{1}{\text{Probability}}$ is fundamental! If $k = \# \text{students}$ and $p_1, \dots, p_k$ are probabilities of choosing students $1, \dots, k$ respectively, then here $p_1 = p_2 = \dots = p_k = \frac{1}{128}$. We can rewrite:

$$\text{Answer} = \sum_{i=1}^k p_i \log \frac{1}{p_i}$$

Now suppose I choose a student not uniformly at random, but with probabilities $p_1 = \frac{1}{4}$, $p_2 = \frac{1}{8}, \dots, p_7 = \frac{1}{8}$, $p_8, \dots, p_{128} = 0$. Then average number of questions you need to identify the chosen student is still $\sum_{i=1}^k p_i \log \frac{1}{p_i}$ (where $0 \cdot \log \frac{1}{0} := 0$). This equals $11/4$ for this problem.

The quantity $\sum_{i=1}^k p_i \log \frac{1}{p_i}$ is called the **entropy** of the distribution $(p_1, \dots, p_k)$.

There are many cool properties, uses, and results about entropy. Measures amount of uncertainty: more uncertainty $\iff$ more questions needed.

Recall that we wanted to choose attributes that reduce amount of uncertainty of the labels. We will now work towards that.

Let $\{P(Y = y)\}_{y \in \mathcal{Y}}$ be the distribution of the labels in S . E.g., if 40% of labels in S are “1” and 60% are “0” then $P(Y = 1) = 0.4$ and $P(Y = 0) = 0.6$.

Then the entropy of these labels is:

$$\sum_{y \in \mathcal{Y}} P(Y = y) \log \frac{1}{P(Y = y)},$$

denoted as $H(Y)$.

We want to choose an attribute so that splitting on this attribute will reduce the entropy as much as possible.

Consider a random variable W . (We will later set W to be the attribute that was split.)

Specific conditional entropy:

$$H(Y \mid W = \omega) = \sum_{y \in \mathcal{Y}} P(Y = y \mid W = \omega) \log \frac{1}{P(Y = y \mid W = \omega)}$$

- The conditioning $|W = \omega$ defines left or right child node
- $P(Y|W = \omega)$ Distribution of labels when you take subset of S that has attribute $W = \omega$

Conditional entropy:

$$H(Y | W) = \sum_{\omega \in \mathcal{W}} P(W = \omega) \cdot H(Y | W = \omega)$$

This is the average entropy across the child nodes (after the current node is split).

Mutual information: Amount of reduction in entropy (before vs. after splitting):

$$I(Y; W) = H(Y) - H(Y | W)$$

BestAttribute chooses the attribute that maximizes mutual information (equivalently, minimizes conditional entropy).