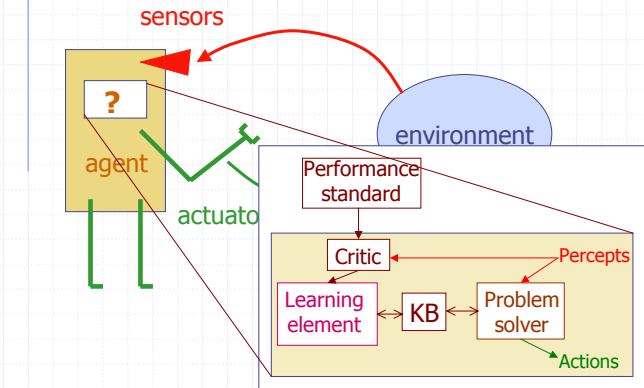


Learning - Decision Trees

Russell and Norvig:
Chapter 18, Sections 18.1 through 18.4
CMSC 421 – Fall 2002

material from Jean-Claude Latombe
and Daphne Koller

Learning Agent



Types of Learning

- ◆ Supervised Learning - classification, prediction
- ◆ Unsupervised Learning – clustering, segmentation, pattern discovery
- ◆ Reinforcement Learning – learning MDPs, online learning

Supervised Learning

- ◆ A general framework
- ◆ Logic-based/discrete learning:
 - learn a function $f(X) \rightarrow (0,1)$
 - Decision trees
 - Version space method
- ◆ Probabilistic/Numeric learning
 - learn a function $f(X) \rightarrow \mathbf{R}$
 - Neural nets

Supervised Learning

- ◆ Someone gives you a bunch of examples, telling you what each one is
- ◆ Eventually, you figure out the mapping from properties (**features**) of the examples and their type

Logic-Inference Formulation

Unlike in the function-learning formulation, h must be a logical sentence, but its inference may benefit from the background knowledge

- ◆ Inductive inference: Find h (inductive hypothesis) such that
 - KB and h are consistent
 - $KB, h \models \Delta$

Note that $h = \Delta$ is a trivial, but uninteresting solution (data caching)

Rewarded Card Example

- ◆ Deck of cards, with each card designated by $[r,s]$, its rank and suit, and some cards "rewarded"
- ◆ **Background knowledge KB:**
 - $((r=1) \vee \dots \vee (r=10)) \Leftrightarrow \text{NUM}(r)$
 - $((r=J) \vee (r=Q) \vee (r=K)) \Leftrightarrow \text{FACE}(r)$
 - $((s=S) \vee (s=C)) \Leftrightarrow \text{BLACK}(s)$
 - $((s=D) \vee (s=H)) \Leftrightarrow \text{RED}(s)$
- ◆ **Training set Δ :**

$$\text{REWARD}([4,C]) \wedge \text{REWARD}([7,C]) \wedge \text{REWARD}([2,S]) \wedge \neg \text{REWARD}([5,H]) \wedge \neg \text{REWARD}([J,S])$$

Rewarded Card Example

- ◆ **Background knowledge KB:**
 - $((r=1) \vee \dots \vee (r=10)) \Leftrightarrow \text{NUM}(r)$
 - $((r=J) \vee (r=Q) \vee (r=K)) \Leftrightarrow \text{FACE}(r)$
 - $((s=S) \vee (s=C)) \Leftrightarrow \text{BLACK}(s)$
 - $((s=D) \vee (s=H)) \Leftrightarrow \text{RED}(s)$
- ◆ **Training set Δ :**

$$\text{REWARD}([4,C]) \wedge \text{REWARD}([7,C]) \wedge \text{REWARD}([2,S]) \wedge \neg \text{REWARD}([5,H]) \wedge \neg \text{REWARD}([J,S])$$
- ◆ **Possible hypothesis:**

$$h \equiv (\text{NUM}(r) \wedge \text{BLACK}(s) \Leftrightarrow \text{REWARD}([r,s]))$$

There are several possible inductive hypotheses

Learning a Predicate

- ◆ Set E of objects (e.g., cards)
- ◆ Goal predicate $\text{CONCEPT}(x)$, where x is an object in E, that takes the value True or False (e.g., REWARD)
- ◆ Observable predicates $A(x)$, $B(x)$, ... (e.g., NUM, RED)
- ◆ Training set: values of CONCEPT for some combinations of values of the observable predicates

A Possible Training Set

Ex. #	A	B	C	D	E	CONCEPT
1	True	True	False	True	False	False
2	True	False	False	False	False	True
3	False	False	True	True	True	False
4	True	True	True	False	True	True
5	False	True	True	False	False	False
6	True	True	False	True	True	False
7	False	False	True	False	True	False
8	True	False	True	False	True	True

Note that the training set does not say whether an observable predicate A, ..., E is pertinent or not

Learning a Predicate

- ◆ Set E of objects (e.g., cards)
- ◆ Goal predicate $\text{CONCEPT}(x)$, where x is an object in E, that takes the value True or False (e.g., REWARD)
- ◆ Observable predicates $A(x)$, $B(x)$, ... (e.g., NUM, RED)
- ◆ Training set: values of CONCEPT for some combinations of values of the observable predicates
- ◆ Find a representation of CONCEPT in the form:

$$\text{CONCEPT}(x) \Leftrightarrow S(A, B, \dots)$$
 where $S(A, B, \dots)$ is a sentence built with the observable predicates, e.g.:

$$\text{CONCEPT}(x) \Leftrightarrow A(x) \wedge (\neg B(x) \vee C(x))$$

Example set

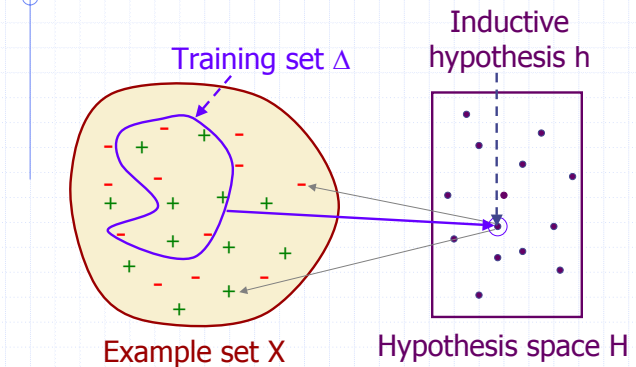
- ◆ An **example** consists of the values of CONCEPT and the observable predicates for some object x
- ◆ A example is **positive** if CONCEPT is True, else it is **negative**
- ◆ The set E of all examples is the **example set**
- ◆ The **training set** is a **subset** of E

a small one!

Hypothesis Space

- ◆ An **hypothesis** is any sentence h of the form:
 $\text{CONCEPT}(x) \Leftrightarrow S(A,B, \dots)$
 where $S(A,B,\dots)$ is a sentence built with the observable predicates
- ◆ The set of all hypotheses is called the **hypothesis space** H
- ◆ An hypothesis h **agrees** with an example if it gives the correct value of **CONCEPT**

Inductive Learning Scheme



Size of the Hypothesis Space

- ◆ n observable predicates
- ◆ 2^n entries in truth table
- ◆ In the absence of any restriction (bias), there are 2^{2^n} hypotheses to choose from
- ◆ $n = 6 \rightarrow 2 \times 10^{19}$ hypotheses!

Multiple Inductive Hypotheses

Rewarded Card Example

◆ Background knowledge KB:
 $((r=1) \vee \dots \vee (r=10)) \Leftrightarrow \text{NUM}([r,s])$

Need for a system of preferences – called a bias – to compare possible hypotheses

◆ Possible inductive hypothesis:
 $h = (\text{NUM}(x) \wedge \text{BLACK}(x) \Leftrightarrow \text{REWARD}(x))$

$h1 \equiv \text{NUM}(x) \wedge \text{BLACK}(x) \Leftrightarrow \text{REWARD}(x)$

$h2 \equiv \text{BLACK}([r,s]) \wedge \neg(r=J) \Leftrightarrow \text{REWARD}([r,s])$

$h3 \equiv ([r,s]=[4,C]) \wedge ([r,s]=[7,C]) \wedge [r,s]=[2,S]) \wedge \neg([r,s]=[5,H]) \wedge \neg([r,s]=[J,S]) \Leftrightarrow \text{REWARD}([r,s])$

agree with all the examples in the training set

Keep-It-Simple (KIS) Bias

◆ Motivation

If the bias allows only sentences S that are conjunctions of $k \ll n$ predicates picked from the n observable predicates, then the size of

◆ H is $O(n^k)$

- Use much fewer observable predicates than suggested by the training set
- Constrain the learnt predicate, e.g., to use only "high-level" observable predicates such as NUM, FACE, BLACK, and RED and/or to be a conjunction of literals

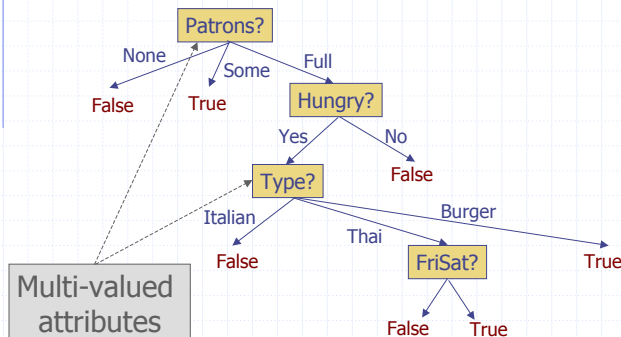
Predicate-Learning Methods

◆ Decision tree

◆ Version space

Decision Tree

WillWait predicate (Russell and Norvig)



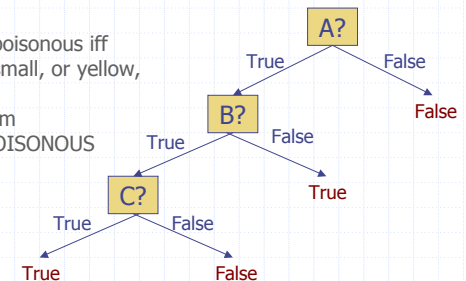
Predicate as a Decision Tree

The predicate $\text{CONCEPT}(x) \Leftrightarrow A(x) \wedge (\neg B(x) \vee C(x))$ can be represented by the following decision tree:

Example:

A mushroom is poisonous iff it is yellow and small, or yellow, big and spotted

- x is a mushroom
- $\text{CONCEPT} = \text{POISONOUS}$
- $A = \text{YELLOW}$
- $B = \text{BIG}$
- $C = \text{SPOTTED}$



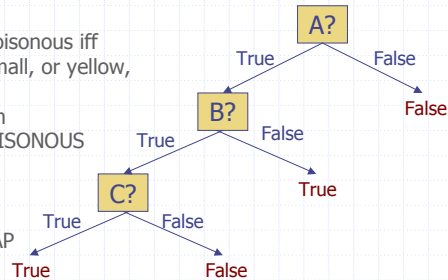
Predicate as a Decision Tree

The predicate $\text{CONCEPT}(x) \Leftrightarrow A(x) \wedge (\neg B(x) \vee C(x))$ can be represented by the following decision tree:

Example:

A mushroom is poisonous iff it is yellow and small, or yellow, big and spotted

- x is a mushroom
- $\text{CONCEPT} = \text{POISONOUS}$
- $A = \text{YELLOW}$
- $B = \text{BIG}$
- $C = \text{SPOTTED}$
- $D = \text{FUNNEL-CAP}$
- $E = \text{BULKY}$

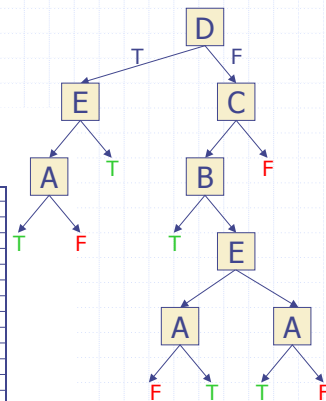


Training Set

Ex. #	A	B	C	D	E	CONCEPT
1	False	False	True	False	True	False
2	False	True	False	False	False	False
3	False	True	True	True	True	False
4	False	False	True	False	False	False
5	False	False	False	True	True	False
6	True	False	True	False	False	True
7	True	False	False	True	False	True
8	True	False	True	False	True	True
9	True	True	True	False	True	True
10	True	True	True	True	True	True
11	True	True	False	False	False	False
12	True	True	False	False	True	False
13	True	False	True	True	True	True

Possible Decision Tree

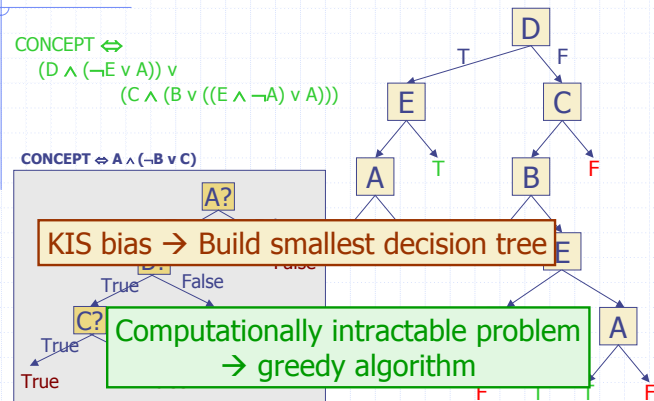
Ex. #	A	B	C	D	E	CONCEPT
1	False	False	True	False	True	False
2	False	True	False	False	False	False
3	False	True	True	True	True	False
4	False	False	True	False	False	False
5	False	False	False	True	True	False
6	True	False	True	False	False	True
7	True	False	False	True	False	True
8	True	False	True	False	True	True
9	True	True	True	False	True	True
10	True	True	True	True	True	True
11	True	True	False	False	False	False
12	True	True	False	False	True	False
13	True	False	True	True	True	True



Possible Decision Tree

$$\text{CONCEPT} \Leftrightarrow (D \wedge (\neg E \vee A)) \vee (C \wedge (B \vee ((E \wedge \neg A) \vee A)))$$

$$\text{CONCEPT} \Leftrightarrow A \wedge (\neg B \vee C)$$



Getting Started

The distribution of the training set is:

True: 6, 7, 8, 9, 10, 13

False: 1, 2, 3, 4, 5, 11, 12

Getting Started

The distribution of training set is:

True: 6, 7, 8, 9, 10, 13

False: 1, 2, 3, 4, 5, 11, 12

Without testing any observable predicate, we could predict that CONCEPT is False (majority rule) with an estimated probability of error $P(E) = 6/13$

Getting Started

The distribution of training set is:

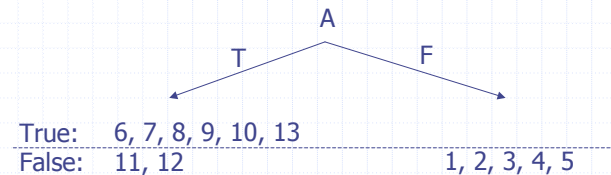
True: 6, 7, 8, 9, 10, 13

False: 1, 2, 3, 4, 5, 11, 12

Without testing any observable predicate, we could report that CONCEPT is False (majority rule) with an estimated probability of error $P(E) = 6/13$

Assuming that we will only include one observable predicate in the decision tree, which predicate should we test to minimize the probability of error?

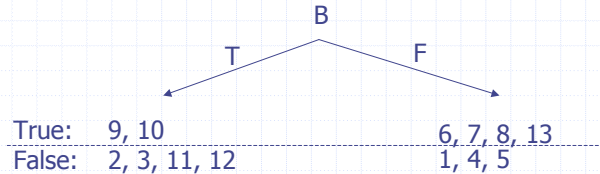
Assume It's A



If we test only A, we will report that CONCEPT is True if A is True (majority rule) and False otherwise

The estimated probability of error is:
 $\Pr(E) = (8/13) \times (2/8) + (5/13) \times 0 = 2/13$

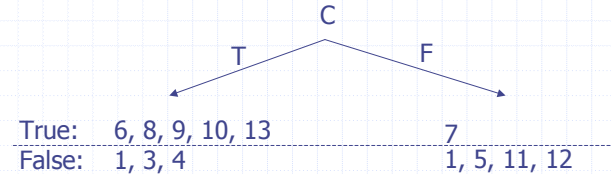
Assume It's B



If we test only B, we will report that CONCEPT is False if B is True and True otherwise

The estimated probability of error is:
 $\Pr(E) = (6/13) \times (2/6) + (7/13) \times (3/7) = 5/13$

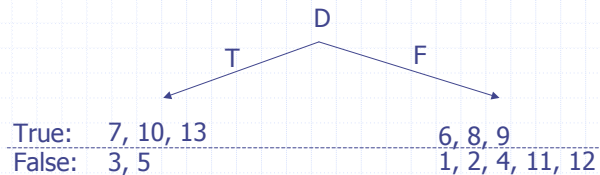
Assume It's C



If we test only C, we will report that CONCEPT is True if C is True and False otherwise

The estimated probability of error is:
 $\Pr(E) = (8/13) \times (3/8) + (5/13) \times (1/5) = 4/13$

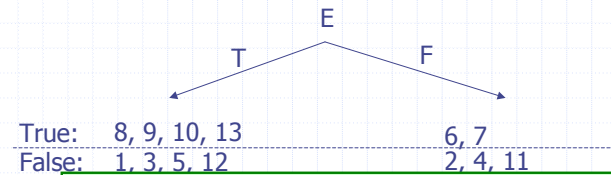
Assume It's D



If we test only D, we will report that CONCEPT is True if D is True and False otherwise

The estimated probability of error is:
 $\Pr(E) = (5/13) \times (2/5) + (8/13) \times (3/8) = 5/13$

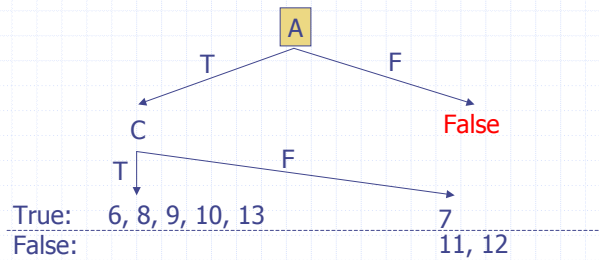
Assume It's E



If we test only E, we will report that CONCEPT is True if E is True and False otherwise. **So, the best predicate to test is A**

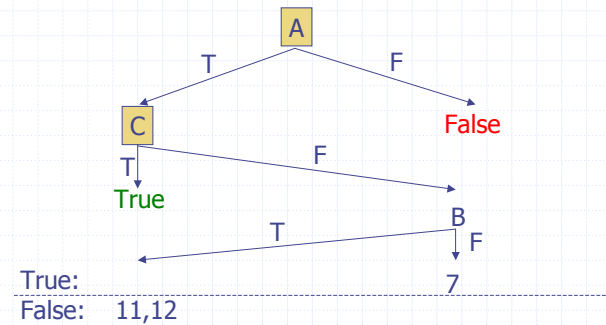
The estimated probability of error is unchanged:
 $\Pr(E) = (8/13) \times (4/8) + (5/13) \times (2/5) = 6/13$

Choice of Second Predicate

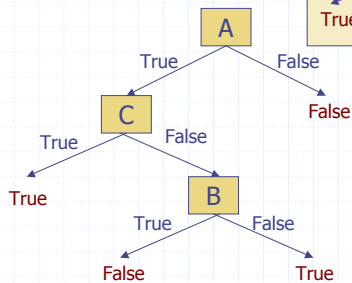


The majority rule gives the probability of error $\Pr(E) = 1/8$

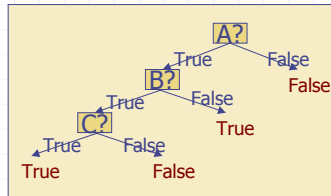
Choice of Third Predicate



Final Tree



$$L \equiv \text{CONCEPT} \Leftrightarrow A \wedge (C \vee \neg B)$$



Learning a Decision Tree

$\text{DTL}(\Delta, \text{Predicates})$

1. If all examples in Δ are positive then return True
2. If all examples in Δ are negative then return False
3. If Predicates is empty then return failure
4. $A \leftarrow$ most discriminating predicate in Predicates
5. Return the tree whose:
 - root is A,
 - left branch is $\text{DTL}(\Delta^+ A, \text{Predicates}-A)$,
 - right branch is $\text{DTL}(\Delta^- A, \text{Predicates}-A)$

Subset of examples that satisfy A

Using Information Theory

- ◆ Rather than minimizing the probability of error, most existing learning procedures try to minimize the expected number of questions needed to decide if an object x satisfies CONCEPT
- ◆ This minimization is based on a measure of the "quantity of information" that is contained in the truth value of an observable predicate

of Questions to Identify an Object

- ◆ Let U be a set of size $|U|$
- ◆ We want to identify any particular object of U with only True/False questions
- ◆ What is the minimum number of questions that will we need on the average?
- ◆ The answer is $\log_2|U|$, since the best we can do at each question is to split the set of remaining objects in half

of Questions to Identify an Object

- ◆ Now, suppose that a question Q splits U into two subsets T and F of sizes $|T|$ and $|F|$
- ◆ What is the minimum number of questions that will we need on the average, assuming that we will ask Q first?

of Questions to Identify an Object

- ◆ Now, suppose that a question Q splits U into two subsets T and F of sizes $|T|$ and $|F|$
- ◆ What is the minimum average number of questions that will we need assuming that we will ask Q first?
- ◆ The answer is:
$$(|T|/|U|)\log_2|T| + (|F|/|U|)\log_2|F|$$

Information Content of an Answer

- ◆ The number of questions saved by asking Q is:

$$I_Q = \log_2|U| - (|T|/|U|)\log_2|T| + (|F|/|U|)\log_2|F|$$
 which is called the **information content** of the answer to Q
- ◆ Posing $p_T = |T|/|U|$ and $p_F = |F|/|U|$, we get:

$$I_Q = \log_2|U| - p_T\log_2(p_T|U|) - p_F\log_2(p_F|U|)$$
- ◆ Since $p_T + p_F = 1$, we have:

$$I_Q = -p_T\log_2p_T - p_F\log_2p_F = \mathbf{I}(p_T, p_F) \leq 1$$

Application to Decision Tree

- ◆ In a decision tree we are not interested in identifying a particular object from a set $U=\Delta$, but in determining if a certain object x verifies or contradicts CONCEPT
- ◆ Let us divide Δ into two subsets:
 - Δ^+ : the positive examples
 - Δ^- : the negative examples
- ◆ Let $p = |\Delta^+|/|\Delta|$ and $q = 1-p$

Application to Decision Tree

- ◆ In a decision tree we are not interested in identifying a particular object from a set Δ , but in determining if a certain object x verifies or contradicts a predicate CONCEPT
- ◆ Let us divide Δ into two subsets:
 - Δ^+ : the positive examples
 - Δ^- : the negative examples
- ◆ Let $p = |\Delta^+|/|\Delta|$ and $q = 1-p$
- ◆ The information content of the answer to the question CONCEPT(x)? would be:

$$I_{\text{CONCEPT}} = \mathbf{I}(p, q) = -p\log_2p - q\log_2q$$

Application to Decision Tree

- ◆ Instead, we can ask $A(x)$? where A is an observable predicate
- ◆ The answer to $A(x)$? divides Δ into two subsets Δ^{+A} and Δ^{-A}
- ◆ Let p_1 be the ratio of objects that verify CONCEPT in Δ^{+A} , and $q_1 = 1-p_1$
- ◆ Let p_2 be the ratio of objects that verify CONCEPT in Δ^{-A} , and $q_2 = 1-p_2$

Application to Decision Tree

- At each recursion, the learning procedure includes in the decision tree the observable predicate that maximizes the gain of information

$$I_{\text{CONCEPT}} - (|\Delta^+|/|\Delta|) I(p_1, q_1) + (|\Delta^-|/|\Delta|) I(p_2, q_2)$$

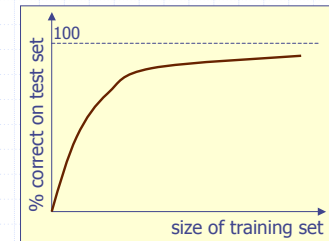
and This predicate is the most discriminating

- The expected information content of the answer to the question CONCEPT(x)? would then be:

$$(|\Delta^+|/|\Delta|) I(p_1, q_1) + (|\Delta^-|/|\Delta|) I(p_2, q_2) \leq I_{\text{CONCEPT}}$$

Miscellaneous Issues

- Assessing performance:
 - Training set and test set
 - Learning curve



Typical learning curve

Miscellaneous Issues

- Assessing performance:

- Training set and test set
- Learning curve

- Overfitting

- Tree pruning
- Cross-validation

Risk of using irrelevant

The resulting decision tree + majority rule may not classify correctly all examples in the training set

Terminate recursion when information gain is too small

Miscellaneous Issues

- Assessing performance:
 - Training set and test set
 - Learning curve
- Overfitting
 - Tree pruning
 - Cross-validation
- Missing data

The value of an observable predicate P is unknown for an example x. Then construct a decision tree for both values of P and select the value that ends up classifying x in the largest class

Miscellaneous Issues

- ◆ Assessing performance:
 - Training set and test set
 - Learning curve
 - ◆ Overfitting
 - Tree pruning
 - Cross-validation
 - ◆ Missing data
 - ◆ Multi-valued and continuous attributes
- Select threshold that maximizes information gain

These issues occur with virtually any learning method

Applications of Decision Tree

- ◆ Medical diagnostic
- ◆ Evaluation of geological systems for assessing gas and oil basins
- ◆ Early detection of problems (e.g., jamming) during oil drilling operations
- ◆ Automatic generation of rules in expert systems

Applications of Decision Tree

- ◆ SGI flight simulator
- ◆ predicting emergency C sections
 - identified new class of high risk patients
- ◆ SKICAT – classifying stars and galaxies from telescope images
 - 40 attributes
 - 8 levels deep
 - could correctly classify images that were too faint for human to classify
 - 16 new high red-shift quasars discovered in at least one order of magnitude less observation time

Summary

- ◆ Inductive learning frameworks
- ◆ Logic inference formulation
- ◆ Hypothesis space and KIS bias
- ◆ Inductive learning of decision trees
- ◆ Using information theory
- ◆ Assessing performance
- ◆ Overfitting

Learning II: Neural Networks

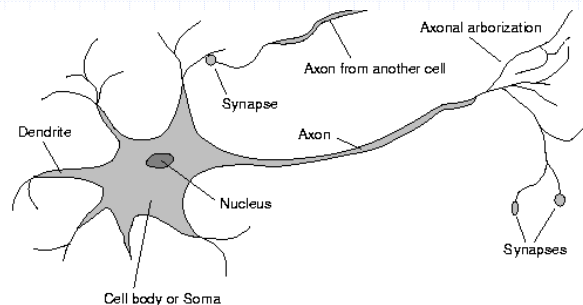
R&N: ch 19

based on material from
Marie desJardins, Ray
Mooney, Daphne Koller

Neural function

- ◆ Brain function (thought) occurs as the result of the firing of **neurons**
- ◆ Neurons connect to each other through **synapses**, which propagate **action potential** (electrical impulses) by releasing **neurotransmitters**
- ◆ Synapses can be **excitatory** (potential-increasing) or **inhibitory** (potential-decreasing), and have varying **activation thresholds**
- ◆ Learning occurs as a result of the synapses' **plasticity**: They exhibit long-term changes in connection strength
- ◆ There are about 10^{11} neurons and about 10^{14} synapses in the human brain

Biology of a neuron



Brain structure

- ◆ Different areas of the brain have different functions
 - Some areas seem to have the same function in all humans (e.g., Broca's region); the overall layout is generally consistent
 - Some areas are more plastic, and vary in their function; also, the lower-level structure and function vary greatly
- ◆ We don't know how different functions are "assigned" or acquired
 - Partly the result of the physical layout / connection to inputs (sensors) and outputs (effectors)
 - Partly the result of experience (learning)
- ◆ We *really* don't understand how this neural structure leads to what we perceive as "consciousness" or "thought"
- ◆ Our neural networks are not nearly as complex or intricate as the actual brain structure

Comparison of computing power

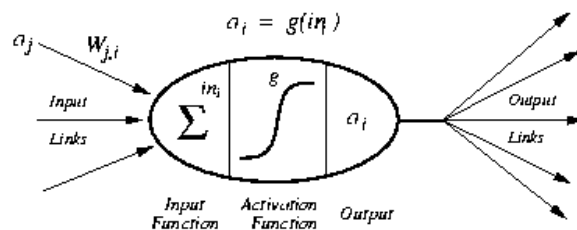
	Computer	Human Brain
Computational units	1 CPU, 10^8 gates	10^{11} neurons
Storage units	10^8 bits RAM, 10^{10} bits disk	10^{11} neurons, 10^{14} synapses
Cycle time	10^{-8} sec	10^{-3} sec
Bandwidth	10^8 bits/sec	10^{14} bits/sec
Neuron updates/sec	10^5	10^{14}

- ◆ Computers are way faster than neurons...
- ◆ But there are a lot more neurons than we can reasonably model in modern digital computers, and they all fire in parallel
- ◆ Neural networks are designed to be massively parallel
- ◆ The brain is effectively a billion times faster

Neural networks

- ◆ Neural networks are made up of **nodes** or **units**, connected by **links**
- ◆ Each link has an associated **weight** and **activation level**
- ◆ Each node has an **input function** (typically summing over weighted inputs), an **activation function**, and an **output**

Neural unit



Model Neuron

- ◆ Neuron modeled a unit j
- ◆ weights on input unit i to j , w_{ji}
- ◆ net input to unit j is:

$$\text{net}_j = \sum_i w_{ij} \cdot o_i$$

- ◆ threshold T_j
- ◆ o_j is 1 if $\text{net}_j > T_j$

Neural Computation

- ◆ McCollough and Pitt (1943) showed how LTU can be used to compute logical functions
 - AND?
 - OR?
 - NOT?
- ◆ Two layers of LTUs can represent any boolean function

Learning Rules

- ◆ Rosenblatt (1959) suggested that if a target output value is provided for a single neuron with fixed inputs, it can incrementally change weights to learn to produce these outputs using the *perceptron learning rule*
 - assumes binary valued input/outputs
 - assumes a single linear threshold unit

Perceptron Learning rule

- ◆ If the target output for unit j is t_j
$$w_{ji} = w_{ji} + \eta(t_j - o_j)o_i$$
- ◆ Equivalent to the intuitive rules:
 - ◆ If output is correct, don't change the weights
 - ◆ If output is low ($o_j=0, t_j=1$), increment weights for all the inputs which are 1
 - ◆ If output is high ($o_j=1, t_j=0$), decrement weights for all inputs which are 1
 - ◆ Must also adjust threshold. Or equivalently assume there is a weight w_{j0} for an extra input unit that has $o_0=1$

Perceptron Learning Algorithm

- ◆ Repeatedly iterate through examples adjusting weights according to the perceptron learning rule until all outputs are correct
 - Initialize the weights to all zero (or random)
 - Until outputs for all training examples are correct
 - for each training example e do
 - compute the current output o_j
 - compare it to the target t_j and update weights
- ◆ each execution of outer loop is an epoch
- ◆ for multiple category problems, learn a separate perceptron for each category and assign to the class whose perceptron most exceeds its threshold
- ◆ Q: when will the algorithm terminate?

Perceptron Video

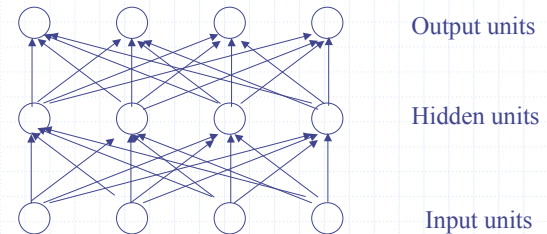
Representation Limitations of a Perceptron

- ◆ Perceptrons can only represent linear threshold functions and can therefore only learn functions which linearly separate the data, I.e. the positive and negative examples are separable by a hyperplane in n -dimensional space

Perceptron Learnability

- ◆ **Perceptron Convergence Theorem:**
If there are a set of weights that are consistent with the training data (I.e. the data is linearly separable), the perceptron learning algorithm will converge (Minicksy & Papert, 1969)
- ◆ Unfortunately, many functions (like parity) cannot be represented by LTU

Layered feed-forward network



"Executing" neural networks

- ◆ Input units are set by some exterior function (think of these as sensors), which causes their output links to be activated at the specified level
- ◆ Working forward through the network, the input function of each unit is applied to compute the input value
 - Usually this is just the weighted sum of the activation on the links feeding into this node
- ◆ The activation function transforms this input function into a final value
 - Typically this is a nonlinear function, often a sigmoid function corresponding to the "threshold" of that node