

INFERÊNCIA EM REDES BAYESIANAS

ANO LECTIVO 2013/2014 - 1^o SEMESTRE (14.4 - 14.5) –
ADAPTADO DE
[HTTP://AIMA.EECS.BERKELEY.EDU/SLIDES-TEX/](http://aima.eecs.berkeley.edu/slides-tex/)

Resumo

- ◇ Inferência exacta por enumeração
- ◇ Inferência exacta por eliminação de variáveis
- ◇ Inferência aproximada por simulação estocástica

Tarefas de inferência

Interrogações simples: calcular distribuição marginal a posteriori $\mathbf{P}(X_i|\mathbf{E} = \mathbf{e})$
e.g., $P(\text{NoGas}|\text{Gauge} = \text{empty}, \text{Lights} = \text{on}, \text{Starts} = \text{false})$

Interrogações conjuntivas: $\mathbf{P}(X_i, X_j|\mathbf{E} = \mathbf{e}) = \mathbf{P}(X_i|\mathbf{E} = \mathbf{e})\mathbf{P}(X_j|X_i, \mathbf{E} = \mathbf{e})$

Decisões óptimas: redes de decisão incluem informação de utilidade;
inferência probabilística necessária para $P(\text{outcome}|\text{action}, \text{evidence})$

Valor da informação: que evidência procurar a seguir?

Análise de sensibilidade: quais os valores de probabilidade mais críticos?

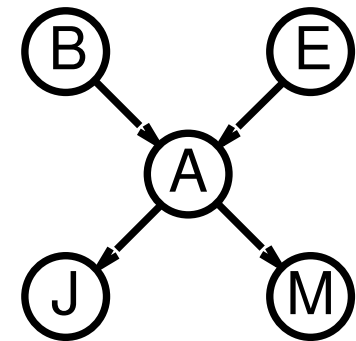
Explicação: por que é que preciso de um novo motor de arranque?

Inferência por enumeração

Forma relativamente inteligente de somar variáveis da distribuição conjunta sem necessitar de construir a sua representação explícita

Interrogação simples na rede do assaltante:

$$\begin{aligned} & \mathbf{P}(B|j, m) \\ &= \mathbf{P}(B, j, m) / P(j, m) \\ &= \alpha \mathbf{P}(B, j, m) \\ &= \alpha \sum_e \sum_a \mathbf{P}(B, e, a, j, m) \end{aligned}$$



Rescrever entradas da distribuição conjunta recorrendo ao produto de entradas da CPT:

$$\begin{aligned} & \mathbf{P}(B|j, m) \\ &= \alpha \sum_e \sum_a \mathbf{P}(B) P(e) \mathbf{P}(a|B, e) P(j|a) P(m|a) \\ &= \alpha \mathbf{P}(B) \sum_e P(e) \sum_a \mathbf{P}(a|B, e) P(j|a) P(m|a) \quad (expr1) \end{aligned}$$

Enumeração recursiva em profundidade primeiro: espaço $O(n)$, tempo $O(d^n)$

Inferência por enumeração

Vamos fazer $B = \text{true}$ em expr1 (ignorando α):

$$\begin{aligned} & \mathbf{P}(B = \text{true}) \sum_{e \in \{\text{true}, \text{false}\}} P(E = e) \sum_{a \in \{\text{true}, \text{false}\}} \mathbf{P}(A = a | B = \text{true}, E = e) P(J = \text{true} | A = a) P(M = \text{true} | A = a) \\ &= 0.001 \times [0.002 \times (0.95 \times 0.9 \times 0.7 + 0.05 \times 0.05 \times 0.01) \\ &\quad + 0.998 \times (0.94 \times 0.9 \times 0.7 + 0.06 \times 0.05 \times 0.01)] = 0,000592243 \end{aligned}$$

Vamos fazer $B = \text{false}$ em expr1 (ignorando α):

$$\begin{aligned} & \mathbf{P}(B = \text{false}) \sum_{e \in \{\text{true}, \text{false}\}} P(E = e) \sum_{a \in \{\text{true}, \text{false}\}} \mathbf{P}(A = a | B = \text{false}, E = e) P(J = \text{true} | A = a) P(M = \text{true} | A = a) \\ &= 0.999 \times [0.002 \times (0.29 \times 0.9 \times 0.7 + 0.71 \times 0.05 \times 0.01) \\ &\quad + 0.998 \times (0.001 \times 0.9 \times 0.7 + 0.999 \times 0.05 \times 0.01)] = 0,001491858 \end{aligned}$$

Portanto normalizando, obtemos:

$$\begin{aligned} \mathbf{P}(B = \text{true} | J = \text{true}, M = \text{true}) &= \frac{0,000592243}{0,000592243 + 0,001491858} = 0,284171835 \\ \mathbf{P}(B = \text{false} | J = \text{true}, M = \text{true}) &= \frac{0,001491858}{0,000592243 + 0,001491858} = 0,715828165 \end{aligned}$$

Algoritmo de enumeração

function ENUMERATION-ASK(X, e, bn) **returns** a distribution over X

inputs: X , the query variable

e , observed values for variables $\mathbf{E}$

bn , a Bayesian network with variables $\{X\} \cup \mathbf{E} \cup \mathbf{Y}$

$Q(X) \leftarrow$ a distribution over X , initially empty

for each value x_i of X **do**

 extend e with value x_i for X

$Q(x_i) \leftarrow$ ENUMERATE-ALL(VARS[bn], e)

return NORMALIZE($Q(X)$)

function ENUMERATE-ALL($vars, e$) **returns** a real number

if EMPTY?($vars$) **then return** 1.0

$Y \leftarrow$ FIRST($vars$)

if Y has value y in e

then return $P(y \mid Pa(Y)) \times$ ENUMERATE-ALL(REST($vars$), e)

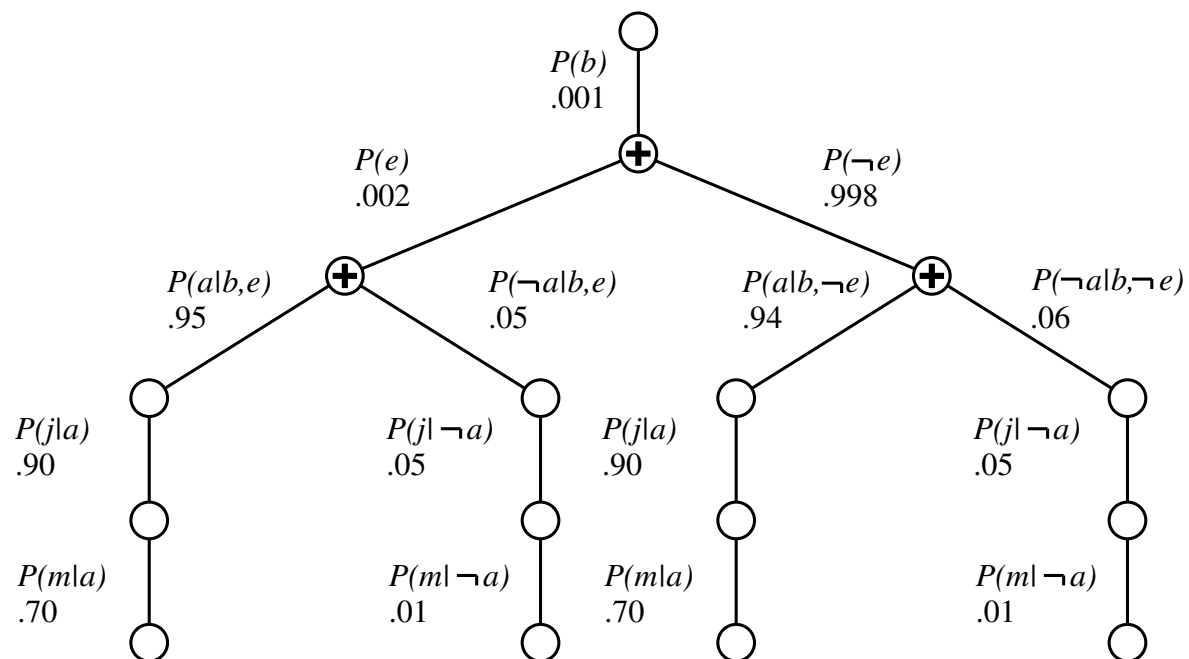
else return $\sum_y P(y \mid Pa(Y)) \times$ ENUMERATE-ALL(REST($vars$), e_y)

 where e_y is e extended with $Y = y$

Árvore de avaliação

Enumeração é ineficiente: cálculos repetidos

e.g., calcula $P(j|a)P(m|a)$ para cada valor de e



Inferência por eliminação de variáveis

Eliminação de variáveis: efectuar somatórios da direita para a esquerda, armazenando resultados intermédios (**factores**) para evitar recomputação

$$\begin{aligned} \mathbf{P}(B|j, m) &= \alpha \underbrace{\mathbf{P}(B)}_B \sum_e \underbrace{P(e)}_E \sum_a \underbrace{\mathbf{P}(a|B, e)}_A \underbrace{P(j|a)}_J \underbrace{P(m|a)}_M \\ &= \alpha \mathbf{P}(B) \sum_e P(e) \sum_a \mathbf{P}(a|B, e) P(j|a) f_M(a) \\ &= \alpha \mathbf{P}(B) \sum_e P(e) \sum_a \mathbf{P}(a|B, e) f_J(a) f_M(a) \\ &= \alpha \mathbf{P}(B) \sum_e P(e) \sum_a f_A(a, B, e) f_J(a) f_M(a) \\ &= \alpha \mathbf{P}(B) \sum_e P(e) f_{\bar{A}JM}(B, e) \text{ (soma-se } A) \\ &= \alpha \mathbf{P}(B) f_{\bar{E}\bar{A}JM}(B) \text{ (soma-se } E) \\ &= \alpha f_B(B) \times f_{\bar{E}\bar{A}JM}(B) \end{aligned}$$

Eliminação de variáveis: operações básicas

Somar para eliminar a variável de um produto de factores:

deslocar todos os factores contantes para fora do somatório

adicionar submatrizes no produto pontual dos factores restantes

$$\sum_x f_1 \times \cdots \times f_k = f_1 \times \cdots \times f_i \sum_x f_{i+1} \times \cdots \times f_k = f_1 \times \cdots \times f_i \times f_{\bar{X}}$$

assumindo que $f_1, \dots, f_i$ não depende de X

Produto pontual dos factores f_1 e f_2 :

$$\begin{aligned} f_1(x_1, \dots, x_j, y_1, \dots, y_k) \times f_2(y_1, \dots, y_k, z_1, \dots, z_l) \\ = f(x_1, \dots, x_j, y_1, \dots, y_k, z_1, \dots, z_l) \end{aligned}$$

E.g., $f_1(a, b) \times f_2(b, c) = f(a, b, c)$

Algoritmo de eliminação de variáveis

```
function ELIMINATION-ASK( $X, e, bn$ ) returns a distribution over  $X$   
  inputs:  $X$ , the query variable  
            $e$ , evidence specified as an event  
            $bn$ , a belief network specifying joint distribution  $\mathbf{P}(X_1, \dots, X_n)$   
   $factors \leftarrow []$ ;  $vars \leftarrow \text{REVERSE}(\text{VARS}[bn])$   
  for each  $var$  in  $vars$  do  
     $factors \leftarrow [\text{MAKE-FACTOR}(var, e) | factors]$   
    if  $var$  is a hidden variable then  $factors \leftarrow \text{SUM-OUT}(var, factors)$   
  return  $\text{NORMALIZE}(\text{POINTWISE-PRODUCT}(factors))$ 
```

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls = true, MaryCalls = true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: []

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): **M**, J, A, B, E

Factores: $[f_M(A)]$

Factor da variável *M*: $f_M(A) = P(M = \text{true} | A)$

<i>A</i>	$f_M(A)$
<i>true</i>	0.70
<i>false</i>	0.01

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: $[f_J(A), f_M(A)]$

Factor da variável *J*: $f_J(A) = P(J = \text{true} | A)$

<i>A</i>	$f_J(A)$	<i>A</i>	$f_M(A)$
<i>true</i>	0.90	<i>true</i>	0.70
<i>false</i>	0.05	<i>false</i>	0.01

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: $[f_A(A, B, E), f_J(A), f_M(A)]$

Factor da variável *A*: $f_A(A, B, E) = P(A|B, E)$

<i>A</i>	<i>B</i>	<i>E</i>	$f_A(A, B, E)$
<i>true</i>	<i>true</i>	<i>true</i>	0.95
<i>true</i>	<i>true</i>	<i>false</i>	0.94
<i>true</i>	<i>false</i>	<i>true</i>	0.29
<i>true</i>	<i>false</i>	<i>false</i>	0.001
<i>false</i>	<i>true</i>	<i>true</i>	0.05
<i>false</i>	<i>true</i>	<i>false</i>	0.06
<i>false</i>	<i>false</i>	<i>true</i>	0.71
<i>false</i>	<i>false</i>	<i>false</i>	0.999

<i>A</i>	$f_J(A)$
<i>true</i>	0.90
<i>false</i>	0.05

<i>A</i>	$f_M(A)$
<i>true</i>	0.70
<i>false</i>	0.01

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: $[f_A(A, B, E), f_J(A), f_M(A)]$

Como a variável A é oculta vai-se eliminá-la (através da soma). Para isso é necessário:

- Efectuar o produto pontual dos factores que têm o parâmetro A

$$f_{AJM}(A, B, E) = f_A(A, B, E) \times f_J(A) \times f_M(A)$$

- Eliminar a variável A obtendo o factor $f_{\bar{A}JM}(B, E)$

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: $[f_A(A, B, E), f_J(A), f_M(A)]$

Produto pontual: $f_{AJM}(A, B, E) = f_A(A, B, E) \times f_{JM}(A)$

$f_{JM}(A) = f_J(A) \times f_M(A)$

<i>A</i>	$f_{JM}(A)$	=	<i>A</i>	$f_J(A)$	×	<i>A</i>	$f_M(A)$
<i>true</i>	$0.90 \times 0.70 = 0.63$		<i>true</i>	0.90		<i>true</i>	0.70
<i>false</i>	$0.05 \times 0.01 = 0.0005$		<i>false</i>	0.05		<i>false</i>	0.01

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: $[f_A(A, B, E), f_J(A), f_M(A)]$

Produto pontual: $f_{AJM}(A, B, E) = f_A(A, B, E) \times f_{JM}(A)$

<i>A</i>	<i>B</i>	<i>E</i>	$f_{AJM}(A, B, E)$
<i>true</i>	<i>true</i>	<i>true</i>	$0.95 \times 0.63 = 0.5985$
<i>true</i>	<i>true</i>	<i>false</i>	0.5922
<i>true</i>	<i>false</i>	<i>true</i>	0.1827
<i>true</i>	<i>false</i>	<i>false</i>	0.00063
<i>false</i>	<i>true</i>	<i>true</i>	0.000025
<i>false</i>	<i>true</i>	<i>false</i>	0.00003
<i>false</i>	<i>false</i>	<i>true</i>	0.000355
<i>false</i>	<i>false</i>	<i>false</i>	0.0004995

=

<i>A</i>	<i>B</i>	<i>E</i>	$f_A(A, B, E)$
<i>true</i>	<i>true</i>	<i>true</i>	0.95
<i>true</i>	<i>true</i>	<i>false</i>	0.94
<i>true</i>	<i>false</i>	<i>true</i>	0.29
<i>true</i>	<i>false</i>	<i>false</i>	0.001
<i>false</i>	<i>true</i>	<i>true</i>	0.05
<i>false</i>	<i>true</i>	<i>false</i>	0.06
<i>false</i>	<i>false</i>	<i>true</i>	0.71
<i>false</i>	<i>false</i>	<i>false</i>	0.999

×

<i>A</i>	$f_{JM}(A)$
<i>true</i>	$0.90 \times 0.70 = 0.63$
<i>false</i>	$0.05 \times 0.01 = 0.0005$

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: $[f_{\bar{A}JM}(B, E)]$

Eliminação por soma da variável A: $f_{\bar{A}JM}(B, E)$

<i>B</i>	<i>E</i>	$f_{\bar{A}JM}(B, E)$
<i>true</i>	<i>true</i>	$0.5985 + 0.000025 = 0.598525$
<i>true</i>	<i>false</i>	$0.5922 + 0.000003 = 0.59223$
<i>false</i>	<i>true</i>	$0.1827 + 0.000355 = 0.183055$
<i>false</i>	<i>false</i>	$0.00063 + 0.0004995 = 0.00113$

$= SUM_A$

<i>A</i>	<i>B</i>	<i>E</i>	$f_{AJM}(A, B, E)$
<i>true</i>	<i>true</i>	<i>true</i>	0.5985
<i>true</i>	<i>true</i>	<i>false</i>	0.5922
<i>true</i>	<i>false</i>	<i>true</i>	0.1827
<i>true</i>	<i>false</i>	<i>false</i>	0.00063
<i>false</i>	<i>true</i>	<i>true</i>	0.000025
<i>false</i>	<i>true</i>	<i>false</i>	0.000003
<i>false</i>	<i>false</i>	<i>true</i>	0.000355
<i>false</i>	<i>false</i>	<i>false</i>	0.0004995

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: $[f_B(B), f_{\overline{AJM}}(B, E)]$

Factor da variável *B*: $f_B(B) = P(B)$

<i>B</i>	$f_B(B)$	<i>B</i>	<i>E</i>	$f_{\overline{AJM}}(B, E)$
		<i>true</i>	<i>true</i>	0.598525
<i>true</i>	0.001	<i>true</i>	<i>false</i>	0.59223
<i>false</i>	0.999	<i>false</i>	<i>true</i>	0.183055
		<i>false</i>	<i>false</i>	0.00113

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores: $[f_E(E), f_B(B), f_{AJM}(B, E)]$

Factor da variável *E*: $f_E(E) = P(E)$

<i>E</i>	$f_E(E)$
<i>true</i>	0.002
<i>false</i>	0.998

<i>B</i>	$f_B(B)$
<i>true</i>	0.001
<i>false</i>	0.999

<i>B</i>	<i>E</i>	$f_{AJM}(B, E)$
<i>true</i>	<i>true</i>	0.598525
<i>true</i>	<i>false</i>	0.59223
<i>false</i>	<i>true</i>	0.183055
<i>false</i>	<i>false</i>	0.00113

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): **M**, **J**, **A**, **B**, **E**

Factores: $[f_B(B), f_{\overline{E}AJM}(B)]$

Eliminação da variável *E*: $f_{\overline{E}AJM}(B) = SUM_E (f_E(E) \times f_{AJM}(B, E))$

<i>B</i>	<i>E</i>	$f_{\overline{E}AJM}(B, E)$		<i>E</i>	$f_E(E)$		<i>B</i>	<i>E</i>	$f_{AJM}(B, E)$
<i>true</i>	<i>true</i>	$0.002 \times 0.598525 = 0.001197$	=	<i>true</i>	0.002	$\times$	<i>true</i>	<i>true</i>	0.598525
<i>true</i>	<i>false</i>	$0.998 \times 0.59223 = 0.591046$		<i>false</i>	0.998		<i>true</i>	<i>false</i>	0.59223
<i>false</i>	<i>true</i>	$0.002 \times 0.183055 = 0.000366$					<i>false</i>	<i>true</i>	0.183055
<i>false</i>	<i>false</i>	$0.998 \times 0.00113 = 0.001127$					<i>false</i>	<i>false</i>	0.00113

Factores finais:

<i>B</i>	$f_B(B)$	<i>B</i>	$f_{\overline{E}AJM}(B)$
<i>true</i>	0.001	<i>true</i>	$0.001197 + 0.591046 = 0.592243$
<i>false</i>	0.999	<i>false</i>	$0.000366 + 0.001127 = 0.001493$

Algoritmo de elim. de variáveis exemplificado

Variável de Interrogação: *Burglary*

Evidência: *JohnCalls* = *true*, *MaryCalls* = *true*

Ordenação das variáveis (das folhas para a raiz): M, J, A, B, E

Factores finais: $[f_B(B), f_{\overline{E}AJM}(B)]$

Cálculos finais: $f_{B\overline{E}AJM}(B) = f_B(B) \times f_{\overline{E}AJM}(B)$

B	$f_{B\overline{E}AJM}(B)$	=	B	$f_B(B)$	×	B	$f_{\overline{E}AJM}(B)$
<i>true</i>	$0.001 \times 0.592243 = 0.000592243$		<i>true</i>	0.001		<i>true</i>	0.592243
<i>false</i>	$0.999 \times 0.001493 = 0.001491858$		<i>false</i>	0.999		<i>false</i>	0.001493

Resultado $P(B|JohnCalls = true, MaryCalls = true) = Normalize(f_{B\overline{E}AJM}(B))$

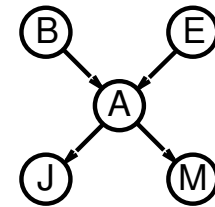
B	$P(B J = true, M = true)$
<i>true</i>	$\frac{0.000592243}{0.000592243 + 0.001491858} = 0.284172$
<i>false</i>	$\frac{0.001491858}{0.000592243 + 0.001491858} = 0.715828$

Variáveis irrelevantes

Considere-se a interrogação $P(JohnCalls|Burglary = true)$

$$P(J|b) = \alpha P(b) \sum_e P(e) \sum_a P(a|b, e) P(J|a) \sum_m P(m|a)$$

A soma de m é igual 1; M é **irrelevante** para a interrogação



Teorema 1: Y é irrelevante a não ser que $Y \in Ancestors(\{X\} \cup \mathbf{E})$

Fazendo $X = JohnCalls$, $\mathbf{E} = \{Burglary\}$, e
 $Ancestors(\{X\} \cup \mathbf{E}) = \{Alarm, Earthquake\}$
conclui-se que M é irrelevante

(Compare-se com o método de encadeamento para trás a partir da interrogação em KBs de cláusulas de Horn)

Complexidade da inferência exacta

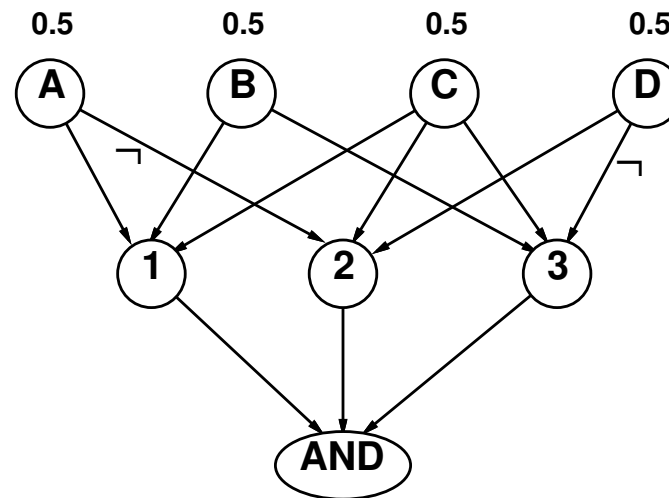
Redes **singularmente conexas** (ou **polytrees**):

- quaisquer dois nós estão ligados no máximo por um caminho (não dirigido)
- complexidade temporal e espacial da eliminação de variáveis é $O(d^k n)$

Redes **multiplamente conexas**:

- pode-se reduzir 3SAT a inferência exacta $\Rightarrow$ NP-difícil
- equivalente **contar** modelos 3SAT $\Rightarrow$ #P-completo

1. $A \vee B \vee C$
2. $C \vee D \vee \neg A$
3. $B \vee C \vee \neg D$



Inferência por simulação estocástica

Ideia básica:

- 1) Efectuar N amostras de uma distribuição de amostragem S
- 2) Calcular um probabilidade a posteriori aproximada $\hat{P}$
- 3) Mostrar que processo converge para a probabilidade real P

0.5

Métodos:

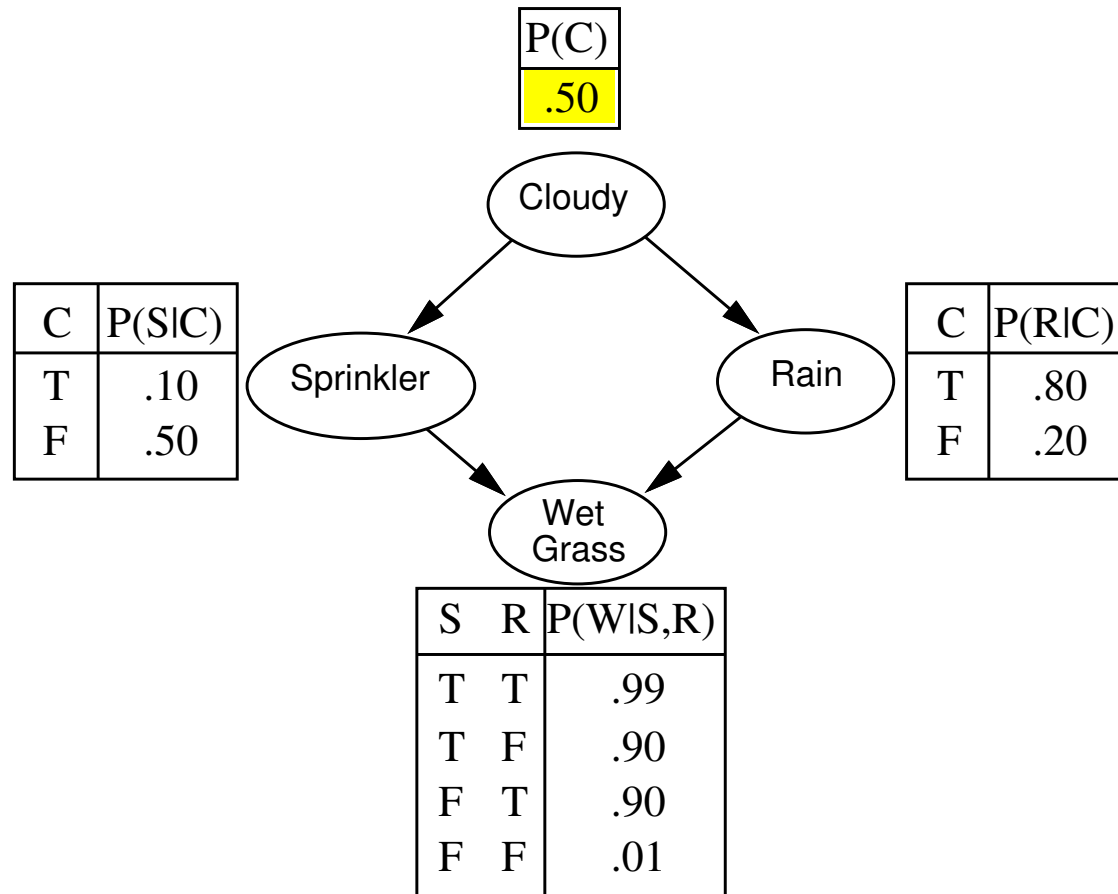
- Amostragem a partir de uma rede vazia
- Amostragem por rejeição: rejeitar amostras que não estão de acordo com evidência
- Pesagem por Verosimilhança: utilizar evidência para pesar amostras
- Markov chain Monte Carlo (MCMC): amostrar a partir de um processo estocástico cuja distribuição estacionária é a probabilidade a posteriori real



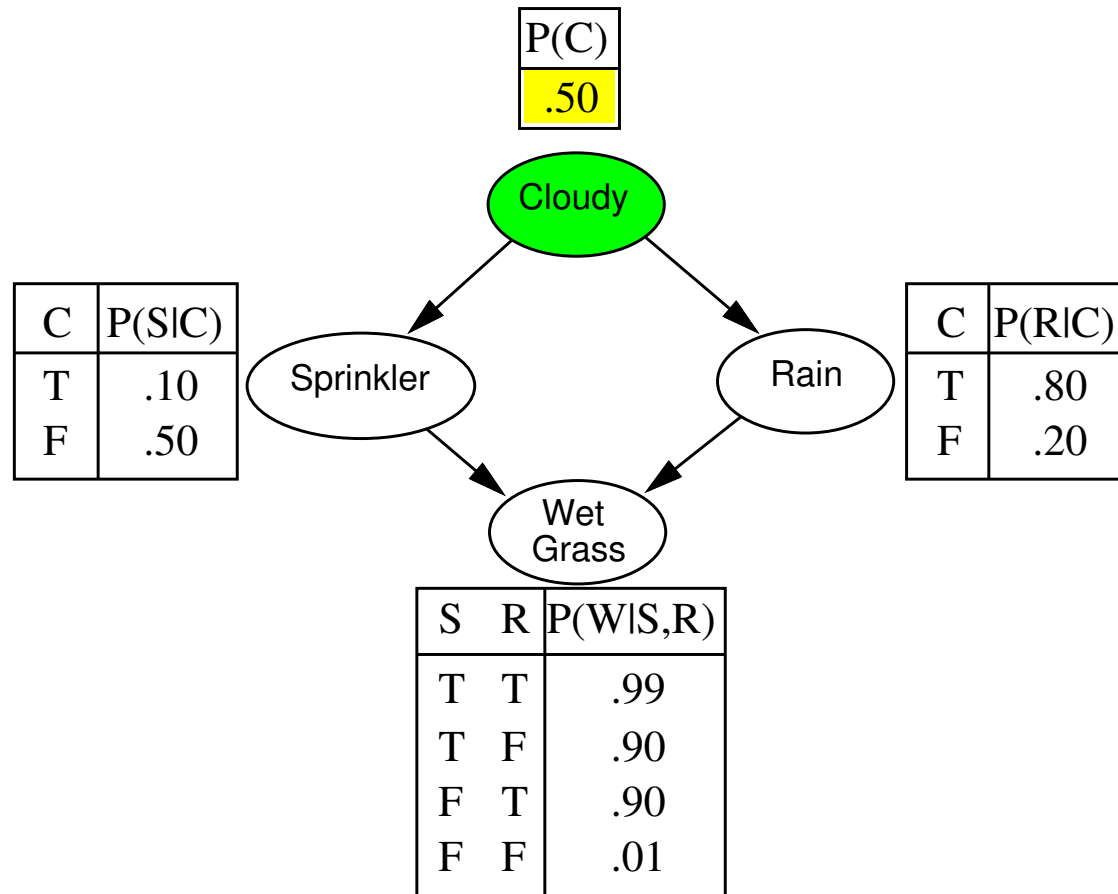
Amostragem a partir de uma rede vazia

```
function PRIOR-SAMPLE( $bn$ ) returns an event sampled from  $bn$   
  inputs:  $bn$ , a belief network specifying joint distribution  $\mathbf{P}(X_1, \dots, X_n)$   
   $\mathbf{x} \leftarrow$  an event with  $n$  elements  
  for  $i = 1$  to  $n$  do  
     $x_i \leftarrow$  a random sample from  $\mathbf{P}(X_i \mid Parents(X_i))$   
  return  $\mathbf{x}$ 
```

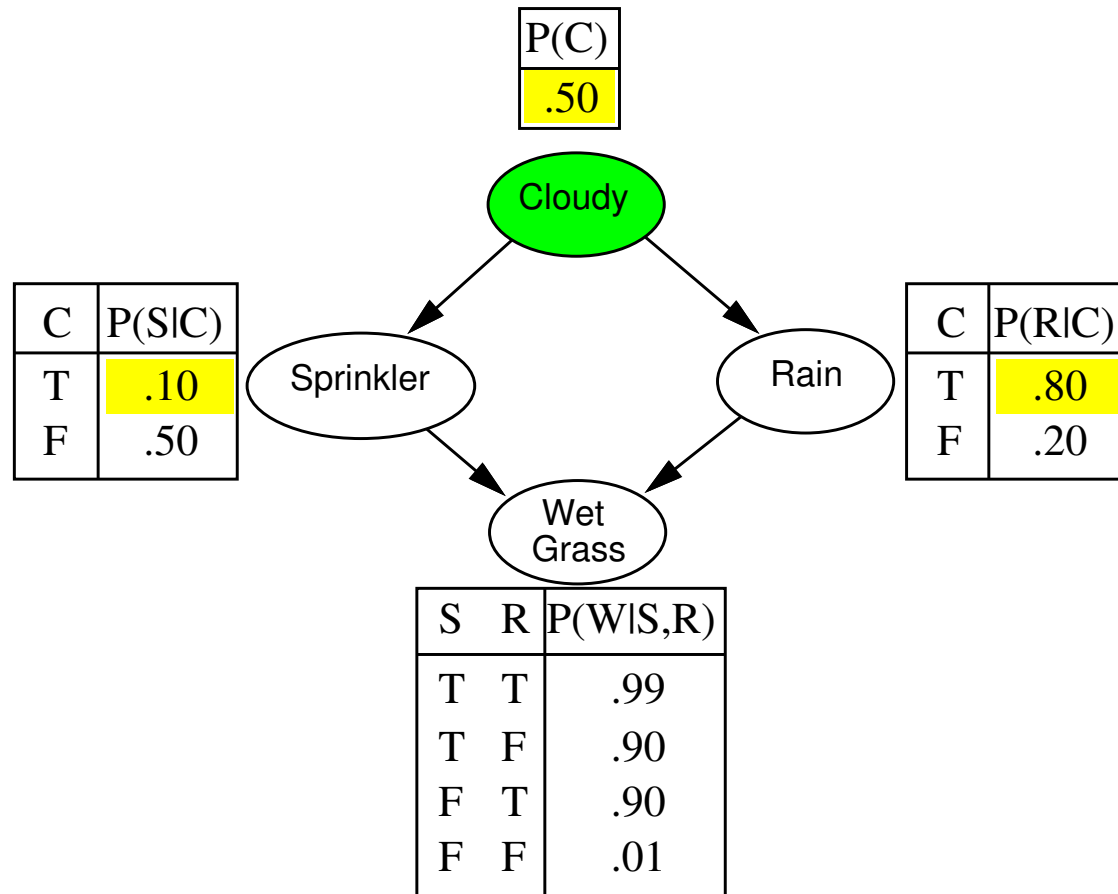
Exemplo



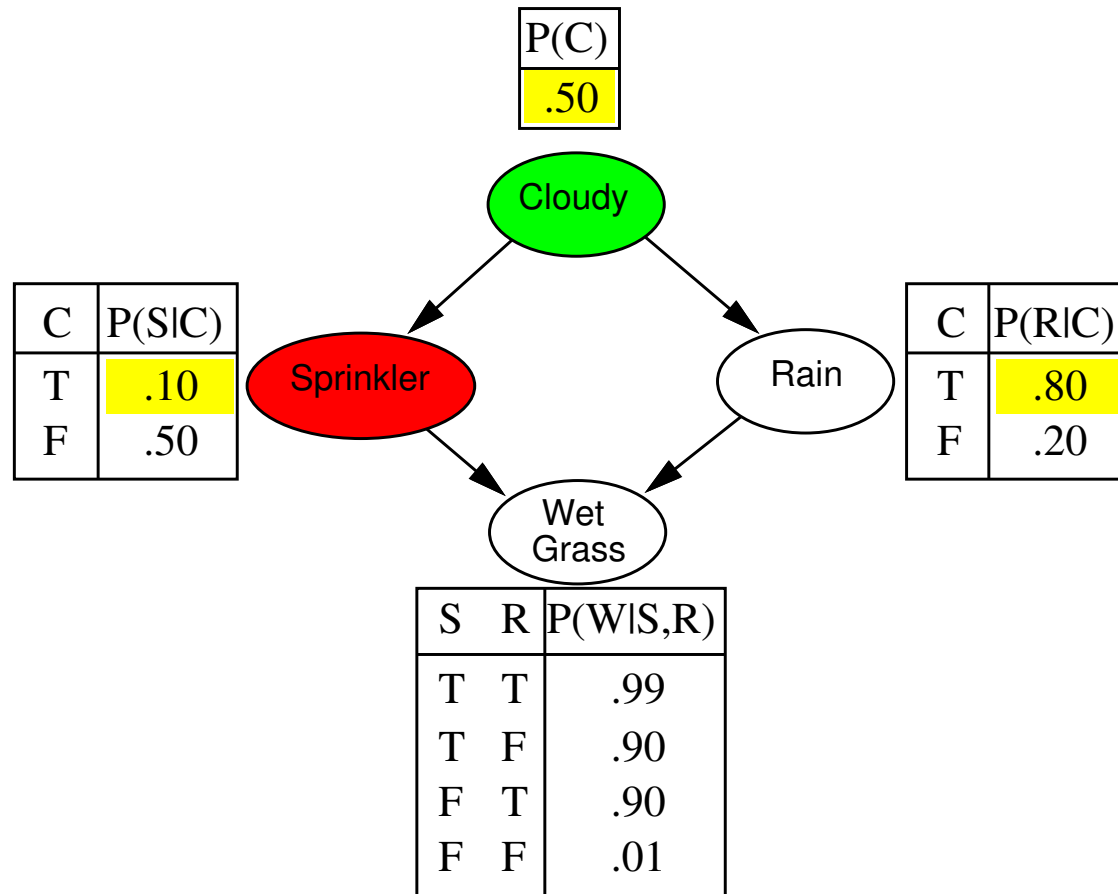
Exemplo



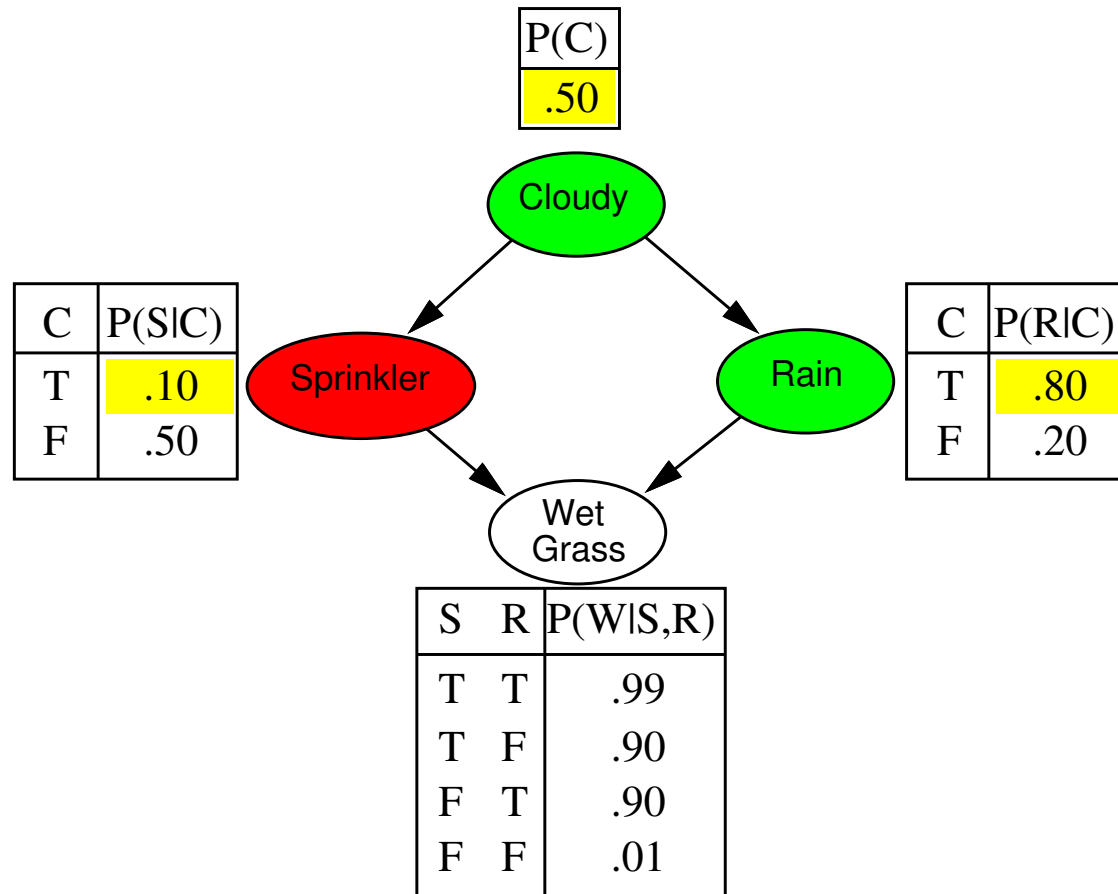
Exemplo



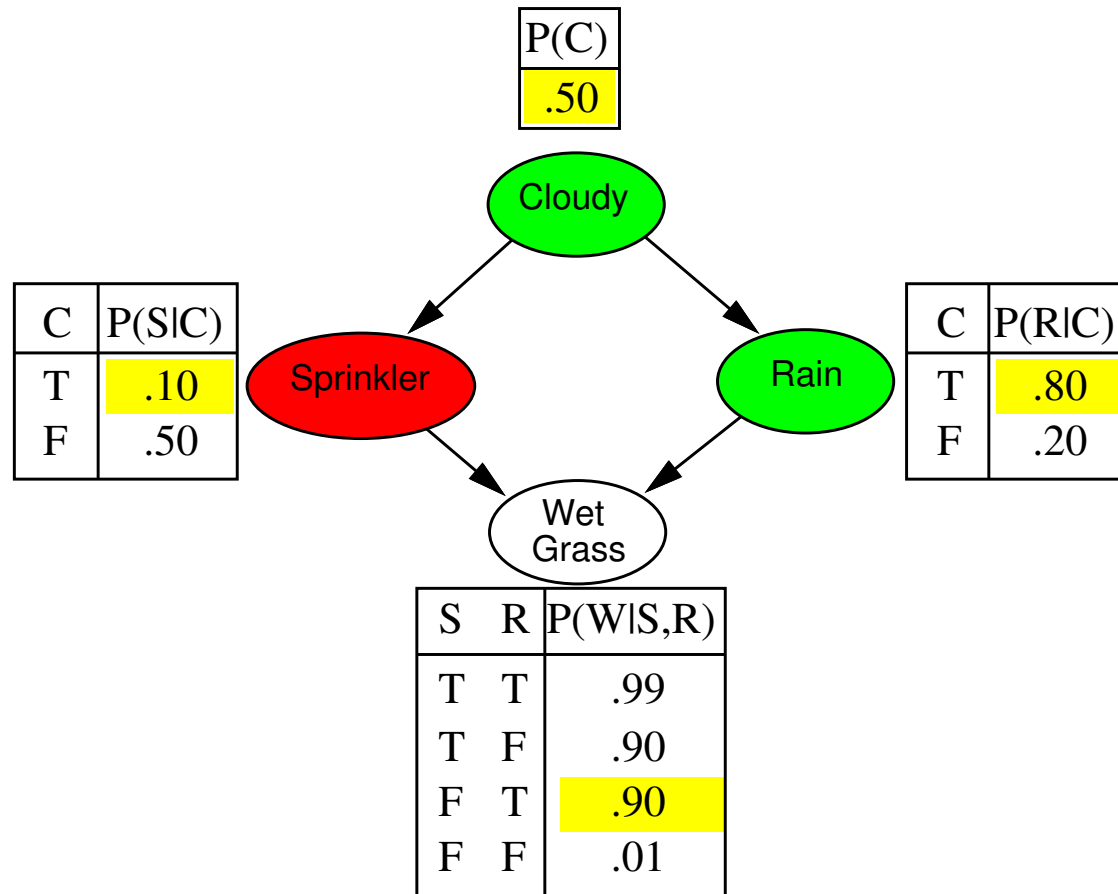
Exemplo



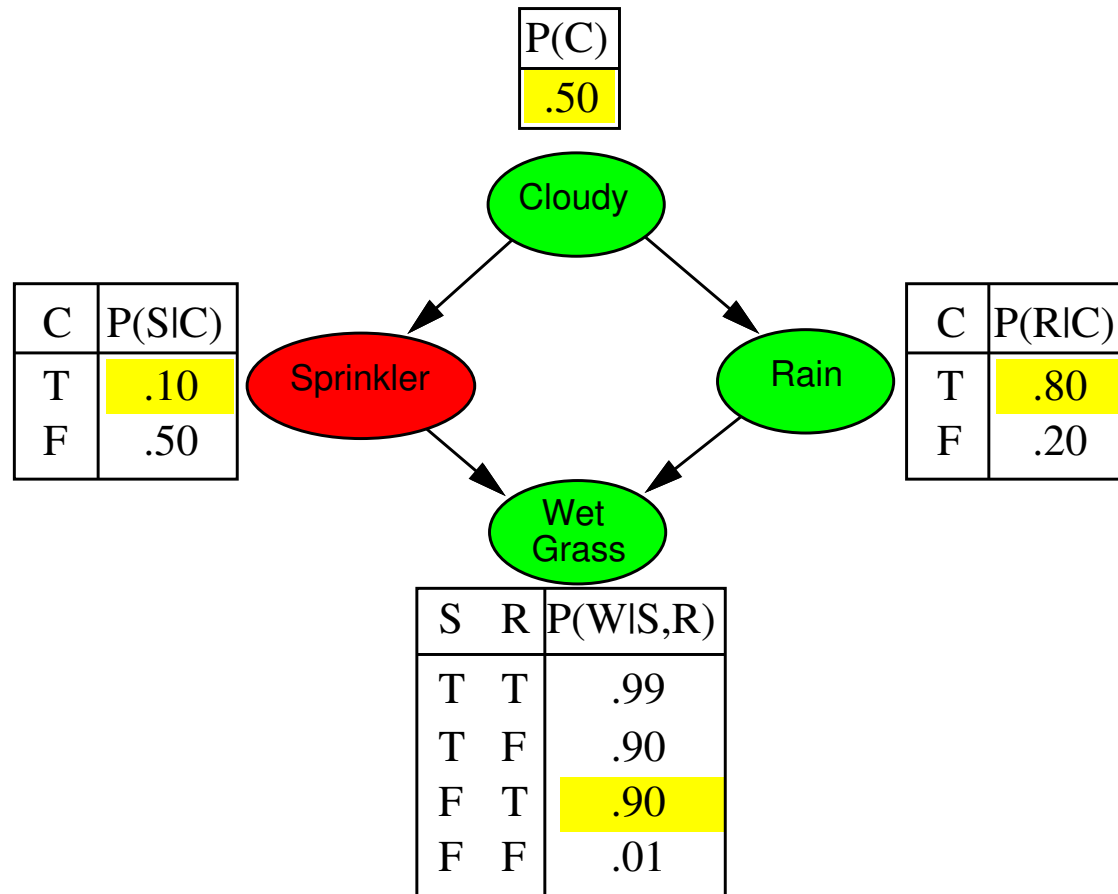
Exemplo



Exemplo



Exemplo



Amostragem a partir de uma rede vazia (cont.)

Probabilidade de PRIORSAMPLE gerar um acontecimento particular

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

i.e., a probabilidade a priori real

$$\text{E.g., } S_{PS}(t, f, t, t) = 0.5 \times 0.9 \times 0.8 \times 0.9 = 0.324 = P(t, f, t, t)$$

Seja $N_{PS}(x_1 \dots x_n)$ o número de amostras geradas para o acontecimento $x_1, \dots, x_n$

Então obtemos

$$\begin{aligned} \lim_{N \rightarrow \infty} \hat{P}(x_1, \dots, x_n) &= \lim_{N \rightarrow \infty} N_{PS}(x_1, \dots, x_n) / N \\ &= S_{PS}(x_1, \dots, x_n) \\ &= P(x_1 \dots x_n) \end{aligned}$$

Ou seja, estimativas derivadas de PRIORSAMPLE são **consistentes**

Abreviadamente: $\hat{P}(x_1, \dots, x_n) \approx P(x_1 \dots x_n)$

Amostragem por rejeição

$\hat{\mathbf{P}}(X|\mathbf{e})$ estimado das amostras que concordam com $\mathbf{e}$

```
function REJECTION-SAMPLING( $X, \mathbf{e}, bn, N$ ) returns an estimate of  $P(X|\mathbf{e})$ 
  local variables:  $\mathbf{N}$ , a vector of counts over  $X$ , initially zero
  for  $j = 1$  to  $N$  do
     $\mathbf{x} \leftarrow \text{PRIOR-SAMPLE}(bn)$ 
    if  $\mathbf{x}$  is consistent with  $\mathbf{e}$  then
       $\mathbf{N}[x] \leftarrow \mathbf{N}[x] + 1$  where  $x$  is the value of  $X$  in  $\mathbf{x}$ 
  return NORMALIZE( $\mathbf{N}[X]$ )
```

E.g., estimar $\mathbf{P}(\text{Rain}|\text{Sprinkler} = \text{true})$ utilizando 100 amostras

27 amostras têm $\text{Sprinkler} = \text{true}$

Destas, 8 têm $\text{Rain} = \text{true}$ e 19 têm $\text{Rain} = \text{false}$.

$$\hat{\mathbf{P}}(\text{Rain}|\text{Sprinkler} = \text{true}) = \text{NORMALIZE}(\langle 8, 19 \rangle) = \langle 0.296, 0.704 \rangle$$

Semelhante aos procedimentos empíricos de estimativa

Análise da amostragem por rejeição

$$\begin{aligned}\hat{\mathbf{P}}(X|\mathbf{e}) &= \alpha \mathbf{N}_{PS}(X, \mathbf{e}) && \text{(algoritmo defn.)} \\ &= \mathbf{N}_{PS}(X, \mathbf{e}) / N_{PS}(\mathbf{e}) && \text{(normalizado por } N_{PS}(\mathbf{e}) \text{)} \\ &\approx \mathbf{P}(X, \mathbf{e}) / P(\mathbf{e}) && \text{(propriedade de PRIOR SAMPLE)} \\ &= \mathbf{P}(X|\mathbf{e}) && \text{(defn. de probabilidade condicional)}\end{aligned}$$

Logo amostragem por rejeição devolve estimativas consistentes da probabilidade a posteriori

Problema: terrivelmente ineficiente se $P(\mathbf{e})$ é pequeno

$P(\mathbf{e})$ diminui exponencialmente com o número de variáveis evidência!

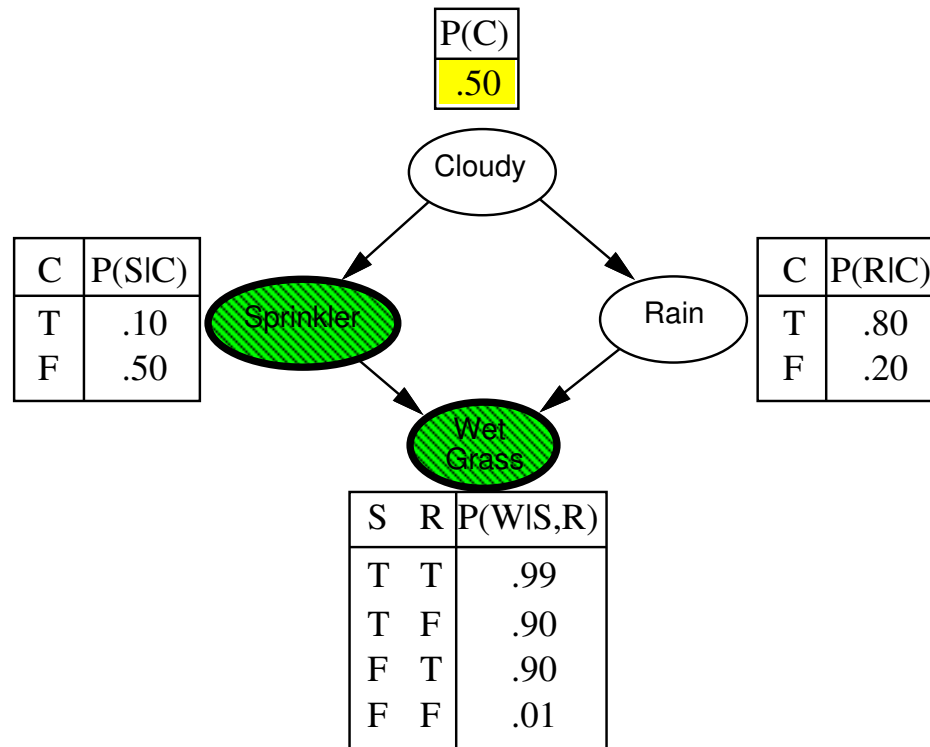
Pesagem por verosimilhança

Ideia: fixar variáveis evidência, amostrar apenas as outras variáveis, e pesar cada amostra pela verosimilhança de acordo com a evidência

```
function LIKELIHOOD-WEIGHTING( $X, e, bn, N$ ) returns an estimate of  $P(X|e)$   
  local variables:  $\mathbf{W}$ , a vector of weighted counts over  $X$ , initially zero  
  
  for  $j = 1$  to  $N$  do  
     $\mathbf{x}, w \leftarrow \text{WEIGHTED-SAMPLE}(bn)$   
     $\mathbf{W}[x] \leftarrow \mathbf{W}[x] + w$  where  $x$  is the value of  $X$  in  $\mathbf{x}$   
  return NORMALIZE( $\mathbf{W}[X]$ )
```

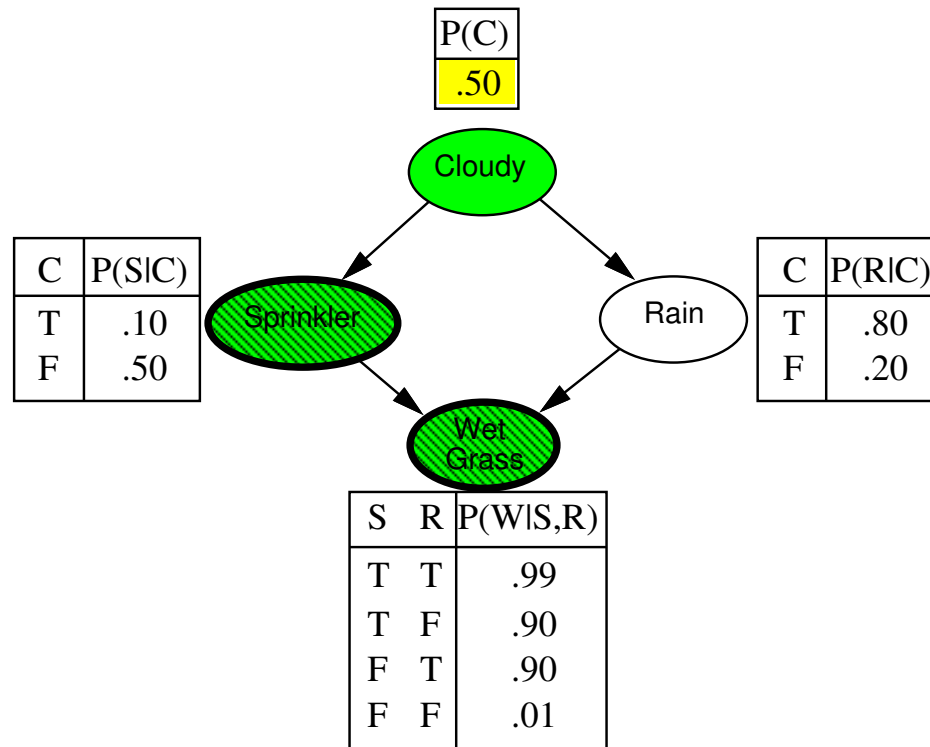
```
function WEIGHTED-SAMPLE( $bn, e$ ) returns an event and a weight  
  
   $\mathbf{x} \leftarrow$  an event with  $n$  elements;  $w \leftarrow 1$   
  for  $i = 1$  to  $n$  do  
    if  $X_i$  has a value  $x_i$  in  $e$   
      then  $w \leftarrow w \times P(X_i = x_i \mid \text{Parents}(X_i))$   
      else  $x_i \leftarrow$  a random sample from  $\mathbf{P}(X_i \mid \text{Parents}(X_i))$   
  return  $\mathbf{x}, w$ 
```

Exemplo de pesagem por verosimilhança



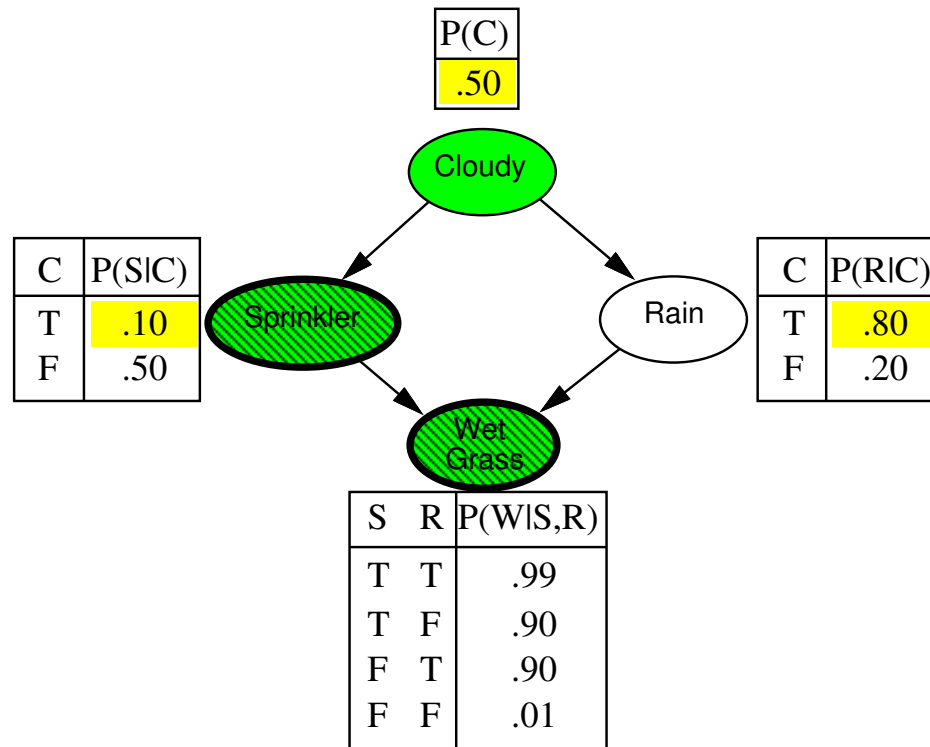
$$w = 1.0$$

Exemplo de pesagem por verosimilhança



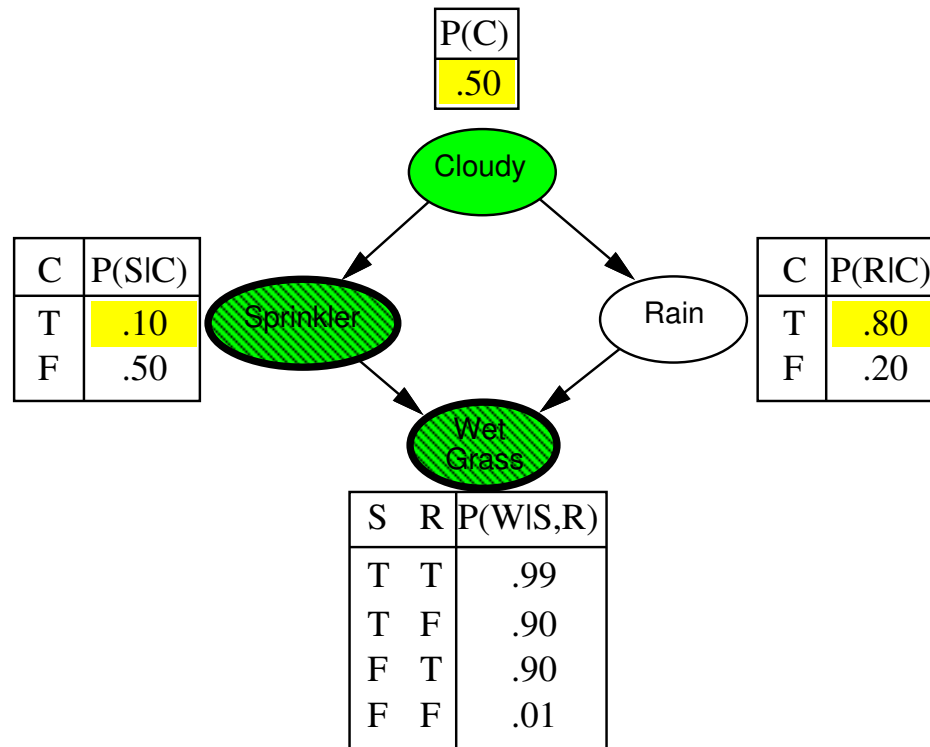
$$w = 1.0$$

Exemplo de pesagem por verosimilhança



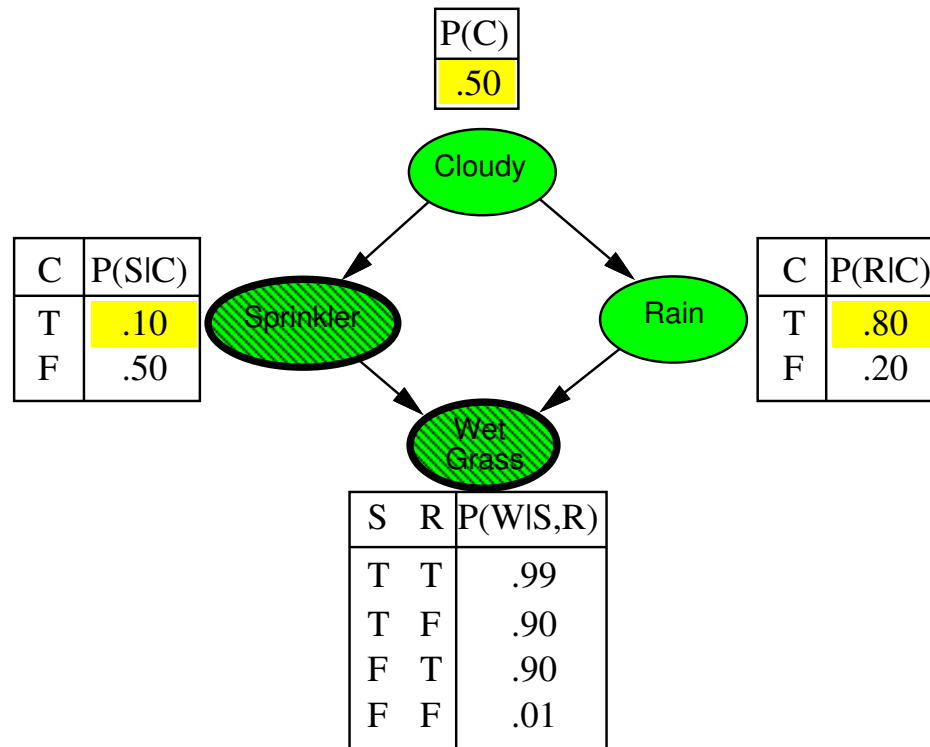
$$w = 1.0$$

Exemplo de pesagem por verosimilhança



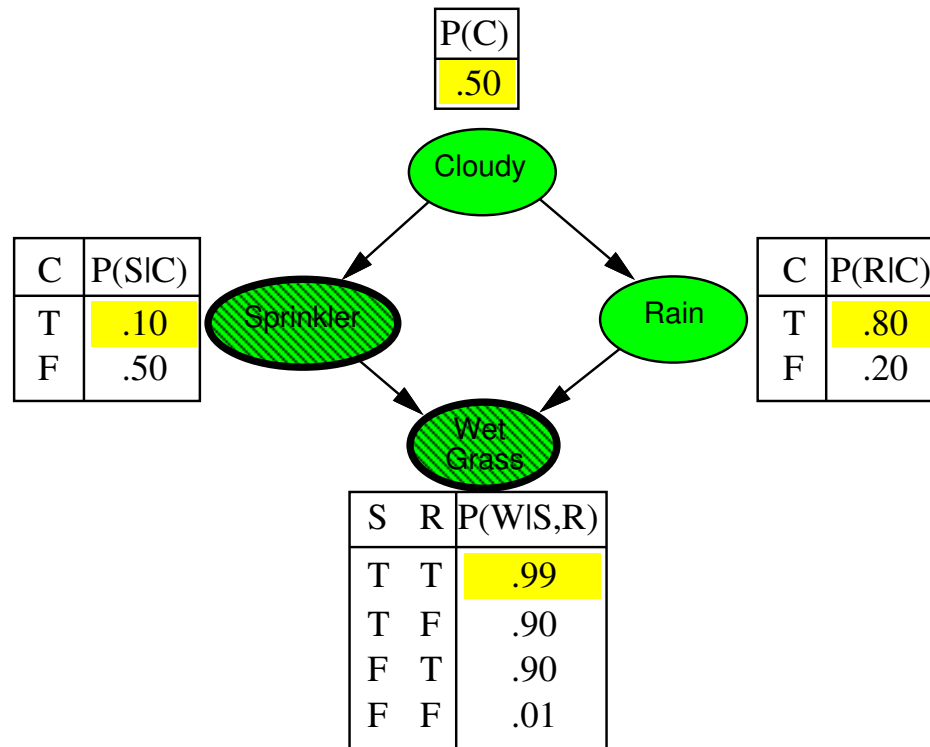
$$w = 1.0 \times 0.1$$

Exemplo de pesagem por verosimilhança



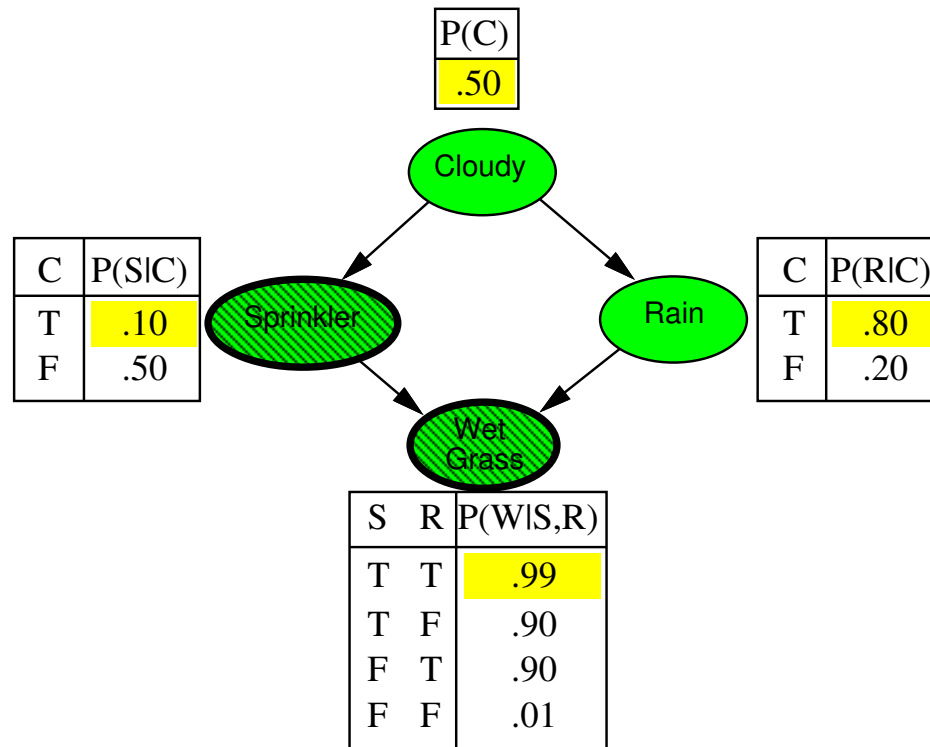
$$w = 1.0 \times 0.1$$

Exemplo de pesagem por verosimilhança



$$w = 1.0 \times 0.1$$

Exemplo de pesagem por verosimilhança



$$w = 1.0 \times 0.1 \times 0.99 = 0.099$$

Análise de pesagem por verosimilhança

Probabilidade de amostragem de WEIGHTEDSAMPLE é

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

Nota: só entra em conta apenas com a evidência dos **antecessores**

⇒ algures “entre” a
distribuição a priori e a posteriori

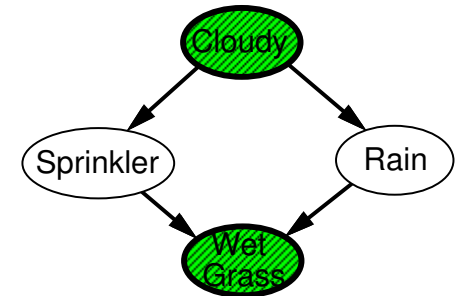
Peso para uma dada amostra $\mathbf{z}, \mathbf{e}$ é

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

Probabilidade de amostragem pesada é

$$\begin{aligned} & S_{WS}(\mathbf{z}, \mathbf{e}) w(\mathbf{z}, \mathbf{e}) \\ &= \prod_{i=1}^l P(z_i | \text{Parents}(Z_i)) \prod_{i=1}^m P(e_i | \text{Parents}(E_i)) \\ &= P(\mathbf{z}, \mathbf{e}) \text{ (pela semântica global da rede)} \end{aligned}$$

Portanto a pesagem por verosimilhança retorna estimativas consistentes mas o desempenho continua a degradar-se com muitas variáveis evidência porque algumas amostras têm quase todo o peso total



Markov Chain Monte Carlo (MCMC)

“Estado” da rede = atribuição corrente a todas as variáveis.

Gerar o estado seguinte amostrando uma variável dado o seu Markov blanket
Altera-se uma variável de cada vez, por amostragem, mantendo a evidência.

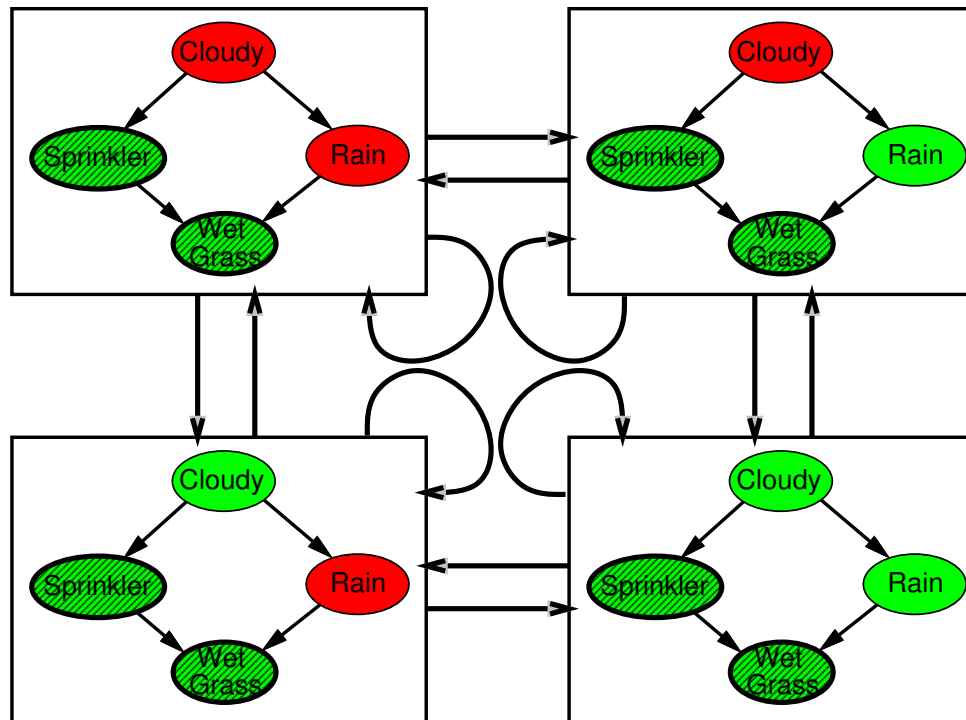
```
function MCMC-Ask( $X, e, bn, N$ ) returns an estimate of  $P(X|e)$ 
  local variables:  $N[X]$ , a vector of counts over  $X$ , initially zero
                   $Z$ , the nonevidence variables in  $bn$ 
                   $x$ , the current state of the network, initially copied from  $e$ 

  initialize  $x$  with random values for the variables in  $Z$ 
  for  $j = 1$  to  $N$  do
     $N[x] \leftarrow N[x] + 1$  where  $x$  is the value of  $X$  in  $x$ 
    for each  $Z_i$  in  $Z$  do
      sample the value of  $Z_i$  in  $x$  from  $P(Z_i|MB(Z_i))$ 
      given the values of  $MB(Z_i)$  in  $x$ 
  return NORMALIZE( $N[X]$ )
```

Pode-se também escolher aleatoriamente a variável a amostrar

A cadeia de Markov

Com *Sprinkler = true*, *WetGrass = true*, existem quatro estados:



Vagueia durante algum tempo, efectua a média do que observa

MCMC exemplo (cont.)

Estimar $\mathbf{P}(\textit{Rain}|\textit{Sprinkler} = \textit{true}, \textit{WetGrass} = \textit{true})$

Amostrar *Cloudy* ou *Rain* dado o seu Markov blanket, repetir.

Contar p número de vezes que *Rain* é verdadeiro e falso nas amostras.

E.g., visita 100 estados

31 tem *Rain = true*, 69 tem *Rain = false*

$$\begin{aligned}\hat{\mathbf{P}}(\textit{Rain}|\textit{Sprinkler} = \textit{true}, \textit{WetGrass} = \textit{true}) \\ = \text{NORMALIZE}(\langle 31, 69 \rangle) = \langle 0.31, 0.69 \rangle\end{aligned}$$

Teorema: cadeia aproxima-se da **distribuição estacionária** :
a fracção de tempo gasto em cada espaço é exactamente
proporcional à sua probabilidade a posteriori

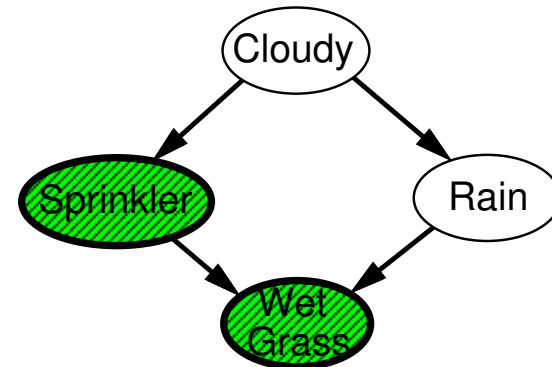
Amostragem do Markov blanket

Markov blanket de *Cloudy* é

Sprinkler e *Rain*

Markov blanket de *Rain* é

Cloudy, *Sprinkler*, e *WetGrass*



Probabilidade dado o Markov blanket é obtida como se segue:

$$P(x'_i | MB(X_i)) = P(x'_i | Parents(X_i)) \prod_{Z_j \in Children(X_i)} P(z_j | Parents(Z_j))$$

Facilmente implementável em sistemas paralelos de troca de mensagens, cérebros

Problemas computacionais principais:

- 1) Dificuldade em detectar que se atingiu a convergência
- 2) Pode ser dispendioso se Markov blanket é grande:

$P(X_i | MB(X_i))$ não varia muito

Sumário

Inferência exacta por eliminação de variáveis:

- tempo polonomial em polytrees, NP-difícil em grafos arbitrários
- espaço = tempo, muito sensível à topologia

Inferência aproximada por PV, MCMC:

- PV comporta-se mal quando existe muita evidência
- PV, MCMC geralmente insensíveis à topologia
- Convergência pode ser muito lenta com probabilidades perto de 0 ou 1
- Pode tratar combinações arbitrárias de variáveis discretas e contínuas

Extensões para linguagem de primeira ordem:

- Relational Probability Models (RPMs)
- Open Universe Probability Models (OUPMs)
- P-log extensão de ASP para lidar com probabilidades