

CS786

Lecture 18: July 2nd, 2012

Bayesian Parameter Estimation
[KF Chapter 17]

CS786 (c) 2012 P. Poupart

1

Bayesian Learning

- Idea: compute posterior over hypotheses given data
$$\Pr(H|d) \propto \Pr(H) \Pr(d|H)$$
- Recall candy example: $H \in \{h_1, \dots, h_5\}$
 - h_1 : 100% lime, h_2 : 75% lime, ..., h_5 : 0% lime
 - Finitely many hypotheses
- What if we have infinitely many hypotheses?
 - E.g., continuum of parameters

CS786 (c) 2012 P. Poupart

2

Simplified Candy Example

- One-node Bayesian network:
- Parameter: $\theta = \Pr(F = \text{cherry})$
- Continuum of hypotheses: $0 \leq \theta \leq 1$
- What is a good family of distributions over the $[0,1]$ interval?



CS786 (c) 2012 P. Poupart

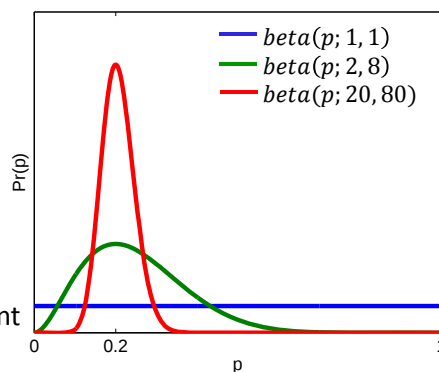
3

Beta Distribution

- Distribution over Bernoulli
- Definition: $\text{beta}(p; \alpha, \beta) = k p^{\alpha-1} (1-p)^{\beta-1}$

$$\text{where } k = \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha)\Gamma(\beta)}$$

- Mean: $\frac{\alpha}{\alpha+\beta}$
- Precision: $\alpha + \beta$
- Interpretation:
 - $\alpha - 1$: freq. of p -event
 - $\beta - 1$: freq. of $(1-p)$ -event



CS786 (c) 2012 P. Poupart

4

Conjugate Prior

- Prior: Let $\Pr(\theta) = \text{beta}(\theta; \alpha, \beta)$
- Suppose that we eat c cherries and l limes
- Posterior: $\Pr(\theta|c, l) \propto \Pr(\theta) \theta^c (1 - \theta)^l$
 $= k \theta^{\alpha-1} (1 - \theta)^{\beta-1} \theta^c (1 - \theta)^l$
 $= k \theta^{\alpha+c-1} (1 - \theta)^{\beta+l-1}$
 $\propto \text{beta}(\theta; \alpha + c, \beta + l)$
- Definition: a family of distributions is said to be a **conjugate prior** when it is closed under Bayesian learning for a certain type of likelihood distribution.
 - E.g., the Beta distribution is a conjugate prior for the Bernoulli distribution

CS786 (c) 2012 P. Poupart

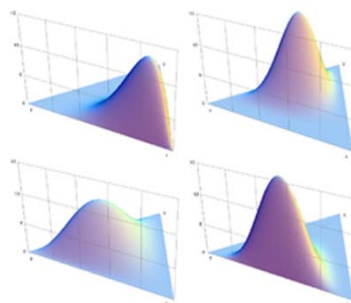
5

Dirichlet Distribution

- Generalization of the Beta to multivalued variables
- Let $\boldsymbol{\theta} = \langle \theta_1, \theta_2, \dots, \theta_m \rangle$ such that $\sum_{i=1}^m \theta_i = 1$
- **Dirichlet:** $\text{Dir}(\boldsymbol{\theta}; \mathbf{n}) = k \prod_{i=1}^m \theta_i^{n_i-1}$

Where $\mathbf{n} = \langle n_1, n_2, \dots, n_m \rangle$

and $k = \frac{\Gamma(\sum_i n_i)}{\prod_i \Gamma(n_i)}$

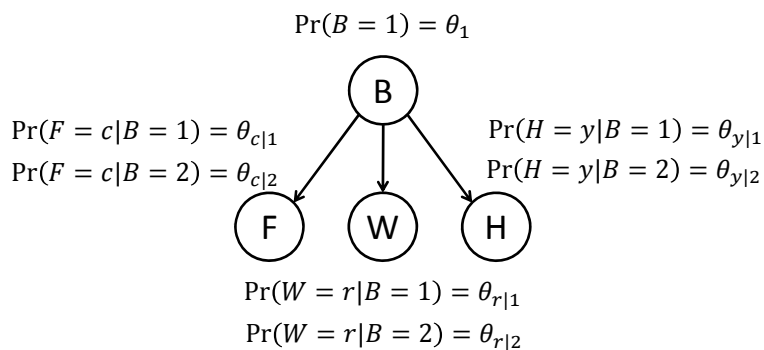


CS786 (c) 2012 P. Poupart

6

Bayesian Network

- Candy example

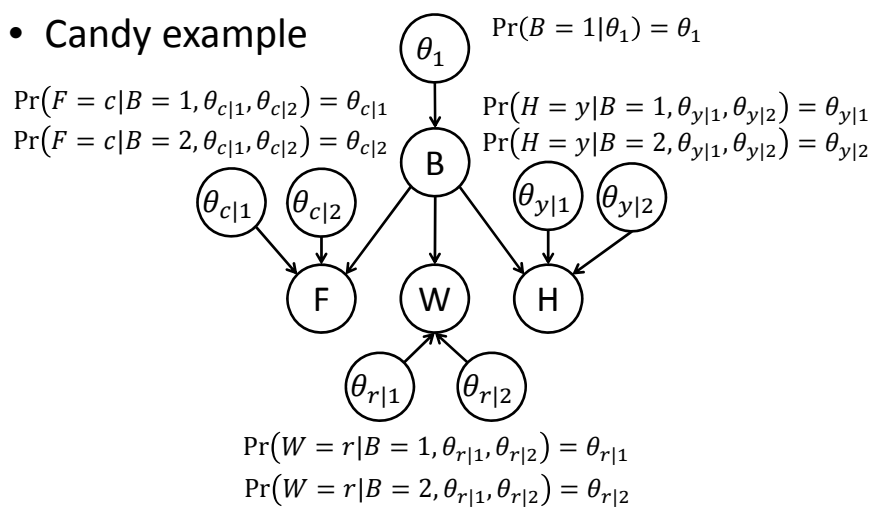


CS786 (c) 2012 P. Poupart

7

Augmented Bayesian Network

- Candy example

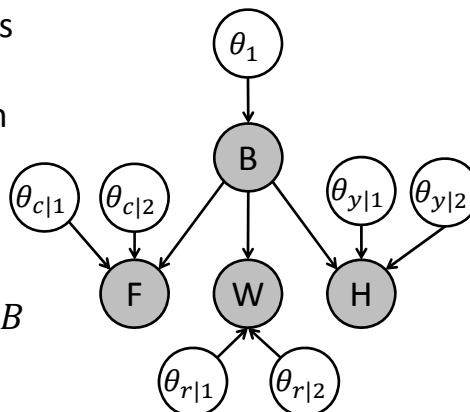


CS786 (c) 2012 P. Poupart

8

Supervised Learning

- D-separation allows us to estimate the parameters of each CPT in isolation
- E.g., $\theta_{c|1}$ and $\theta_{c|2}$ depends on F and B only



CS786 (c) 2012 P. Poupart

9

Supervised Learning

- Start with Dirichlet priors over all θ 's.
e.g., $\Pr(\theta_{c|1}) = \text{beta}(\theta_{c|1}; \alpha, \beta)$
- Compute posterior over each parameter
e.g., $\Pr(\theta_{c|1} | \text{data}) \propto \Pr(\theta_{c|1}) \Pr(\text{data} | \theta_{c|1})$

$$= k \theta_{c|1}^{\alpha-1} (1 - \theta_{c|1})^{\beta-1} \theta_{c|1}^{\#(c,1)} (1 - \theta_{c|1})^{\#(l,1)}$$

$$= k \theta_{c|1}^{\alpha+\#(c,1)-1} (1 - \theta_{c|1})^{\beta+\#(l,1)-1}$$

$$\propto \text{beta}(\theta_{c|1}; \alpha + \#(c,1), \beta + \#(l,1))$$

CS786 (c) 2012 P. Poupart

10

Prediction

- Query: $\Pr(F = c | B = 1)$?

- Average over hypotheses

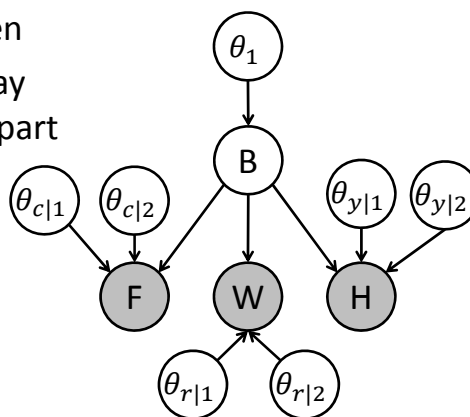
$$\begin{aligned}
 & \Pr(F = c | B = 1) \\
 &= \int_{\theta_{c|1}} \Pr(F = c | B = 1, \theta_{c|1}) \Pr(\theta_{c|1} | \text{data}) d\theta_{c|1} \\
 &= \int_{\theta_{c|1}} \theta_{c|1} \Pr(\theta_{c|1} | \text{data}) d\theta_{c|1} \\
 &= E_{\text{beta}(\theta_{c|1}; \alpha + \#(c,1), \beta + \#(l,1))}[\theta_{c|1}] \\
 &= \frac{\alpha + \#(c,1)}{\alpha + \beta + \#(c,1) + \#(l,1)} \quad \leftarrow \text{smoothed relative frequencies}
 \end{aligned}$$

CS786 (c) 2012 P. Poupart

11

Unsupervised Learning

- Suppose B is hidden
- Each parameter may depend on a large part of the network
- E.g. $\theta_{c|1}$ depends on everything



CS786 (c) 2012 P. Poupart

12

Unsupervised Learning

- Posterior distribution of each parameter:
 - **Closed form:** mixture of Dirichlets
 - **Intractable:** # of mixture components grows exponentially with the amount of data.
- Solution: approximations
 - Gibbs sampling
 - Expectation propagation

CS786 (c) 2012 P. Poupart

13

Gibbs Sampling

- Sample each hidden variable in turn given values for all the other variables

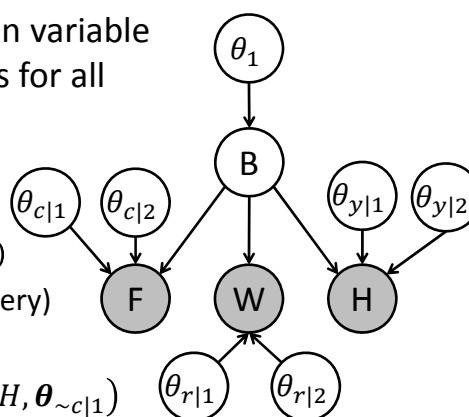
- E.g. sample

$$B \sim \Pr(B|F, W, H, \theta)$$

(regular inference query)

$$\theta_{c|1} \sim \Pr(\theta_{c|1}|F, W, H, \theta_{\sim c|1})$$

(sample from posterior Dirichlet)



CS786 (c) 2012 P. Poupart

14