

Decision Trees

Doug Downey
EECS 348 Spring 2013

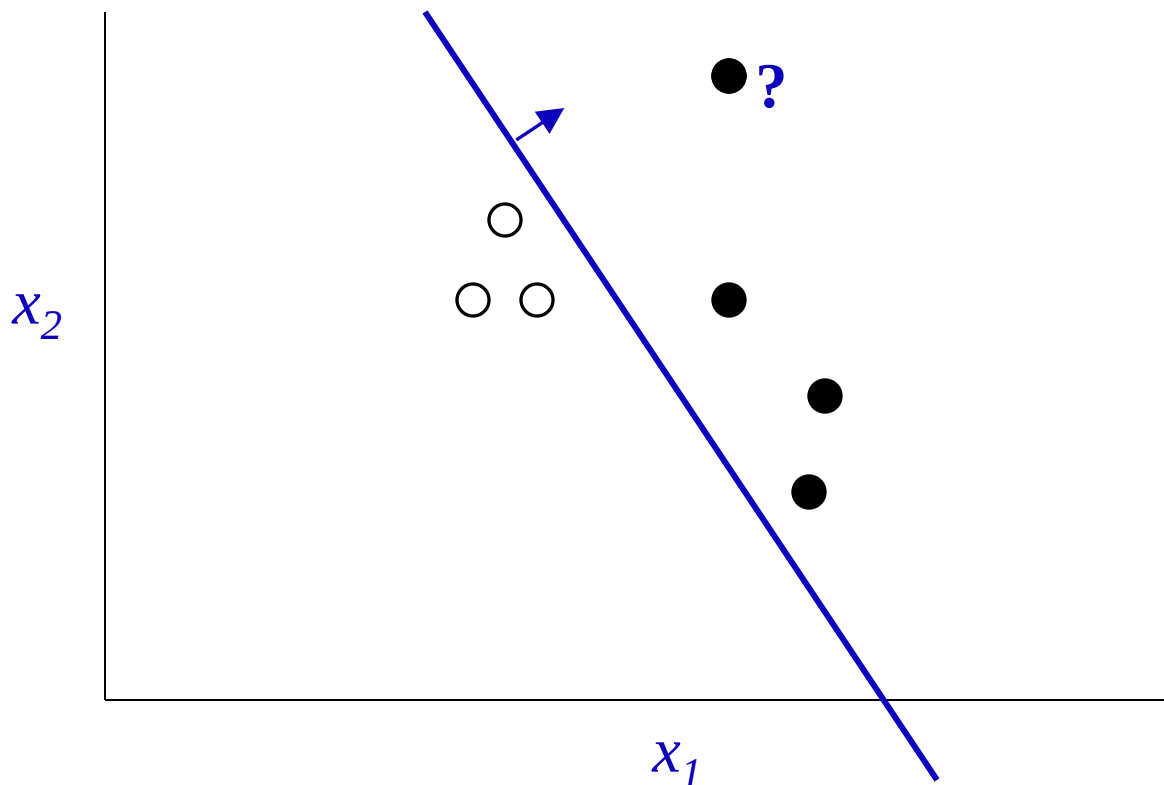
with slides from Pedro Domingos, Bryan Pardo

Outline

- Classical AI Limitations
 - Knowledge Acquisition Bottleneck, Brittleness
- “Modern” directions:
 - Situatedness, embodiment
 - Probability
 - Learning from data (machine learning)

Recall: example

Learn function from $\mathbf{x} = (x_1, \dots, x_d)$ to $f(\mathbf{x}) \in \{0, 1\}$
given **labeled** examples $(\mathbf{x}, f(\mathbf{x}))$



Instances

- E.g. Days, in terms of weather:

Sky	Temp	Humid	Wind	Water	Forecast
sunny	warm	normal	strong	warm	same
sunny	warm	high	strong	warm	same
rainy	cold	high	strong	warm	change
sunny	warm	high	strong	cool	change

Functions

- “Days on which my friend Aldo enjoys his favorite water sport”

INPUT $\mathbf{x} = (x_1, \dots, x_6)$

OUTPUT



Sky	Temp	Humid	Wind	Water	Forecast	f(x)
sunny	warm	normal	strong	warm	same	1
sunny	warm	high	strong	warm	same	1
rainy	cold	high	strong	warm	change	0
sunny	warm	high	strong	cool	change	1

Machine Learning!

- Learn function from training examples, **predict** the output for a new instance

INPUT $\mathbf{x} = (x_1, \dots, x_6)$

OUTPUT

Sky	Temp	Humid	Wind	Water	Forecast	f(x)
sunny	warm	normal	strong	warm	same	1
sunny	warm	high	strong	warm	same	1
rainy	cold	high	strong	warm	change	0
sunny	warm	high	strong	cool	change	1
rainy	warm	high	strong	cool	change	?

General Machine Learning Task

DEFINE:

- Set X of **instances** (of n -tuples $\mathbf{x} = \langle x_1, \dots, x_n \rangle$)
 - E.g., days described by **attributes** (or **features**):
Sky, Temp, Humidity, Wind, Water, Forecast
- **Target function** f , e.g.:
 - EnjoySport $X \rightarrow Y = \{0,1\}$

GIVEN:

- **Training examples** D
 - examples of the target function: $\langle \mathbf{x}, f(\mathbf{x}) \rangle$

FIND:

- A **hypothesis** h such that $h(\mathbf{x})$ approximates $f(\mathbf{x})$.

Task: Will I wait for a table?

Example	Attributes										Target
	<i>Alt</i>	<i>Bar</i>	<i>Fri</i>	<i>Hun</i>	<i>Pat</i>	<i>Price</i>	<i>Rain</i>	<i>Res</i>	<i>Type</i>	<i>Est</i>	<i>WillWait</i>
X_1	<i>T</i>	<i>F</i>	<i>F</i>	<i>T</i>	<i>Some</i>	<i>\$\$\$</i>	<i>F</i>	<i>T</i>	<i>French</i>	<i>0-10</i>	<i>T</i>
X_2	<i>T</i>	<i>F</i>	<i>F</i>	<i>T</i>	<i>Full</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Thai</i>	<i>30-60</i>	<i>F</i>
X_3	<i>F</i>	<i>T</i>	<i>F</i>	<i>F</i>	<i>Some</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Burger</i>	<i>0-10</i>	<i>T</i>
X_4	<i>T</i>	<i>F</i>	<i>T</i>	<i>T</i>	<i>Full</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Thai</i>	<i>10-30</i>	<i>T</i>
X_5	<i>T</i>	<i>F</i>	<i>T</i>	<i>F</i>	<i>Full</i>	<i>\$\$\$</i>	<i>F</i>	<i>T</i>	<i>French</i>	<i>>60</i>	<i>F</i>
X_6	<i>F</i>	<i>T</i>	<i>F</i>	<i>T</i>	<i>Some</i>	<i>\$\$</i>	<i>T</i>	<i>T</i>	<i>Italian</i>	<i>0-10</i>	<i>T</i>
X_7	<i>F</i>	<i>T</i>	<i>F</i>	<i>F</i>	<i>None</i>	<i>\$</i>	<i>T</i>	<i>F</i>	<i>Burger</i>	<i>0-10</i>	<i>F</i>
X_8	<i>F</i>	<i>F</i>	<i>F</i>	<i>T</i>	<i>Some</i>	<i>\$\$</i>	<i>T</i>	<i>T</i>	<i>Thai</i>	<i>0-10</i>	<i>T</i>
X_9	<i>F</i>	<i>T</i>	<i>T</i>	<i>F</i>	<i>Full</i>	<i>\$</i>	<i>T</i>	<i>F</i>	<i>Burger</i>	<i>>60</i>	<i>F</i>
X_{10}	<i>T</i>	<i>T</i>	<i>T</i>	<i>T</i>	<i>Full</i>	<i>\$\$\$</i>	<i>F</i>	<i>T</i>	<i>Italian</i>	<i>10-30</i>	<i>F</i>
X_{11}	<i>F</i>	<i>F</i>	<i>F</i>	<i>F</i>	<i>None</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Thai</i>	<i>0-10</i>	<i>F</i>
X_{12}	<i>T</i>	<i>T</i>	<i>T</i>	<i>T</i>	<i>Full</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Burger</i>	<i>30-60</i>	<i>T</i>

Classification of examples is **positive** (T) or **negative** (F)

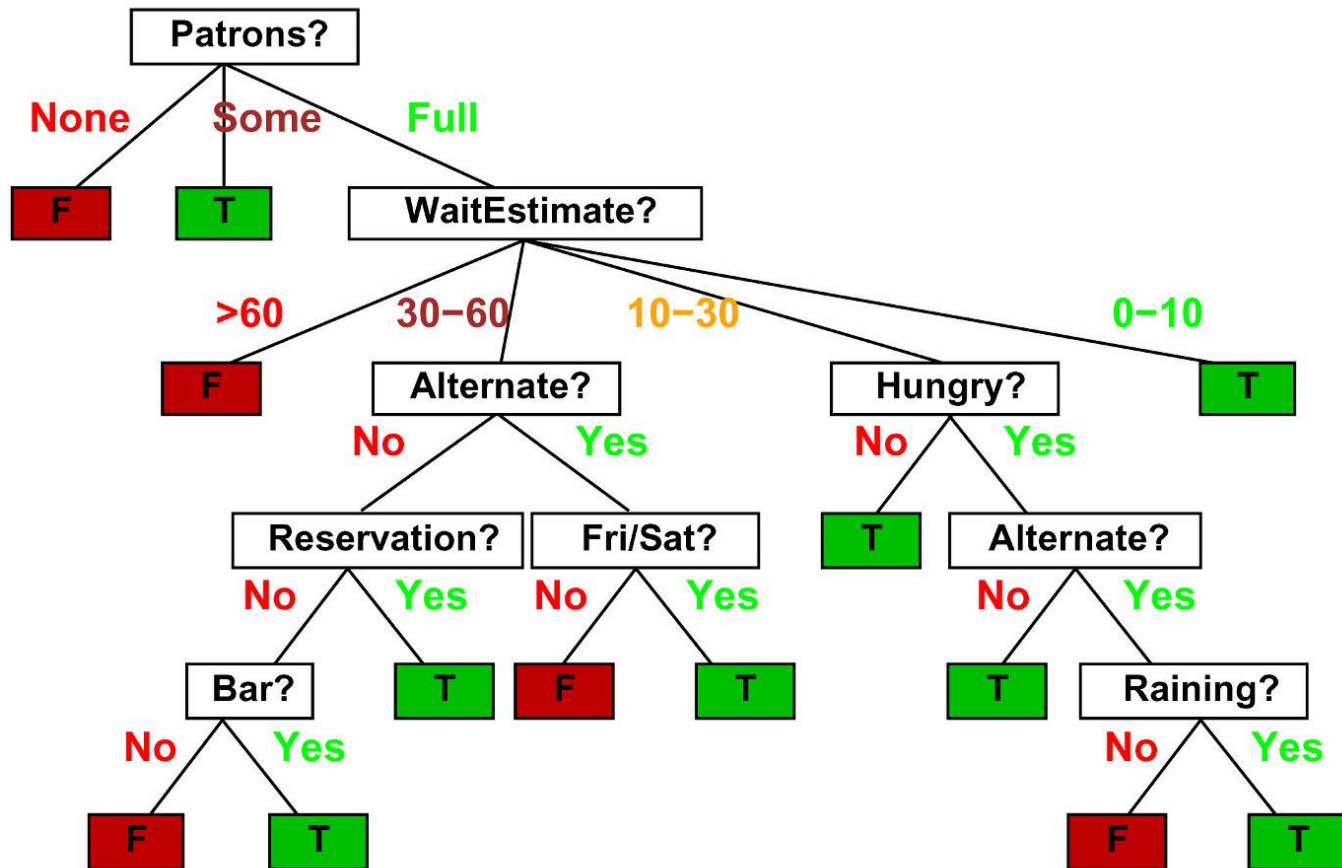
Hypothesis Spaces

- **Hypothesis space** H is a subset of all $g: X \rightarrow Y$ e.g.:
 - Linear separators
 - Conjunctions of constraints on attributes (humidity must be low, and outlook != rain)
 - Etc.
- In machine learning, we restrict ourselves to H

Decision Trees!

One possible representation for hypotheses

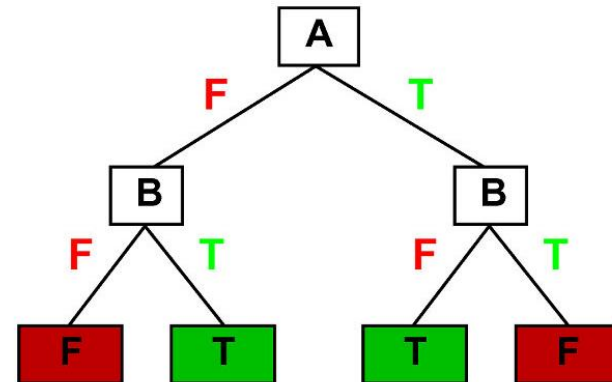
E.g., here is the “true” tree for deciding whether to wait:



Expressiveness of D-Trees

Decision trees can express any function of the input attributes.
E.g., for Boolean functions, truth table row $\rightarrow$ path to leaf:

A	B	A xor B
F	F	F
F	T	T
T	F	T
T	T	F

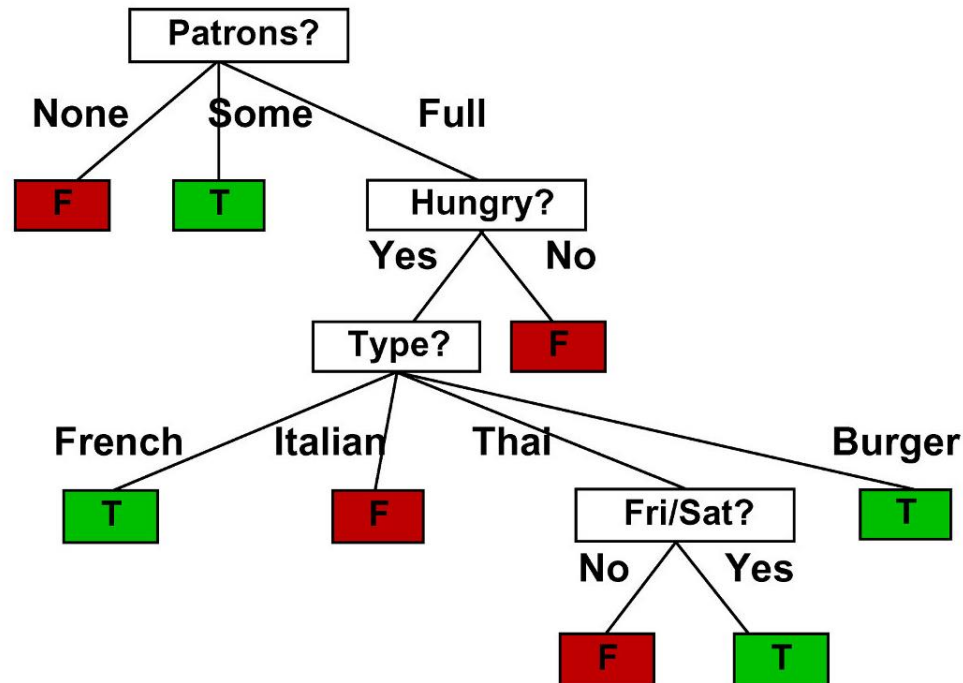


Trivially, there is a consistent decision tree for any training set
w/ one path to leaf for each example (unless f nondeterministic in x)
but it probably won't generalize to new examples

Prefer to find more **compact** decision trees

A learned decision tree

Decision tree learned from the 12 examples:



Substantially simpler than “true” tree—a more complex hypothesis isn’t justified by small amount of data

Inductive Bias

- To learn, we **must** prefer some functions to others
 - **Selection bias**
 - use a **restricted** hypothesis space, e.g.:
 - linear separators
 - 2-level decision trees
 - **Preference bias**
 - use the whole function space, but state a **preference** over concepts, e.g.:
 - *Lowest-degree* polynomial that separates the data
 - *shortest* decision tree that fits the data 

Decision Tree Learning (ID3)

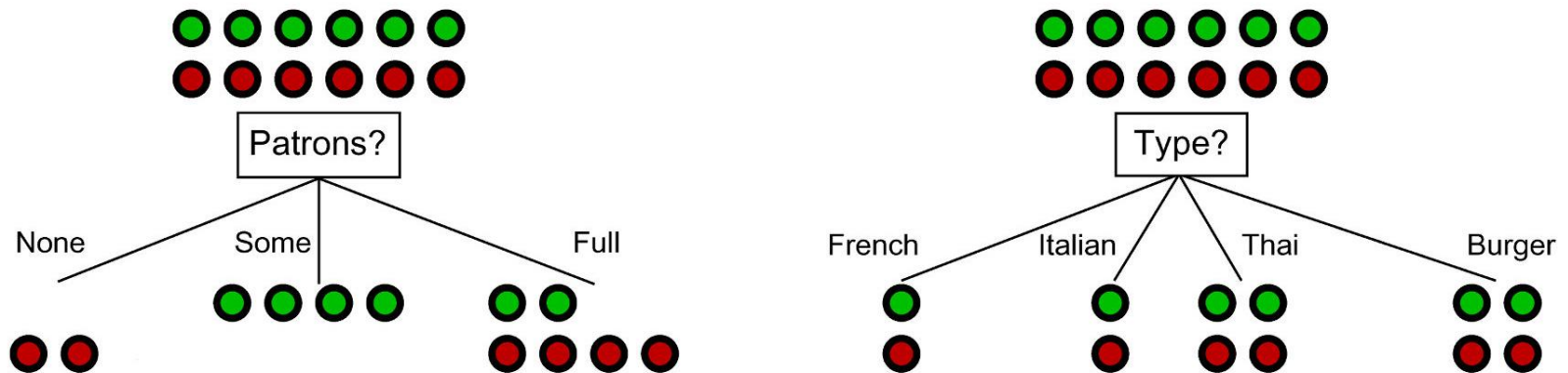
Aim: find a small tree consistent with the training examples

Idea: (recursively) choose “most significant” attribute as root of (sub)tree

```
function DTL(examples, attributes, default) returns a decision tree
  if examples is empty then return default
  else if all examples have the same classification then return the classification
  else if attributes is empty then return MODE(examples)
  else
    best ← CHOOSE-ATTRIBUTE(attributes, examples)
    tree ← a new decision tree with root test best
    for each value  $v_i$  of best do
      examplesi ← {elements of examples with best =  $v_i$ }
      subtree ← DTL(examplesi, attributes – best, MODE(examples))
      add a branch to tree with label  $v_i$  and subtree subtree
  return tree
```

Choosing an attribute

Idea: a good attribute splits the examples into subsets that are (ideally) “all positive” or “all negative”



Patrons? is a better choice—gives **information** about the classification

Information

Information answers questions

The more clueless I am about the answer initially, the more information is contained in the answer

Scale: 1 bit = answer to Boolean question with prior $\langle 0.5, 0.5 \rangle$

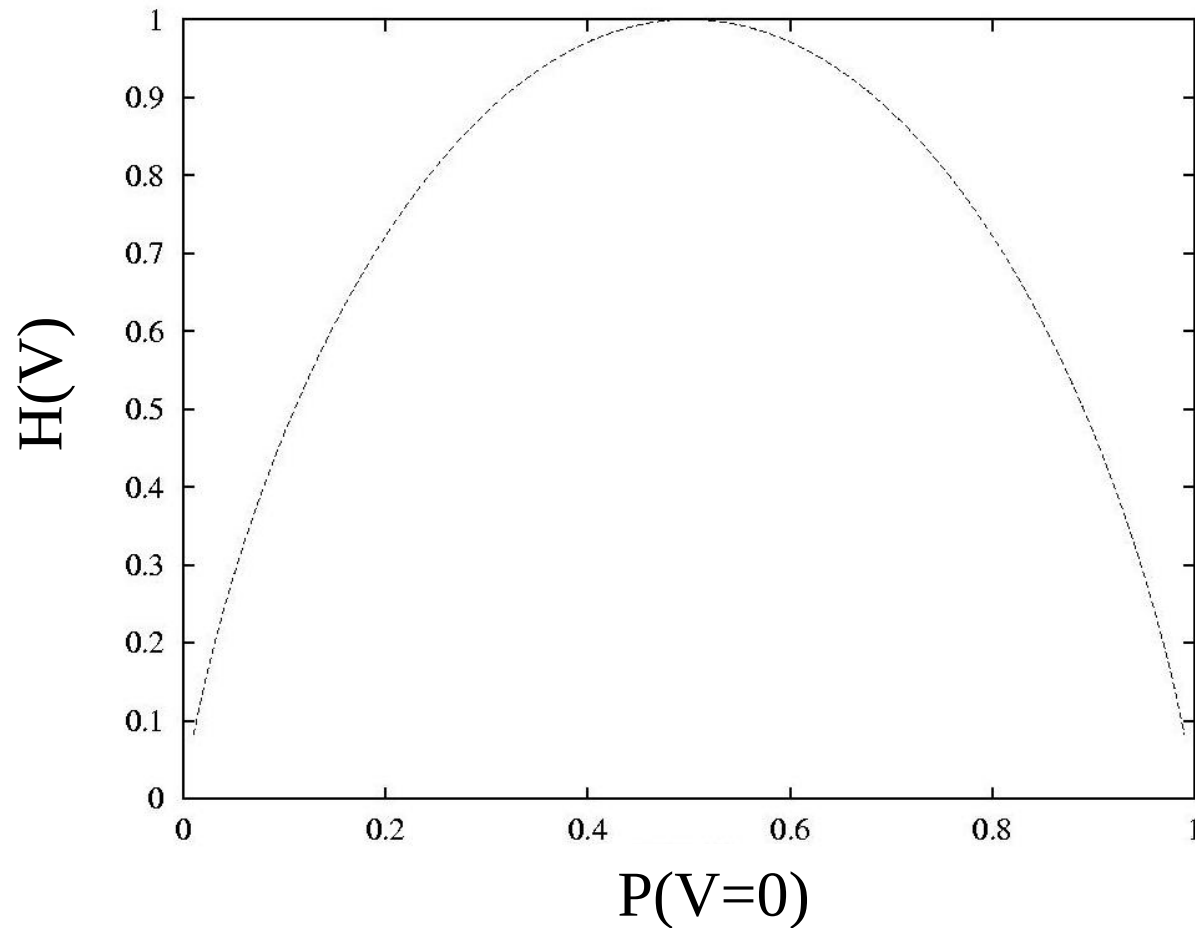
Information in an answer when prior is $\langle P_1, \dots, P_n \rangle$ is

$$H(\langle P_1, \dots, P_n \rangle) = \sum_{i=1}^n -P_i \log_2 P_i$$

(also called **entropy** of the prior)

Entropy

The entropy $H(V)$ of a Boolean random variable V as the probability of $V = 0$ varies from 0 to 1



Using Information

Suppose we have p positive and n negative examples at the root

$\Rightarrow H(\langle p/(p+n), n/(p+n) \rangle)$ bits needed to classify a new example

E.g., for 12 restaurant examples, $p = n = 6$ so we need 1 bit

An attribute splits the examples E into subsets E_i , each of which (we hope) needs less information to complete the classification

Let E_i have p_i positive and n_i negative examples

$\Rightarrow H(\langle p_i/(p_i+n_i), n_i/(p_i+n_i) \rangle)$ bits needed to classify a new example

$\Rightarrow$ **expected** number of bits per example over all branches is

$$\sum_i \frac{p_i + n_i}{p + n} H(\langle p_i/(p_i + n_i), n_i/(p_i + n_i) \rangle)$$

For *Patrons?*, this is 0.459 bits, for *Type* this is (still) 1 bit

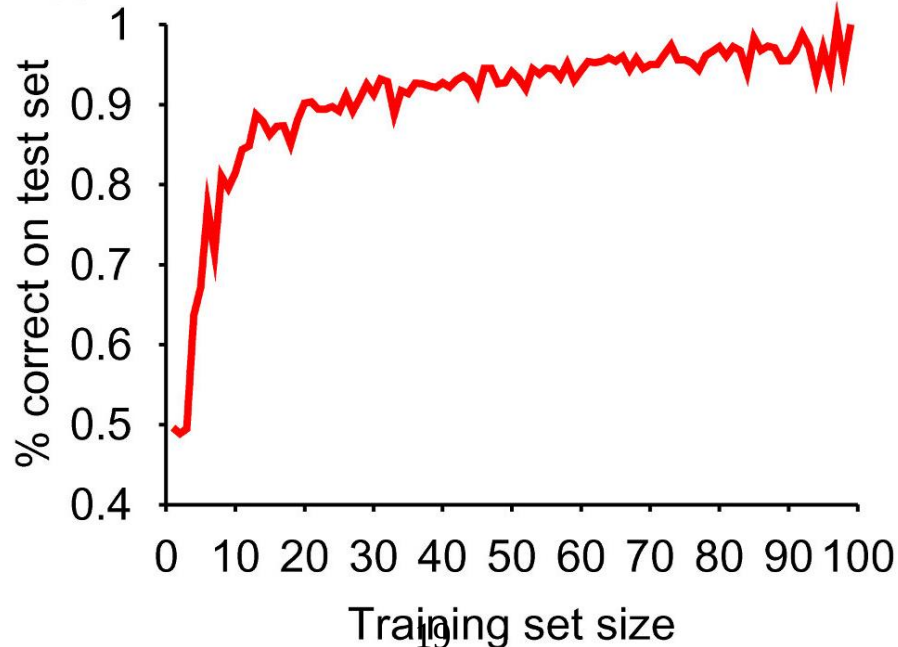
$\Rightarrow$ choose the attribute that minimizes the remaining information needed

Measuring Performance

How do we know that $h \approx f$? (Hume's **Problem of Induction**)

- 1) Use theorems of computational/statistical learning theory
- 2) Try h on a new **test set** of examples
(use **same distribution over example space** as training set)

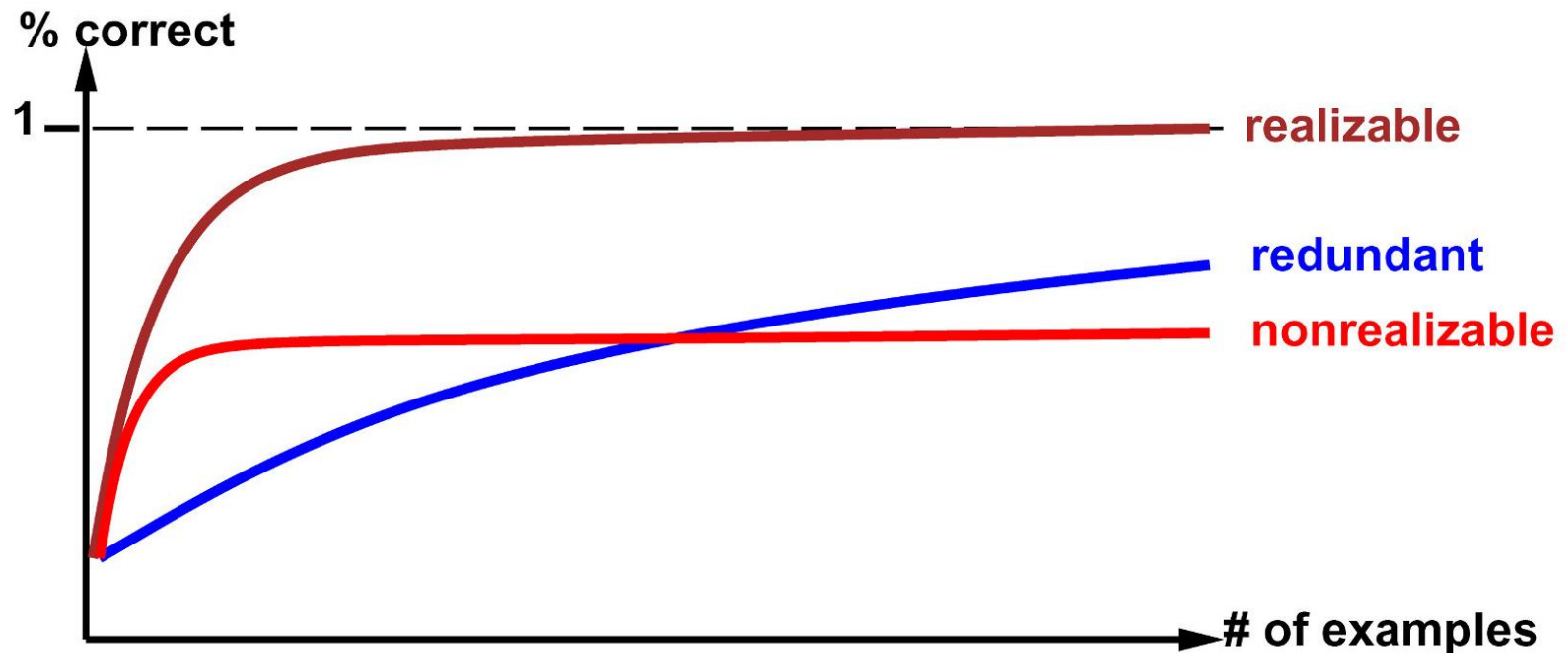
Learning curve = % correct on test set as a function of training set size



What the learning curve tells us

Learning curve depends on

- **realizable** (can express target function) vs. **non-realizable**
non-realizability can be due to missing attributes
or restricted hypothesis class (e.g., thresholded linear function)
- redundant expressiveness (e.g., loads of irrelevant attributes)

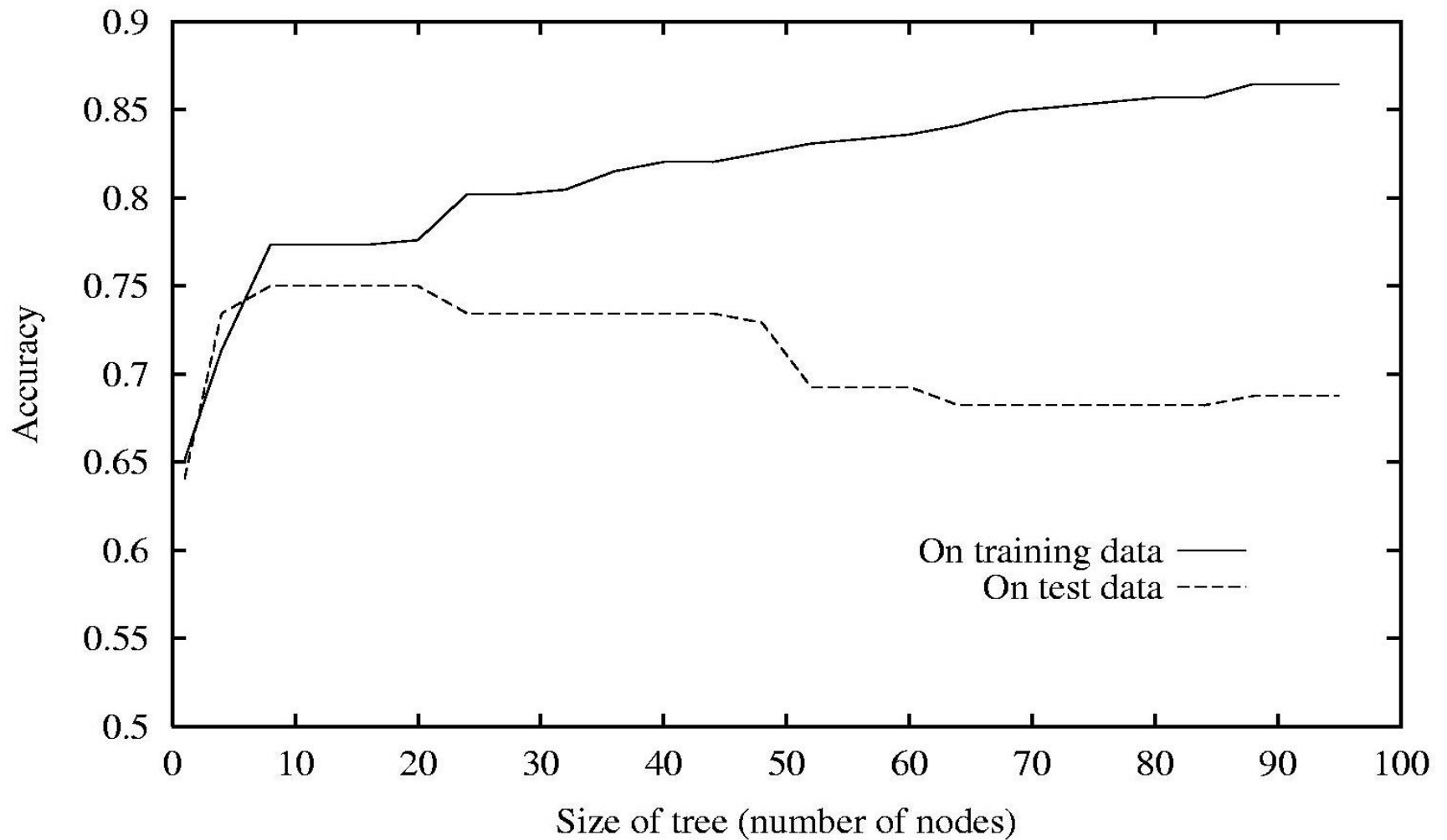


Rule #2 of Machine Learning

The *best* hypothesis almost never achieves 100% accuracy on the training data.

(Rule #1 was: you can't learn anything without inductive bias)

Overfitting



Overfitting is due to “noise”

- Sources of noise:
 - Erroneous training data
 - concept variable incorrect (annotator error)
 - Attributes mis-measured
 - Much more significant:
 - Irrelevant attributes
 - Target function not deterministic in attributes

Irrelevant attributes

- If many attributes are noisy, information gains can be spurious, e.g.:
 - 20 noisy attributes
 - 10 training examples
 - Expected # of different depth-3 trees that split the training data perfectly using *only* noisy attributes:
13.4

Non-determinism

- In general:
 - We can't measure all the variables we need to do perfect prediction.
 - => Target function is not uniquely determined by attribute values

Non-determinism: Example

Humidity	EnjoySport
0.90	0
0.87	1
0.80	0
0.75	0
0.70	1
0.69	1
0.65	1
0.63	1

Decent hypothesis:

Humidity $> 0.70 \rightarrow$ No
Otherwise $\rightarrow$ Yes

Overfit hypothesis:

Humidity $> 0.89 \rightarrow$ No
Humidity > 0.80
 $\wedge$ Humidity $\leq 0.89 \rightarrow$ Yes
Humidity > 0.70
 $\wedge$ Humidity $\leq 0.80 \rightarrow$ No
Humidity $\leq 0.70 \rightarrow$ Yes

Avoiding Overfitting

- Approaches
 - Stop splitting when information gain is low or when split is not statistically significant.
 - Grow full tree and then **prune** it when done
- How to pick the “best” tree?
 - Performance on training data?
 - Performance on validation data?
 - Complexity penalty?

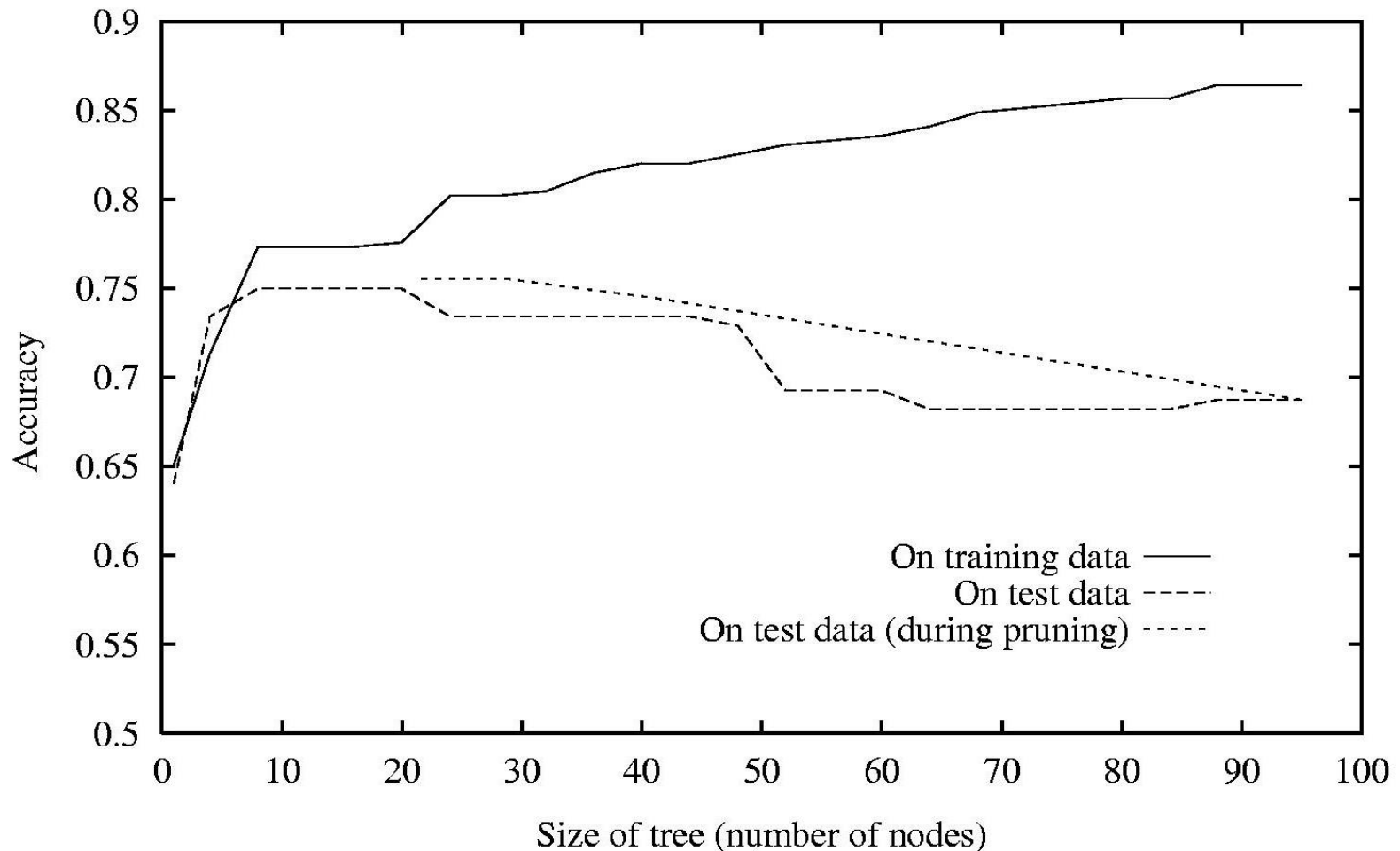
Reduced-Error Pruning

Split data into *training* and *validation* set

Do until further pruning is harmful:

1. Evaluate impact on *validation* set of pruning each possible node (plus those below it)
2. Greedily remove the one that most improves *validation* set accuracy

Effect of Reduced Error Pruning



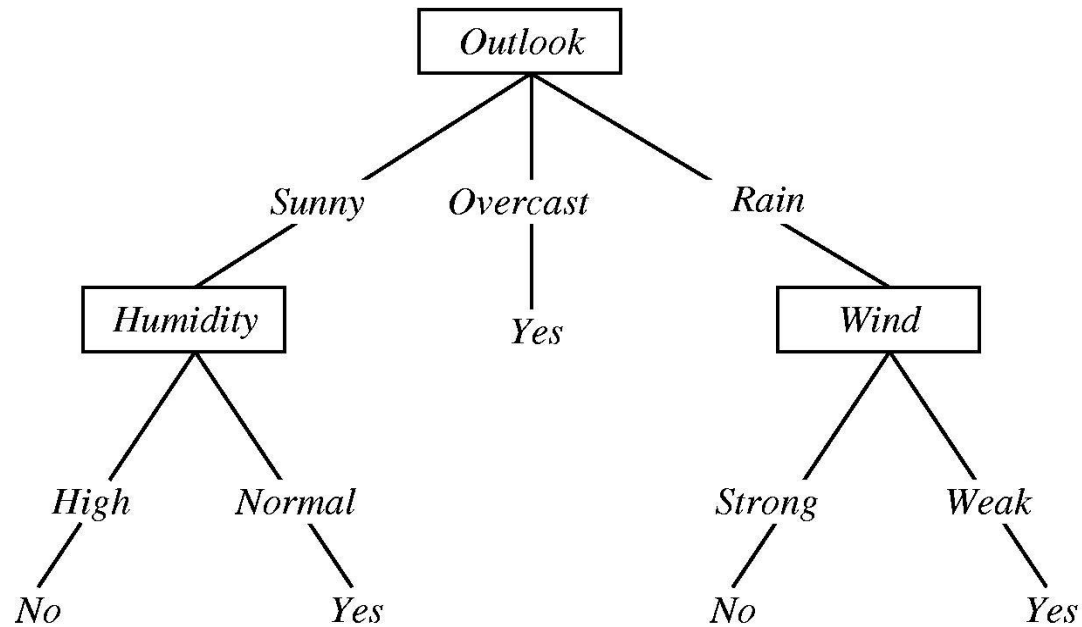
C4.5 Algorithm

- Builds a decision tree from labeled training data
- Also by Ross Quinlan
- Generalizes ID3 by
 - Allowing continuous value attributes
 - Allows missing attributes in examples
 - Prunes tree after building to improve generality

Rule post pruning

- Used in C4.5
- Steps
 1. Build the decision tree
 2. Convert it to a set of logical rules
 3. Prune each rule independently
 4. Sort rules into desired sequence for use

Converting A Tree to Rules



IF (*Outlook = Sunny*) *AND* (*Humidity = High*)
THEN *PlayTennis = No*

IF (*Outlook = Sunny*) *AND* (*Humidity = Normal*)
THEN *PlayTennis = Yes*

...

Scaling Up

- ID3, C4.5, etc. assume data fits in main memory
(OK for up to hundreds of thousands of examples)
- SPRINT, SLIQ: multiple sequential scans of data
(OK for up to millions of examples)
- VFDT: at most one sequential scan
(OK for up to billions of examples)

Unknown Attribute Values

What if some examples are missing values of A ?

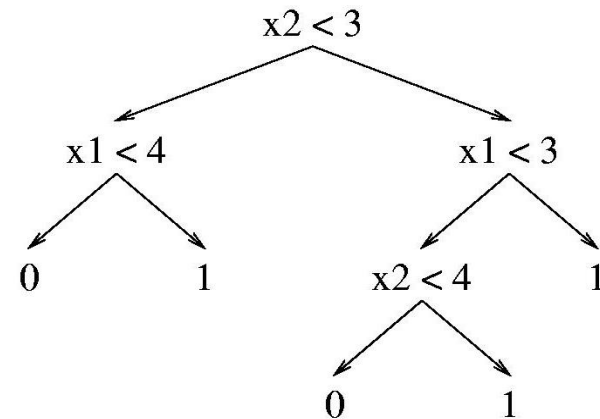
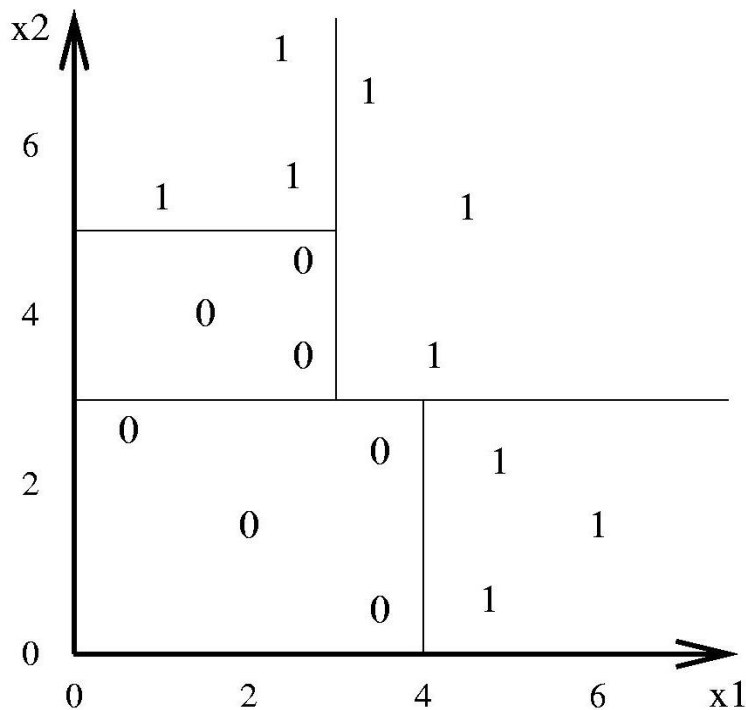
Use training example anyway, sort through tree

- If node n tests A , assign most common value of A among other examples sorted to node n
- Assign most common value of A among other examples with same target value
- Assign probability p_i to each possible value v_i of A
Assign fraction p_i of example to each descendant in tree

Classify new examples in same fashion

Decision Tree Boundaries

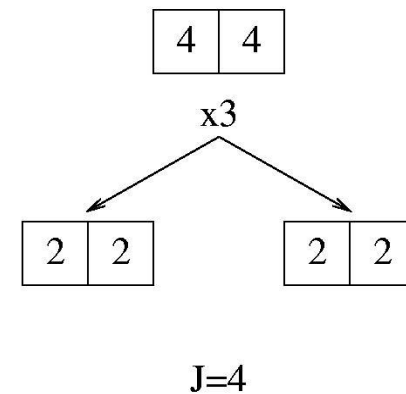
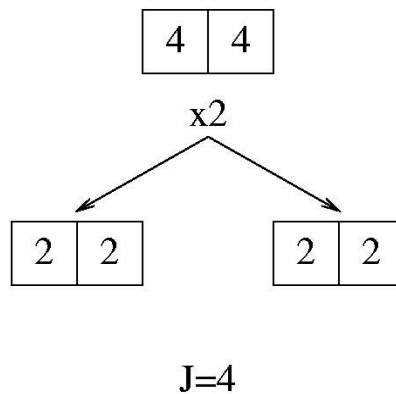
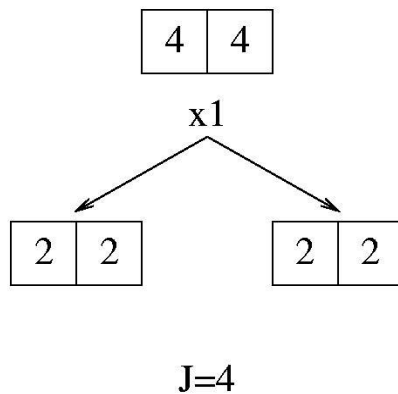
Decision trees divide the feature space into axis-parallel rectangles, and label each rectangle with one of the K classes.



Learning Parity with Noise

When learning exclusive-or (2-bit parity), all splits look equally good. If extra random boolean features are included, they also look equally good. Hence, decision tree algorithms cannot distinguish random noisy features from parity features.

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



Decision Trees Inductive Bias

- How to solve 2-bit parity:
 - Two step look-ahead, or
 - Split on *pairs* of attributes at once
- For k -bit parity, why not just do k -step look ahead?
Or split on k attribute values?

=> Parity functions are the “victims” of the decision tree’s inductive bias.

Take away about decision trees

- Used as classifiers
- Supervised learning algorithms (ID3, C4.5)
- (mostly) Batch processing
- Good for situations where
 - The classification categories are finite
 - The data can be represented as vectors of attributes