

BAYES NETS 4

CH. 14.5: APPROXIMATE INFERENCE

Adapted from slides kindly shared by Stuart Russell

Appreciations

- ◇ More good questions and feedback on Piazza
- ◇ The IETF (Internet Engineering Task Force): an evidence-based standards body!

Share some of yours?

Announcements

Project P1 grades are up on D2L, with extra credit, early bonus, late add-ons, etc.

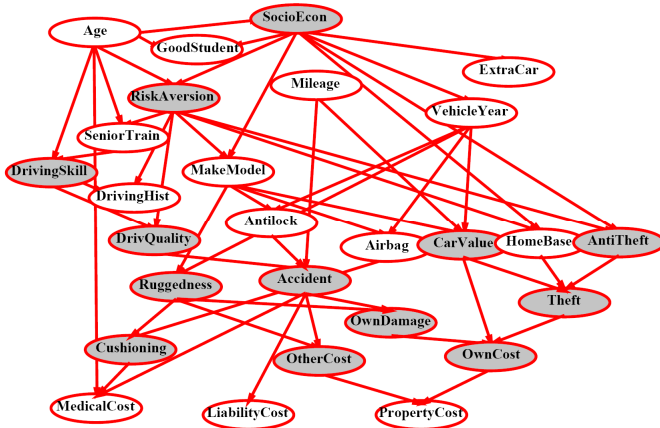
Project P3 Reinforcement due Thu Nov 29th at 17:00

Outline

◇ Bayes Nets: Approximate Inference

Credit to Dan Klein, Stuart Russell and Andrew Moore for most of today's slides

Approximate Inference

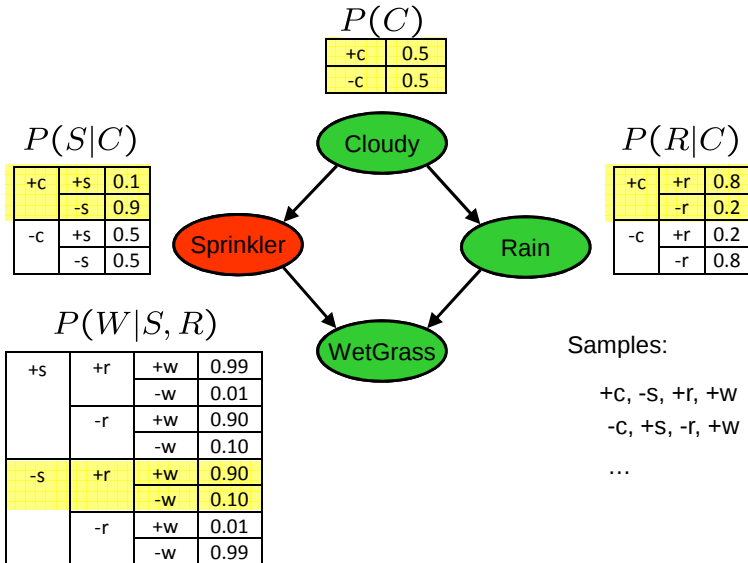


Approximate Inference

- Simulation has a name: sampling
- Sampling is a hot topic in machine learning, and it's really simple
- Basic idea:
 - Draw N samples from a sampling distribution S
 - Compute an approximate posterior probability
 - Show this converges to the true probability P
- Why sample?
 - Learning: get samples from a distribution you don't know
 - Inference: getting a sample is faster than computing the right answer (e.g. with variable elimination)



Prior Sampling



Prior Sampling

- This process generates samples with probability:

$$S_{PS}(x_1 \dots x_n) = \prod_{i=1}^n P(x_i | \text{Parents}(X_i)) = P(x_1 \dots x_n)$$

...i.e. the BN's joint probability

- Let the number of samples of an event be $N_{PS}(x_1 \dots x_n)$
- Then
$$\begin{aligned} \lim_{N \rightarrow \infty} \hat{P}(x_1, \dots, x_n) &= \lim_{N \rightarrow \infty} N_{PS}(x_1, \dots, x_n) / N \\ &= S_{PS}(x_1, \dots, x_n) \\ &= P(x_1 \dots x_n) \end{aligned}$$
- I.e., the sampling procedure is **consistent**

Example

- First: Get a bunch of samples from the BN:

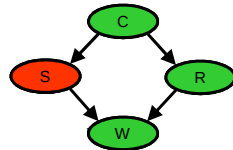
+C, -S, +r, +W

+C, +S, +r, +W

-C, +S, +r, -W

+C, -S, +r, +W

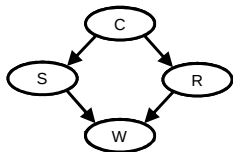
-C, -S, -r, +W



- Example: we want to know $P(W)$
 - We have counts $\langle +w:4, -w:1 \rangle$
 - Normalize to get approximate $P(W) = \langle +w:0.8, -w:0.2 \rangle$
 - This will get closer to the true distribution with more samples
 - Can estimate anything else, too
 - What about $P(C|+w)$? $P(C|+r, +w)$? $P(C|-r, -w)$?
 - Fast: can use fewer samples if less time (what's the drawback?)

Rejection Sampling

- Let's say we want $P(C)$
 - No point keeping all samples around
 - Just tally counts of C as we go
- Let's say we want $P(C|+s)$
 - Same thing: tally C outcomes, but ignore (reject) samples which don't have $S=+s$
 - This is called **rejection sampling**
 - It is also consistent for conditional probabilities (i.e., correct in the limit)



+c, -s, +r, +w
+c, +s, +r, +w
-c, +s, +r, -w
+c, -s, +r, +w
-c, -s, -r, +w

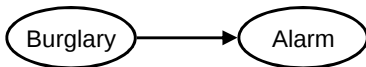
Sampling Example

- There are 2 cups.
 - The first contains 1 penny and 1 quarter
 - The second contains 2 quarters
- Say I pick a cup uniformly at random, then pick a coin randomly from that cup. It's a quarter (yes!). What is the probability that the other coin in that cup is also a quarter?

Likelihood Weighting

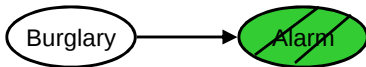
- Problem with rejection sampling:

- If evidence is unlikely, you reject a lot of samples
- You don't exploit your evidence as you sample
- Consider $P(B|+a)$



-b, -a
-b, -a
-b, -a
-b, -a
+b, +a

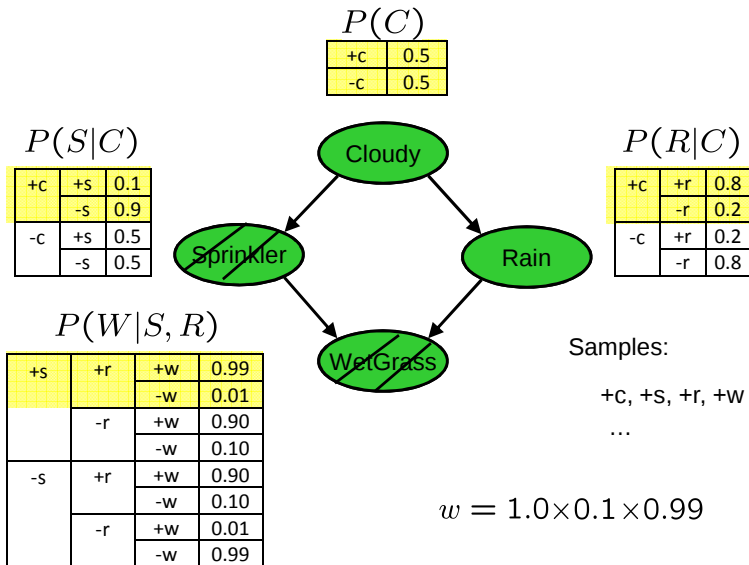
- Idea: fix evidence variables and sample the rest



-b +a
-b, +a
-b, +a
-b, +a
+b, +a

- Problem: sample distribution not consistent!
- Solution: weight by probability of evidence given parents

Likelihood Weighting



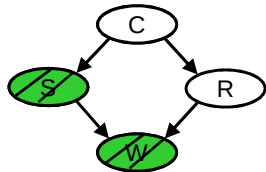
Likelihood Weighting

- Sampling distribution if z sampled and e fixed evidence

$$S_{WS}(z, e) = \prod_{i=1}^l P(z_i | \text{Parents}(Z_i))$$

- Now, samples have weights

$$w(z, e) = \prod_{i=1}^m P(e_i | \text{Parents}(E_i))$$

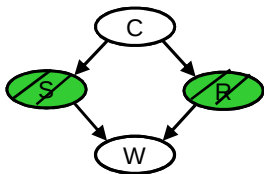


- Together, weighted sampling distribution is consistent

$$\begin{aligned} S_{WS}(z, e) \cdot w(z, e) &= \prod_{i=1}^l P(z_i | \text{Parents}(z_i)) \prod_{i=1}^m P(e_i | \text{Parents}(e_i)) \\ &= P(z, e) \end{aligned}$$

Likelihood Weighting

- Likelihood weighting is good
 - We have taken evidence into account **as we generate the sample**
 - E.g. here, W 's value will get picked based on the evidence values of S , R
 - More of our samples will reflect the state of the world suggested by the evidence
- Likelihood weighting doesn't solve all our problems
 - Evidence influences the choice of downstream variables, but not upstream ones (C isn't more likely to get a value matching the evidence)
- We would like to consider evidence when we sample every variable



Markov Chain Monte Carlo*

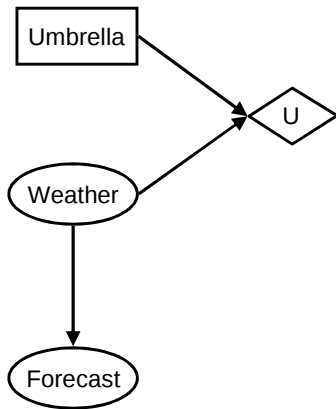
- *Idea*: instead of sampling from scratch, create samples that are each like the last one.
- *Procedure*: resample one variable at a time, conditioned on all the rest, but keep evidence fixed. E.g., for $P(B|+c)$:



- *Properties*: Now samples are not independent (in fact they're nearly identical), but sample averages are still consistent estimators!
- *What's the point*: both upstream and downstream variables condition on evidence.

Decision Networks

- MEU: choose the action which maximizes the expected utility given the evidence
- Can directly operationalize this with decision networks
 - Bayes nets with nodes for utility and actions
 - Lets us calculate the expected utility for each action
- New node types:
 - Chance nodes (just like BNs)
 - Actions (rectangles, cannot have parents, act as observed evidence)
 - Utility node (diamond, depends on action and chance nodes)

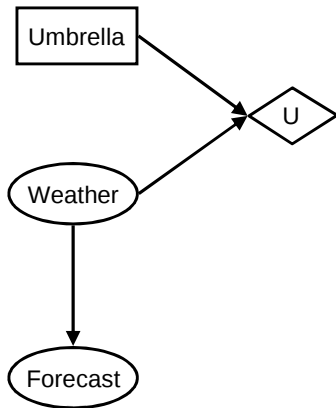


[DEMO: Ghostbusters]

Decision Networks

■ Action selection:

- Instantiate all evidence
- Set action node(s) each possible way
- Calculate posterior for all parents of utility node, given the evidence
- Calculate expected utility for each action
- Choose maximizing action



Example: Decision Networks

Umbrella = leave

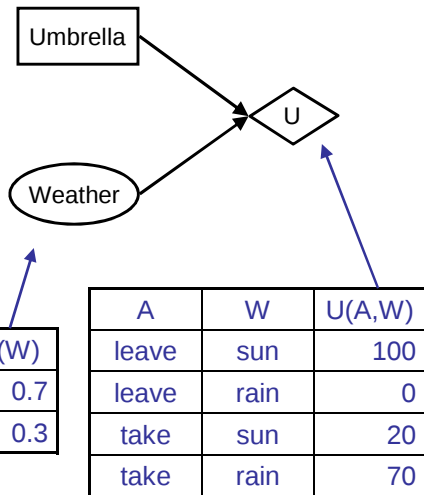
$$\begin{aligned} EU(\text{leave}) &= \sum_w P(w)U(\text{leave}, w) \\ &= 0.7 \cdot 100 + 0.3 \cdot 0 = 70 \end{aligned}$$

Umbrella = take

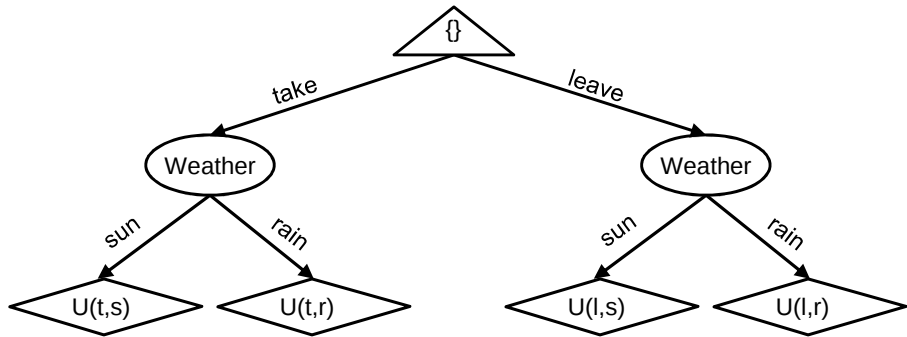
$$\begin{aligned} EU(\text{take}) &= \sum_w P(w)U(\text{take}, w) \\ &= 0.7 \cdot 20 + 0.3 \cdot 70 = 35 \end{aligned}$$

Optimal decision = leave

$$MEU(\emptyset) = \max_a EU(a) = 70$$

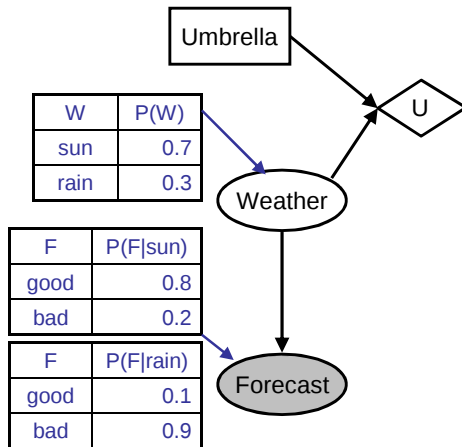


Decisions as Outcome Trees



- Almost exactly like expectimax / MDPs
- What's changed?

Evidence in Decision Networks



Find $P(W|F=\text{bad})$

Select for evidence

W	P(W)
sun	0.7
rain	0.3

$P(W)$

W	P(F=bad W)
sun	0.2
rain	0.9

$P(\text{bad}|W)$

First we join $P(W)$ and $P(\text{bad}|W)$

Then we normalize

W	P(W, F=bad)
sun	0.14
rain	0.27

$P(W, \text{bad})$



W	P(W F=bad)
sun	0.34
rain	0.66

$P(W|F = \text{bad})$

Example: Decision Networks

Umbrella = leave

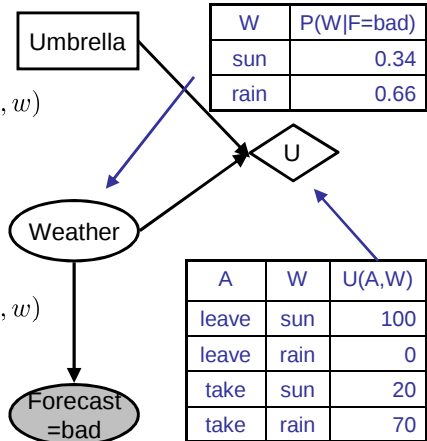
$$\begin{aligned} EU(\text{leave}|\text{bad}) &= \sum_w P(w|\text{bad})U(\text{leave}, w) \\ &= 0.34 \cdot 100 + 0.66 \cdot 0 = 34 \end{aligned}$$

Umbrella = take

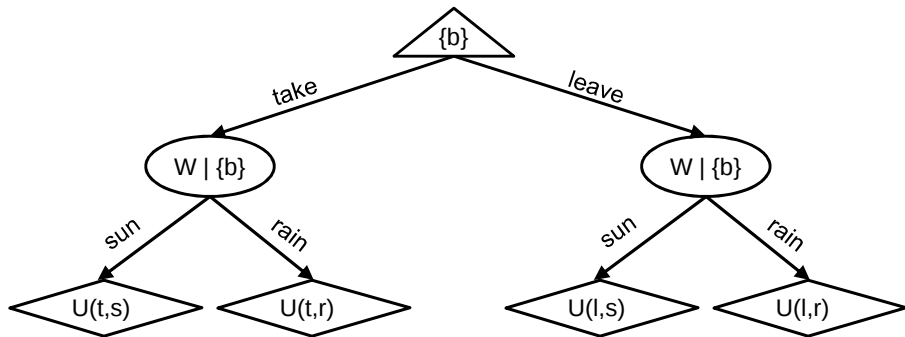
$$\begin{aligned} EU(\text{take}|\text{bad}) &= \sum_w P(w|\text{bad})U(\text{take}, w) \\ &= 0.34 \cdot 20 + 0.66 \cdot 70 = 53 \end{aligned}$$

Optimal decision = take

$$MEU(F = \text{bad}) = \max_a EU(a|\text{bad}) = 53$$

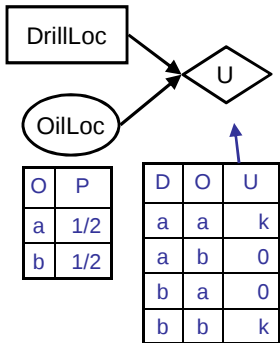


Decisions as Outcome Trees



Value of Information

- Idea: compute value of acquiring evidence
 - Can be done directly from decision network
- Example: buying oil drilling rights
 - Two blocks A and B, exactly one has oil, worth k
 - You can drill in one location
 - Prior probabilities 0.5 each, & mutually exclusive
 - Drilling in either A or B has $EU = k/2$, $MEU = k/2$
- Question: what's the value of information of O ?
 - Value of knowing which of A or B has oil
 - Value is expected gain in MEU from new info
 - Survey may say "oil in a" or "oil in b," prob 0.5 each
 - If we know OilLoc, MEU is k (either way)
 - Gain in MEU from knowing OilLoc?
 - $VPI(OilLoc) = k/2$
 - Fair price of information: $k/2$



Value of Information

- Assume we have evidence $E=e$. Value if we act now:

$$MEU(e) = \max_a \sum_s P(s|e) U(s, a)$$

- Assume we see that $E' = e'$. Value if we act then:

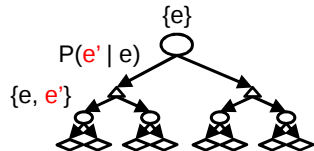
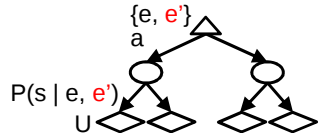
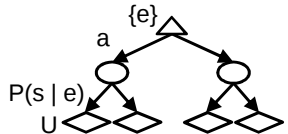
$$MEU(e, e') = \max_a \sum_s P(s|e, e') U(s, a)$$

- BUT E' is a random variable whose value is unknown, so we don't know what e' will be
- Expected value if E' is revealed and then we act:

$$MEU(e, E') = \sum_{e'} P(e'|e) MEU(e, e')$$

- Value of information: how much MEU goes up by revealing E' first then acting, over acting now:

$$VPI(E'|e) = MEU(e, E') - MEU(e)$$



VPI Example: Weather

MEU with no evidence

$$\text{MEU}(\emptyset) = \max_a \text{EU}(a) = 70$$

MEU if forecast is bad

$$\text{MEU}(F = \text{bad}) = \max_a \text{EU}(a|\text{bad}) = 53$$

MEU if forecast is good

$$\text{MEU}(F = \text{good}) = \max_a \text{EU}(a|\text{good}) = 95$$

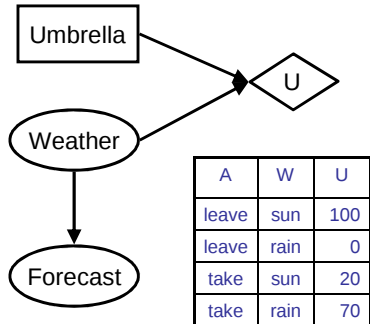
Forecast distribution

F	P(F)
good	0.59
bad	0.41



$$0.59 \cdot (95) + 0.41 \cdot (53) - 70$$

$$77.8 - 70 = 7.8$$



$$\text{VPI}(E|e') = \left(\sum_{e'} P(e'|e) \text{MEU}(e, e') \right) - \text{MEU}(e)$$

VPI Properties

- Nonnegative

$$\forall E', e : \text{VPI}(E'|e) \geq 0$$

- Nonadditive – consider, e.g., obtaining E_j twice

$$\text{VPI}(E_j, E_k|e) \neq \text{VPI}(E_j|e) + \text{VPI}(E_k|e)$$

- Order-independent

$$\begin{aligned}\text{VPI}(E_j, E_k|e) &= \text{VPI}(E_j|e) + \text{VPI}(E_k|e, E_j) \\ &= \text{VPI}(E_k|e) + \text{VPI}(E_j|e, E_k)\end{aligned}$$

Quick VPI Questions

- The soup of the day is either clam chowder or split pea, but you wouldn't order either one. What's the value of knowing which it is?
- There are two kinds of plastic forks at a picnic. One kind is slightly sturdier. What's the value of knowing which?
- You're playing the lottery. The prize will be \$0 or \$100. You can play any number between 1 and 100 (chance of winning is 1%). What is the value of knowing the winning number?