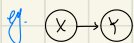




Maximum Likelihood Estimation for Bayesian Network



X	Y	
x^1	y^1	$\theta_{y^1 x^1}$
x^1	y^2	$\theta_{y^2 x^1}$
x^2	y^1	$\theta_{y^1 x^2}$
x^2	y^2	$\theta_{y^2 x^2}$

given a tuple $D = \{ \langle x[m], y[m] \rangle \}$

$$L(\theta; D) = \prod_{m=1}^M P(x[m], y[m]; \theta)$$

$$= \prod_{m=1}^M P(x[m]; \theta) \cdot P(y[m] | x[m]; \theta)$$

$$= \left(\prod_{m=1}^M P(x[m]; \theta) \right) \cdot \left(\prod_{m=1}^M P(y[m] | x[m]; \theta) \right)$$

each of these terms is a "local likelihood" that measures how well the variables is predicted given its parents

$$\ell(\theta; D) = \sum_{m=1}^M \log P(x[m]; \theta) + \sum_{m=1}^M \log P(y[m] | x[m]; \theta)$$

$$= M E[X] \cdot \log \theta_{x^1} + M E[X] \cdot \log \theta_{x^2} + \sum_{x,y \in \text{val}(X,Y)} M E[X,Y] \cdot \log \theta_{y|x}$$

$$\min_{\theta} -C^T \log \theta$$

$$I\theta = 1$$

$$\ell(\theta; D) = -C^T \log \theta + V(I^T \theta - 1)$$

$$\nabla_{\theta} \ell = -C \cdot \frac{1}{\theta} + V \cdot 1 \stackrel{!}{=} 0$$

$$V = \frac{C}{\theta}$$

$$\theta = \frac{C}{V}$$

$$g(V) = -C^T (\log C - \log V) + V(I^T \frac{C}{V} - 1)$$

$$= -C^T \log C + C^T \log V + V \cdot I^T \frac{C}{V} - V$$

$$= -C^T \log C + C^T \log V + I^T C - V$$

$$\nabla_V g = \frac{I^T C}{V} - 1 \stackrel{!}{=} 0$$

$$V = I^T C$$

$$\theta = C / I^T C$$

$$\theta_{x^1} = \frac{M E[X]}{M E[X] + M E[X]}$$

$$\theta_{x^2} = \frac{M E[X]}{M E[X] + M E[X]}$$

$$\theta_{y|x} = \frac{M E[Y]}{\sum_{y \in \text{val}(Y)} M E[Y, x]} = \frac{M E[Y, x]}{M E[X]}$$

$$\hat{\theta}_{x,u} = \frac{M E[X, u]}{\sum_{x \in \text{val}(X)} M E[X, u]} = \frac{M E[X, u]}{M E[u]}$$

$\nexists x \notin \text{val}(X) \quad u \in \text{val}(\text{params}(X))$

Gaussian Bayesian Networks

$$P(X|u) = N(\beta_0 + \beta_1 u_1 + \dots + \beta_k u_k; \sigma^2)$$

$$\ell(\theta; D, u) = \sum_{i=1}^n \log N(x[i]; \beta_0 + \beta_1 u_1 + \dots + \beta_k u_k, \sigma^2)$$

$$= \sum_{i=1}^n -\frac{1}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} (\beta_0 + \beta_1 u_1 + \dots + \beta_k u_k - x[i])^2$$

$$\nabla_{\beta} \ell = \sum_{i=1}^n -\frac{1}{\sigma^2} (\beta_0 + \beta_1 u_1 + \dots + \beta_k u_k - x[i]) \stackrel{!}{=} 0$$

$$\sum_{i=1}^n x[i] = \sum_{i=1}^n \beta_0 + \dots + \beta_k u_k[i] = \beta_0 \cdot n + \beta_1 \sum_{i=1}^n u_1[i] + \dots + \beta_k \sum_{i=1}^n u_k[i]$$

$$E[X] = \beta_0 + \beta_1 \cdot E[u_1] + \dots + \beta_k \cdot E[u_k]$$

$$\nabla_{\beta} \ell = \sum_{m=1}^M -\frac{1}{\sigma^2} (\beta_0 + \beta_1 u_m[m] - x_m[m]) \cdot u_m[m] = 0$$

$$\frac{1}{M} \sum_{m=1}^M x_m[m] \cdot u_m[m] = \frac{1}{M} \sum_{m=1}^M \beta_0 \cdot u_m[m] + \dots + \frac{1}{M} \sum_{m=1}^M \beta_K u_m[m] \cdot u_m[m]$$

$$E_0[X \cdot u_i] = \beta_0 E_0[u_i] + \beta_1 E_0[u_i \cdot u_i] + \dots + \beta_K E_0[u_i u_K]$$

$$\begin{bmatrix} 1 & E[u_1] & \dots & E[u_K] \\ E[u_1] & E[u_1 u_1] & \dots & E[u_1 u_K] \\ E[u_2] & E[u_2 u_1] & \dots & E[u_2 u_K] \\ \vdots & \vdots & \ddots & \vdots \\ E[u_K] & E[u_K u_1] & \dots & E[u_K u_K] \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \\ \beta_2 \\ \vdots \\ \beta_K \end{bmatrix} = \begin{bmatrix} E[X] \\ E[X u_1] \\ E[X u_2] \\ \vdots \\ E[X u_K] \end{bmatrix}$$

Solve for β by solving linear equation

$$\text{Cov}[X, u_i] = E_0[X \cdot u_i] - E_0[X] E_0[u_i]$$

$$= \beta_0 E_0[u_i] + \beta_1 E_0[u_i \cdot u_i] + \dots + \beta_K E_0[u_i u_K]$$

$$- \beta_0 E_0[u_i] - \beta_1 E_0[u_i] E_0[u_1] - \dots - \beta_K E_0[u_K] E_0[u_i]$$

$$= \beta_1 \text{Cov}[u_1, u_i] + \dots + \beta_K \text{Cov}[u_K, u_i]$$

$$\ell(\theta; D) = \sum_{i=1}^M -\frac{1}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} (\beta_0 + \beta_1 u_i[1] - x_i[1])^2$$

$$\nabla_{\beta} \ell = \sum_{m=1}^M -\frac{1}{2} \frac{2(\beta_0 + \beta_1 u_m[1] - x_m[1])}{2\sigma^2} = -\frac{1}{\sigma^2} (\beta_0 + \beta_1 u_m[1] - x_m[1]) = 0$$

$$\sum_{m=1}^M \frac{1}{\sigma^2} (\beta_0 + \beta_1 u_m[1] - x_m[1]) = 0$$

$$\hat{\beta}^* = \text{Cov}[X, X]^{-1} \sum_{i=1}^M \beta_i \text{Cov}[u_i, X]$$

MLE as M-Projection

$$\text{let } D \text{ be a dataset, } \ell(\theta; D) = \log L(\theta; D) = M \cdot E_{x \sim P_0} [\log P(x; \theta)]$$

$$\ell(\theta; D) = \sum_{m=1}^M \log P(X_m[m], \theta)$$

$$= \sum_{x \in \text{val}(X)} \left[\sum_m \mathbb{1}\{X_m[m] = x\} \right] \cdot \log P(x; \theta)$$

$$= \sum_{x \in \text{val}(X)} M \cdot \hat{P}_0(x) \log P(x; \theta)$$

$$= M \cdot E_{x \sim \hat{P}_0} [\log P(x; \theta)]$$

$$D(P \| P') = \int p(x) \log \frac{p(x)}{p'(x)} dx = E_{x \sim p} [\log \frac{p(x)}{p'(x)}] = -H(p) - E_{x \sim p} [\log p'(x)]$$

$$\ell(\theta; D) = M \cdot (-H_{\hat{P}_0}(X) - D(\hat{P}_0 \| P))$$

The maximum likelihood estimation θ in a parametric family is the param of M-projection of $\hat{P}_0$ on P_0

Joint Probabilistic Model

encode the prior knowledge about θ with a probability distribution, which represents how likely of different choices of θ are

then create a joint distribution of parameter θ and observation $\{X[m]\}$

eg. coin toss

$$P(X[m]=1|\theta) = \theta$$

$$P(X[m]=0|\theta) = 1-\theta$$

$$\theta \sim \text{Uniform}(0,1)$$

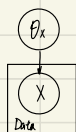
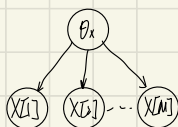


plate model



Ground Bayesian Network

$$X[i] \perp X[j] | \theta_x$$

$$P(\theta, X[1] \dots X[M]) = P(\theta) \cdot P(X[1] \dots X[M] | \theta)$$

$$= P(\theta) \cdot \prod_{m=1}^M P(X[m] | \theta)$$

$$= P(\theta) \cdot \theta^{M[1]} \cdot (1-\theta)^{M[0]}$$

posterior:
$$P(\theta | X[1] \dots X[M]) = \frac{P(\theta, X[1] \dots X[M])}{P(X[1] \dots X[M])}$$

In prediction, instead of selecting a single value for θ we use it for predicting the probability over the next toss

$$P(X[M+1] | X[1] \dots X[M]) = \int P(X[M+1], \theta | X[1] \dots X[M]) d\theta$$

$$= \int P(X[M+1] | \theta, X[1] \dots X[M]) \cdot P(\theta | X[1] \dots X[M]) d\theta$$

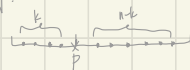
$$= \int P(X[M+1] | \theta) \cdot P(\theta | X[1] \dots X[M]) d\theta$$

$$= \int P(X[M+1] | \theta) \cdot \frac{P(\theta) \cdot P(X[1] \dots X[M] | \theta)}{P(X[1] \dots X[M])} d\theta$$

$$= \frac{1}{P(X[1] \dots X[M])} \int P(X[M+1] | \theta) \cdot \theta^{M[1]} (1-\theta)^{M[0]} d\theta$$

$$\int_0^1 p^k (1-p)^{n-k} dp = \frac{k! (n-k)!}{(n+1)!}$$

Strong proof:



throw k white balls, $n-k$ black balls and 1 red.

$$P(\text{white} | \text{red} | \text{black})$$

$$P(X[M+1]=1 | X[1] \dots X[M]) = \frac{1}{P(X[1] \dots X[M])} \int_0^1 \theta^{M[1]+1} (1-\theta)^{M[0]} d\theta = \frac{1}{2} \frac{(M[1]+1)! M[0]!}{(M[1]+M[0]+2)!}$$

$$P(X[M+1]=0 | X[1] \dots X[M]) = \frac{1}{2} \int_0^1 \theta^{M[0]} (1-\theta)^{M[1]+1} d\theta = \frac{M[1]! (M[0]+1)!}{(M[1]+M[0]+2)!}$$

$$\begin{cases} a+b=1 \\ \frac{a}{b} = \frac{(M[0]+1)! \cdot M[1]!}{M[1]! \cdot (M[0]+1)!} = \frac{M[1]+1}{M[0]+1} \end{cases}$$

$$\Rightarrow P(X[M+1]=1 | X[1] \dots X[M]) = \frac{M[0] + 1}{M[0] + M[1] + 2}$$

Laplace smooth

Priors

appropriate distribution is beta distribution

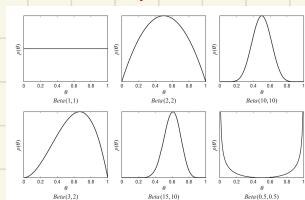
$$\theta \sim \text{Beta}(\alpha_1, \alpha_0) = \frac{1}{\Gamma(\alpha_1)\Gamma(\alpha_0)} (1-\theta)^{\alpha_0-1} \theta^{\alpha_1-1} \quad \theta \in [0,1]$$

where $\frac{\Gamma(\alpha_1+\alpha_0)}{\Gamma(\alpha_1)\Gamma(\alpha_0)}$ is a normalizing constant, $\Gamma(x) = \int_0^{+\infty} t^{x-1} e^{-t} dt$ is the Gamma function

$$\text{Gamma function } \Gamma(x) = \int_0^{+\infty} x^x e^{-x} dx \quad (x > 0)$$

$$\Gamma(n) = (n-1)! \quad n \text{ integer } > 0$$

$$\Gamma(\frac{1}{2}) = \sqrt{\pi}$$



$$\begin{aligned} \Gamma(\frac{1}{2}) &= \int_0^{+\infty} x^{-\frac{1}{2}} e^{-x} dx \quad u = x^{\frac{1}{2}} \\ &= \int_0^{+\infty} u^{-1} e^{-u^2} du \\ &= \frac{1}{2} \int_0^{+\infty} e^{-u^2} du \\ &= \frac{1}{2} \sqrt{\pi} \end{aligned}$$

$$\begin{aligned} \Gamma(n+1) &= \int_0^{+\infty} x^n e^{-x} dx = \int_0^{+\infty} x^n dx \\ &= - \left[x^n e^{-x} \right]_0^{+\infty} - \int_0^{+\infty} e^{-x} dx \\ &= n \int_0^{+\infty} e^{-x} x^{n-1} dx \\ \Gamma(n+1) &= n \cdot \Gamma(n) \\ \Gamma(1) &= 1 \end{aligned}$$

$$\int_0^1 \theta^{\alpha_1-1} (1-\theta)^{\alpha_0-1} d\theta = \frac{(\alpha_1-1)! (\alpha_0-1)!}{(\alpha_1-1 + \alpha_0-1 + 1)!} \quad \int_0^1 \frac{1}{\theta} d\theta = \frac{(\alpha_1+\alpha_0-1)!}{(\alpha_1-1)! (\alpha_0-1)!} = \frac{\Gamma(\alpha_1+\alpha_0)}{\Gamma(\alpha_1)\Gamma(\alpha_0)} \quad (\text{for discrete } \alpha_1, \alpha_0)$$

eg. $(X|\theta) \sim \text{Binomial}(M, \theta) \sim \text{Beta}(a, b)$ find $P(\theta|X)$

$$P(\theta|X=k) = \frac{P(\theta) \cdot P(X=k|\theta)}{P(X)} = \frac{\binom{M}{k} \theta^k (1-\theta)^{M-k} \frac{1}{\Gamma(a)\Gamma(b)} \theta^{a-1} (1-\theta)^{b-1}}{\int_0^1 P(\theta) P(X=k|\theta) d\theta} \propto \theta^{a+k-1} (1-\theta)^{b+M-k-1}$$

$$(\theta|X) \sim \text{Beta}(a+X, b+M-X)$$

prior: $\theta \sim \text{Beta}$, $X|\theta \sim \text{Binomial}$ posterior: $\theta|X \sim \text{Beta}$

Beta is "conjugate prior" to Binomial

$$\begin{aligned} P(X[M+1]|X[1] \dots X[M]) &= \int_0^1 P(X[M+1]|\theta | X[1] \dots X[M]) d\theta \\ &= \int_0^1 P(X[M+1]|\theta, X[1] \dots X[M]) P(\theta | X[1] \dots X[M]) d\theta \end{aligned}$$

$$\begin{aligned} P(X[M+1]=1 | X[1] \dots X[M]) &= \int_0^1 \theta \cdot \frac{1}{\Gamma(a)\Gamma(b)} (1-\theta)^{b+M-1} d\theta \\ &= \frac{1}{\Gamma(a)\Gamma(b)} \int_0^1 \theta^{a+1-1} (1-\theta)^{b+M-1} d\theta = \frac{(a+1)! (b+M-1)!}{(a+b+M)!} \end{aligned}$$

$$\begin{aligned} P(X[M+1]=0 | X[1] \dots X[M]) &= \int_0^1 (1-\theta) \frac{1}{\Gamma(a)\Gamma(b)} (1-\theta)^{b+M-1} d\theta \\ &= \frac{1}{\Gamma(a)\Gamma(b)} \int_0^1 \theta^{a-1} (1-\theta)^{b+M} d\theta = \frac{a! (b+M+1)!}{(a+b+M+1)!} \end{aligned}$$

$$P(X[M+1]=1 | X[1] \dots X[M]) = \frac{a+X}{a+b+M}$$

$$x = \sum_{i=1}^M \mathbb{1}_{\{X[i]=1\}}$$

intuitively, a or b corresponds to the num of imaginary heads and tails we "have seen" before the experiment $\{X[1] \dots X[M]\}$

• Prior and Posteriors

$$D = \{X[1] \dots X[M]\} \text{ iid } X[i] \sim P_\theta$$

$$P(\theta|D) = \frac{P(D|\theta) \cdot P(\theta)}{P(D)} = \frac{\overbrace{P(D|\theta)}^{\text{likelihood}} \cdot \overbrace{P(\theta)}^{\text{prior}}}{\int P(D|\theta) P(\theta) d\theta} \quad \text{marginal likelihood}$$

(a priori probability of seeing the particular dataset D given prior beliefs $P(\theta)$)

A Dirichlet distribution is specified by a set of hyperparameters $\alpha_1 \dots \alpha_k$

$$\theta \sim \text{Dirichlet}(\alpha_1 \dots \alpha_k) \text{ if } P(\theta) \propto \prod \theta_k^{\alpha_k - 1}$$

If $P(\theta)$ is Dirichlet $(\alpha_1 \dots \alpha_k)$, then $P(\theta|D)$ is Dirichlet $(\alpha_1 + M[1], \dots, \alpha_k + M[k])$

A family of priors $P(\theta; \alpha)$ is *conjugate* to a model $P(X|\theta)$ if for any $D = \{X[1] \dots X[M]\}$ of iid samples $X[i] \sim P(X|\theta)$, $\forall \alpha$, there are hyperparameters α' that describe the posterior

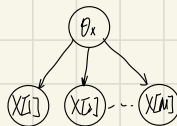
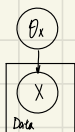
$$P(\theta; \alpha') \propto P(D|\theta) P(\theta; \alpha) \quad (\text{prior and posterior are in the same family})$$

eg. multinomial likelihood and Dirichlet prior

$$\theta \in \mathbb{R}^k, \sum \theta = 1, P(\theta) \propto \prod \theta_k^{\alpha_k - 1}$$

$$P(X=k|\theta) = \theta_k$$

$$L(\theta; D) = P(D|\theta) = \prod_k \theta_k^{M[k]}$$



$$P(\theta|D) = \frac{P(D|\theta) \cdot P(\theta)}{\int P(D|\theta) P(\theta) d\theta} \propto \prod \theta_k^{M[k]} \cdot \theta_k^{\alpha_k - 1}$$

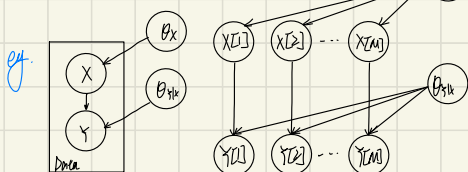
$$\theta|D \sim \text{Dirichlet}(\alpha_1 + M[1], \dots, \alpha_k + M[k])$$

Dirichlet prior are conjugate of multinomial model

Special case: Beta prior are conjugate of Binomial model

• Parameter Independence and Global Decomposition

the instances are independent given unknown parameters



$$X[M], Y[M] \perp X[M'], Y[M'] \mid \theta_x, \theta_{yk}$$

$$\theta_x \perp \theta_{yk} \mid (X[1], Y[1], \dots, X[M], Y[M])$$

knowing the value of one parameter tells nothing about another

let G be a bayesian network with parameter $\theta = (\theta_{x_1|parents}, \dots, \theta_{y|parents})$.

$P(\theta)$ satisfies global parameter independence if it has form $P(\theta) = \prod P(\theta_{x_i|parents})$

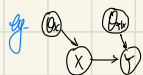
complete data d-separates the parameter for different CPDs

$$P(\theta_x, \theta_y | D) = P(\theta_x | D) \cdot P(\theta_y | D)$$

Once we can store each parameter separately, we can combine the results (like in MLE)

$$\begin{aligned} P(\theta | D) &= \frac{P(D | \theta) \cdot P(\theta)}{P(D)} \\ &= \frac{1}{P(D)} \underbrace{\prod L_i(\theta_{x_i|parents} | D)}_{\text{local likelihoods}} \cdot \underbrace{\prod P(\theta_{x_i|parents})}_{\text{global parameter independence}} \\ &= \frac{1}{P(D)} \prod [L_i(\theta_{x_i|parents} | D) \cdot P(\theta_{x_i|parents})] \end{aligned}$$

let G be a bayesian network over $\mathcal{X}$. $D = \{x_1, \dots, x_n\}$ if $P(\theta)$ satisfies global parameter independence

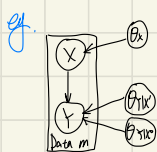


$$\begin{aligned} P(x_1, \dots, x_n, y_1, \dots, y_n | D) &= \int P(x_1, \dots, x_n, y_1, \dots, y_n | \theta, D) d\theta = \int P(x_1, \dots, x_n, y_1, \dots, y_n | \theta) \cdot P(\theta | D) d\theta \\ &= \int \int_{\theta_x \theta_y} P(x_1, \dots, x_n | \theta_x) \cdot P(y_1, \dots, y_n | \theta_y, x_1, \dots, x_n) P(\theta_x | D) P(\theta_y | D) d\theta_x d\theta_y \\ &= \int_{\theta_x} P(x_1, \dots, x_n | \theta_x) P(\theta_x | D) d\theta_x \cdot \int_{\theta_y} P(y_1, \dots, y_n | \theta_y, x_1, \dots, x_n) P(\theta_y | D) d\theta_y \end{aligned}$$

if the priors for parameters of different CPDs are independent

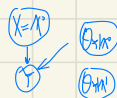
$$P(x_1, \dots, x_n, y_1, \dots, y_n | D) = \prod \int P(x_i | \theta_{x_i|parents}, \theta_{y|parents}) \cdot P(\theta_{x_i|parents} | D) d\theta_{x_i|parents}$$

Local Decomposition

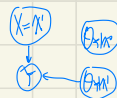


$$P(y | x_1, \dots, x_n, \theta_x, \theta_y) = \begin{cases} \theta_{y|x} & x_1 = x_1 \\ \theta_{y|x} & x_1 = x_1 \end{cases}$$

θ_x and θ_y are dependent given the y (not given x)



θ_x and θ_y are independent given the data



$$P(\theta_x | D) = P(\theta_x | D) \cdot P(\theta_y | D)$$

$$\begin{aligned}
 P(\theta, D) &= P(\theta) \cdot P(D|\theta) \\
 &= P(\theta_x) \cdot P(\theta_{yx}) \cdot \prod_{m=1}^M P(x[m]|\theta_x) \cdot P(y[m]|x[m], \theta_{yx}) \\
 &= P(\theta_x) \cdot L_x(\theta_x; D) \cdot P(\theta_{yx}) \cdot \prod_{m: x[m] \neq x} P(y[m]|x[m], \theta_{yx}) \\
 &\quad P(\theta_{yx}) \cdot \prod_{m: x[m] = x} P(y[m]|x[m], \theta_{yx})
 \end{aligned}$$

if $\theta \sim \text{Dirichlet}(\alpha_1, \dots, \alpha_K)$ $P(\theta) = \prod_k \theta_k^{\alpha_k - 1}$ $P(\theta_{yx}) = \prod_k \text{Dir}(k)$

$$\begin{aligned}
 P(\theta_{yx}|D) &= \int \int P(\theta_{yx}|D) \cdot P(\theta_x|D) \cdot P(\theta_{yx}|D) d\theta_x d\theta_{yx} \\
 &= \int \int P(\theta|D) d\theta_x d\theta_{yx} \\
 &\propto P(\theta_{yx}) \prod_{m: x[m] \neq x} P(y[m]|x[m], \theta_{yx}) \\
 &= \left(\prod_k \text{Dir}(k) \right)^{N_{\text{Dir}}(k)-1} \cdot \theta_{yx}^{N_{\text{Dir}}(x)-1} \cdot \theta_{yx}^{N_{\text{Dir}}(y)} \\
 &= \text{Dirichlet}(\theta_{yx} + N_{\text{Dir}}(y), \theta_{yx} + N_{\text{Dir}}(x))
 \end{aligned}$$

let X be a variable with parents U , $P(\theta_{xu})$ satisfies local parameter independence if $P(\theta_{xu}) = \prod_{u \in \text{children}(x)} P(\theta_{xu}|D)$

If the prior $P(\theta)$ satisfies global and local parameter independence, then $P(\theta|D) = \prod_i \prod_{p \in \text{pa}(x_i)} P(\theta_{x_i p}|D)$

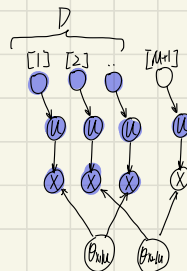
$$\begin{aligned}
 P(\theta|D) &= \prod_i P(\theta_{x_i \text{pa}(x_i)}|D) \quad // \text{global parameter independence} \\
 &= \prod_{p \in \text{pa}(x_i)} P(\theta_{x_i p}|D) \quad // \text{local parameter independence}
 \end{aligned}$$

If $P(\theta_{xu})$ is a Dirichlet distribution $P(\theta_{xu}) \propto \prod_{k \in \text{children}(x)} \theta_{x_k u(k)}^{\alpha_{x_k u(k)} - 1}$
 then $P(\theta_{xu}|D)$ is a Dirichlet distribution $P(\theta_{xu}) \propto \prod_{k \in \text{children}(x)} \theta_{x_k u(k)}^{\alpha_{x_k u(k)} + N(x_k, u)}$

If U is a parent of X $\theta_{xu} \sim \text{Dirichlet}(\alpha_{xu1}, \dots, \alpha_{xuK})$, D is complete data set over X

$$P(x_{i+1} = x_i | U_{i+1} = u, D) = \frac{\alpha_{xu, u} + N(x_i, u)}{\sum_j (\alpha_{xu, j} + N(x_j, u))}$$

$$\begin{aligned}
 P(x_{i+1} = x_i | U_{i+1} = u, D) &= \int P(x_{i+1} = x_i, \theta_{xu} | U_{i+1} = u, D) d\theta_{xu} \\
 &= \int P(x_{i+1} = x_i | \theta_{xu}, U_{i+1} = u, D) P(\theta_{xu} | U_{i+1} = u, D) d\theta_{xu} \\
 &= \int \underbrace{P(x_{i+1} = x_i | \theta_{xu}, U_{i+1} = u)}_{(x_{i+1} \perp D | \theta_{xu}) \text{ blocks the path}} \cdot \underbrace{P(\theta_{xu} | D)}_{(U_{i+1} \perp D | \theta_{xu}) \text{ blocks the path}} d\theta_{xu} \\
 &= \int \theta_{xu} \frac{P(\theta_{xu}) P(D|\theta_{xu})}{P(D)} d\theta_{xu}
 \end{aligned}$$



$$\begin{aligned}
 &= \frac{\gamma}{P(D)} \int \theta_{xu} \prod_{j \in I} \theta_{yju}^{\alpha_{yju}} \prod_{j \in I} \theta_{yju}^{M[D_{y,j}, u]} d\theta_{xu} \quad P(D|\theta_{xu}) \hookrightarrow \text{chain rule} \\
 &= \frac{\gamma}{P(D)} \int \theta_{xu}^{\alpha_{xu} + M[D_{x,u}, u]} \prod_{j \in I} \theta_{yju}^{\alpha_{yju} + M[D_{y,j}, u]} d\theta_{xu} \quad // \text{integrating a polynomial over a simplex} \\
 &= \frac{\gamma}{P(D)} \cdot \frac{\Gamma(\alpha_{xu} + M[D_{x,u}, u] + 1) \cdot \prod_{j \in I} \Gamma(\alpha_{yju} + M[D_{y,j}, u] + 1)}{\Gamma(\alpha_{xu} + M[D_{x,u}, u] + 1 + \sum_{j \in I} \alpha_{yju} + M[D_{y,j}, u])} \\
 &= \frac{1}{Z} \alpha_{xu} + M[D_{x,u}, u]
 \end{aligned}$$

$$P(X_{I \cup J} = x_i | U_{I \cup J} = u, D) = \frac{\alpha_{x_i | u} + M[D_{x_i, u}]}{\sum_j \alpha_{x_j | u} + M[D_{x_j, u}]}$$

// Integrating over a simplex

$$\text{let } S_n = \{\theta_1, \dots, \theta_n \mid \theta_i \geq 0, \sum \theta_i = 1\}$$

$$\begin{aligned}
 &\int_{S_n} \theta_1^{k_1} \dots \theta_n^{k_n} d\theta \\
 I(t) &= \int_{[0,1]^n} \theta_1^{k_1} \dots \theta_n^{k_n} \delta(t - \theta_1 - \dots - \theta_n) d\theta
 \end{aligned}$$

$$\begin{aligned}
 \text{Laplace transform: } I(s) &= \int_0^{1+\epsilon} I(t) e^{-st} dt \\
 &= \int_0^{1+\epsilon} \int_{[0,1]^n} \theta_1^{k_1} \dots \theta_n^{k_n} \delta(t - \theta_1 - \dots - \theta_n) d\theta e^{-st} dt \\
 &= \int_{[0,1]^n} \theta_1^{k_1} \dots \theta_n^{k_n} \int_0^{1+\epsilon} \delta(t - \theta_1 - \dots - \theta_n) e^{-st} dt d\theta \\
 &= \int_{[0,1]^n} \theta_1^{k_1} \dots \theta_n^{k_n} \exp(-s(\theta_1 + \dots + \theta_n)) d\theta \\
 &= \prod_{i=1}^n \int_0^1 \theta_i^{k_i} \exp(-s\theta_i) d\theta_i \\
 &= \prod_{i=1}^n \frac{k_i!}{s^{k_i+1}} \\
 &= \frac{\prod_{i=1}^n k_i!}{s^{\sum k_i + n}} \\
 &= \frac{\prod_{i=1}^n k_i!}{s^{\sum k_i + n}} \cdot \frac{\prod_{i=1}^n k_i!}{(\sum_{i=1}^n k_i + 1)!} \\
 &= \frac{\prod_{i=1}^n k_i!}{(\sum_{i=1}^n k_i + n + 1)!} \cdot \mathcal{L}(t^{\sum k_i + n})(s)
 \end{aligned}$$

$$\begin{aligned}
 &\int_0^{1+\epsilon} t^{\sum k_i + n} \exp(-st) dt \\
 &= \frac{1}{s} \int_0^{1+\epsilon} t^{\sum k_i + n} \exp(-st) dt \\
 &= -\frac{1}{s} \left[t^{\sum k_i + n} \exp(-st) \right]_0^{1+\epsilon} + \int_0^{1+\epsilon} \exp(-st) \cdot t^{\sum k_i + n} dt \\
 &= \frac{1}{s} \int_0^{1+\epsilon} t^{\sum k_i + n} \exp(-st) dt \\
 &= \frac{1}{s^{k_i+1}} \\
 \mathcal{L}(t^k)(s) &= \frac{k!}{s^{k+1}}
 \end{aligned}$$

$$\int \theta_1^{k_1} \dots \theta_n^{k_n} d\theta = I(1) = \frac{\prod_{i=1}^n k_i!}{(\sum_{i=1}^n k_i + n + 1)!} = \frac{\prod_{i=1}^n \Gamma(k_i + 1)}{\Gamma(\sum_{i=1}^n (k_i + 1))}$$

$$\int \theta_1^{k_1} \dots \theta_n^{k_n} d\theta = \frac{\prod_{i=1}^n \Gamma(k_i + 1)}{\Gamma(\sum_{i=1}^n k_i + 1)}$$

• Priors for Bayesian Networks

In a Bayesian network $\mathcal{B}$ over $\mathcal{X}$, for each var $\textcircled{x} \in \mathcal{B}$, x has a set of multinomial distribution θ_{xu} $x \in \text{val}(U)$

can see the prior of θ_{xu} as a Dirichlet prior

$$\theta_{xu} \sim \text{Dirichlet}(\#I_{x,u}, \dots, \#I_{x,u})$$

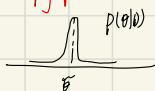
MAP Estimation

the parameters that maximize the posterior probability

$$\hat{\theta} = \underset{\theta}{\operatorname{argmax}} \log p(\theta|D)$$

when we have a large amount of data, the posterior is often sharply peaked around $\hat{\theta}$

$$P(X_{1:n+1}|D) = \int P(X_{1:n+1}|\theta) p(\theta|D) d\theta \approx P(X_{1:n+1}|\hat{\theta})$$



$$\underset{\theta}{\operatorname{argmax}} \log p(\theta|D) = \underset{\theta}{\operatorname{argmax}} \log \frac{p(\theta) p(D|\theta)}{p(D)}$$

$$= \underset{\theta}{\operatorname{argmax}} \underbrace{\log p(\theta)} + \underbrace{\log p(D|\theta)} \quad \text{use the prior to provide regularization over log-likelihood.}$$

$$= \underset{\theta}{\operatorname{argmax}} \log p(\theta) + \left(\sum_{x \in D} \log P(x|\theta) \right) \quad \text{when } D \text{ is a big dataset, } \log p(\theta) \text{ is negligible}$$

eg. Beta prior and Bernoulli model

$$p(\theta; \alpha, \beta) = c \cdot \theta^{\alpha-1} (1-\theta)^{\beta-1}$$

$$p(x) = (1-\theta) \exp\{ \sum_{i=1}^n x_i \}, \quad \ln \frac{p}{1-p} >$$

$$\text{let } \eta = \ln \frac{p}{1-p} \quad \theta = \frac{1}{1+e^\eta} \quad 1-\theta = \frac{1}{1+e^\eta}$$

$$\begin{aligned} p(\eta) &= \frac{d\theta}{d\eta} p(\theta|\eta) = \frac{d\theta}{d\eta} p(\theta|\eta) = 1 \cdot \left(\frac{1}{1+e^\eta} \right)^{\alpha} \cdot e^\eta + c \cdot \left(\frac{1}{1+e^\eta} \right)^{\alpha-1} \left(\frac{1}{1+e^\eta} \right)^{\beta-1} \\ &= \frac{1}{1+e^\eta} \cdot \frac{e^\eta}{1+e^\eta} \cdot c \cdot \left(\frac{1}{1+e^\eta} \right)^{\alpha-1} \left(\frac{1}{1+e^\eta} \right)^{\beta-1} \\ &= c \cdot \left(\frac{1}{1+e^\eta} \right)^{\alpha} \left(\frac{1}{1+e^\eta} \right)^{\beta} \end{aligned}$$

when $\alpha = \beta = 1$, $p(\theta) = 1$ for $\theta \in [0, 1]$, uniform

$$\text{but } p(\eta) = \frac{1}{1+e^\eta} \cdot \frac{1}{1+e^\eta} = \frac{1}{1+e^\eta+e^\eta} = \frac{1}{2+e^\eta+e^\eta} \text{ is far from uniform}$$

both MLE and Bayesian estimation are "representation-independent". (by change of variable in the integral)

but MAP estimation is not

$$\hat{\theta} = \underset{\theta}{\operatorname{argmax}} \log p(\theta) = \underset{\theta}{\operatorname{argmax}} (\alpha-1) \log \theta + (\beta-1) \log (1-\theta) = \frac{\alpha-1}{\alpha+\beta-2}$$

$$\hat{\eta} = \underset{\eta}{\operatorname{argmax}} p(\eta) = \underset{\eta}{\operatorname{argmax}} \frac{d\theta}{d\eta} \quad \theta(\eta) = \frac{1}{1+e^\eta} \quad \hat{\theta} \neq \theta(\hat{\eta})$$

MAP estimation is more sensitive to choices in formalizing the likelihood and the prior than MLE or Bayesian inference