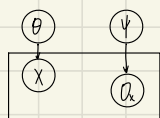


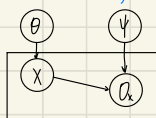


Likelihood of Data and Observation Models

Let $X = \{X_1, \dots, X_n\}$ be some set of random variable, and $O_x = \{O_{x1}, \dots, O_{xn}\}$ be their observability variable
 the observability model is a joint distribution $P_{\text{missy}}(X, O_x) = P(X) \cdot P_{\text{missy}}(O_x | X)$, where $P(X)$ is parameterized by θ and $P_{\text{missy}}(O_x | X)$ is parameterized by ψ



random missing values



deliberate missing values

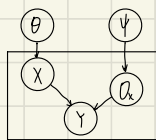
Define a new set of variable $Y = \{Y_1, \dots, Y_n\}$, $\text{Val}(Y_i) = \text{Val}(X_i) \cup \{?\}$

the actual observation is Y , which is a deterministic function of X and O_x

$$Y_i = \begin{cases} X_i & O_{xi} = O' \\ ? & O_{xi} = O'' \end{cases} \quad \text{"?" represents a missing value}$$

eg. $X_i \sim \text{Bernoulli}(\theta)$ $O_{xi} \sim \text{Bernoulli}(\psi)$

$$P(Y=1) = \theta \quad P(Y=0) = (1-\theta)\psi \quad P(Y=?) = 1-\psi$$



Given $Y=?$ O_x must be O''

Given $Y \neq ?$ O_x must be O' , $X=Y$

$$L(\theta, \psi; D) = \psi^{|\text{MCER}| + |\text{MCER}'|} (1-\psi)^{|\text{MCER}''|} \cdot \theta^{|\text{MCER}'|} (1-\theta)^{|\text{MCER}''|}$$

$$\max_{\theta, \psi} L(\theta, \psi; D) \Rightarrow \begin{cases} \max_{\psi} \psi^{|\text{MCER}| + |\text{MCER}'|} (1-\psi)^{|\text{MCER}''|} \\ \max_{\theta} \theta^{|\text{MCER}'|} (1-\theta)^{|\text{MCER}''|} \end{cases}$$

$$\psi^* = \frac{|\text{MCER}'| + |\text{MCER}''|}{|\text{MCER}'| + |\text{MCER}''| + |\text{MCER}''|} \quad \theta^* = \frac{|\text{MCER}'|}{|\text{MCER}'| + |\text{MCER}''|}$$

separable $(X \perp O_x | Y)$ even though knowing Y activates the v-structure (context-specific independence)

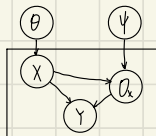
The independence reveals in the likelihood function

eg. $X \sim \text{Bernoulli}(\theta)$ $O_x | X' \sim \text{Bernoulli}(\psi_{\text{aux}})$ $O_x | X' \sim \text{Bernoulli}(\psi_{\text{aux}})$

$$P(Y=1) = \theta \cdot \psi(\alpha|x)$$

$$P(Y=0) = (1-\theta) \cdot \psi(\alpha|x)$$

$$P(Y=?) = \theta \cdot (1-\psi(\alpha|x)) \cdot (1-\theta) \cdot (1-\psi(\alpha|x))$$



knowing $Y=0$, we can infer X from O_x
 $\therefore X \perp O_x | Y$

$$L(\theta, \psi(\alpha|x), \psi(\alpha|x); D) = \theta^{|\text{MCER}'|} \cdot \psi(\alpha|x)^{|\text{MCER}''|} \cdot (1-\theta)^{|\text{MCER}''|} \cdot \psi(\alpha|x)^{|\text{MCER}''|} \cdot [\theta \cdot (1-\psi(\alpha|x)) \cdot (1-\theta) \cdot (1-\psi(\alpha|x))]^{|\text{MCER}''|}$$

Decoupling of Observation Mechanism

A missing data model P_{missy} is missing completely at random (MCAR) if $P_{\text{missy}} = (X \perp O_x)$

In the case of MCAR, the likelihood function decomposes as a product

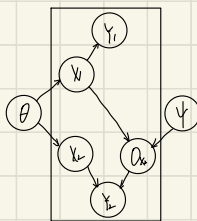
MCAR is sufficient but not necessary for decomposition of likelihood function

eg. always observe the first toss, decide whether or not hide the n th toss based on the first toss

$$p(x_i = x_i^*, y_i = y_i^*) = (1 - \theta_{x_i}) \cdot (1 - \theta_{y_i}) \cdot \psi(x_i = x_i^*)$$

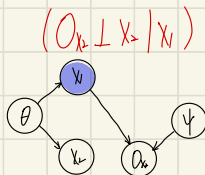
$$P(x_i = x_i^*, y_i = y_i^*) = (1 - \theta_0) \theta_{x_i} \cdot \prod (\theta_{y_i} / x_i = x_i^*)$$

$$P(X = x_i, Y = ?) = (1 - \theta_{10}) \cdot [1 - \psi(0 | x_i, x_i = x_i)]$$



$$L(\theta, \psi; D) = \left[(1 - \theta_u) \cdot (1 - \theta_v) \cdot \psi(O_u | X_u = x_u) \right]^{ML(\psi; \mathcal{Y})} \left[(1 - \theta_u) \cdot \theta_v \cdot \psi(O_u | X_u = x_u) \right]^{ML(\psi; \mathcal{Y})} \left[(1 - \theta_u) \cdot [1 - \psi(O_u | X_u = x_u)] \right]^{ML(\psi; \mathcal{Y})} \\ \left[\theta_u \cdot (1 - \theta_v) \cdot \psi(O_u | X_u = x_u) \right]^{ML(\psi; \mathcal{Y})} \left[\theta_u \cdot \theta_v \cdot \psi(O_u | X_u = x_u) \right]^{ML(\psi; \mathcal{Y})} \left[\theta_u \cdot [1 - \psi(O_u | X_u = x_u)] \right]^{ML(\psi; \mathcal{Y})}$$

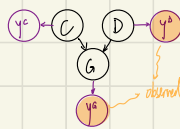
$$= (1 - \theta_{\alpha})^{\overbrace{M[\Psi]}^{M[\Psi]}} \cdot \underbrace{\theta_{\alpha}}_{\theta_{\alpha}}^{\overbrace{M[\Psi]}^{M[\Psi]}} \cdot$$



the conditional independence helps to decouple $P(X)$ from $P(D_i|X)$

let y be a tuple of observations. variables x are partitioned into 2 sets.

$$X_{\text{obs}}^Y = \{X_i : Y_i \neq ?\} \text{ (observed vars)} \text{ and } X_{\text{unobs}}^Y = \{X_i : Y_i = ?\} \text{ (unobserved vars)}$$



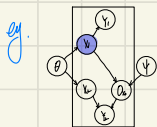
$$X_{\text{die}}^Y = \{D, G\}$$

A missing data model P_{missing} is missing at random (MAR) if

$\forall$ observations y with $\text{Pr}_{\text{missy}}(y) > 0$, $\forall x_{\text{hidden}}^y \in \text{Val}(X_{\text{hidden}}^y)$, $\text{Pr}_{\text{missy}} \models (a_{\perp} \perp x_{\text{hidden}}^y \mid x_{\text{obs}}^y)$

where β_2 are specific values of observation variables, given γ

(given observed vars, whether or not other vars are observed does not depend on the actual value of those vars)



give the value of the first toss.

whether or not the second toss is observed does not depend on the actual value of second toss

$$P_{\text{missing}}(X_{\text{hidden}}^y | X_{\text{obs}}^y, O_z) = P_{\text{missing}}(X_{\text{hidden}}^y | X_{\text{obs}}^y)$$

given observed vars, the observation pattern does not give additional informations about unobserved vars

$$P_{\text{missing}}(y) = \sum_{x_{\text{hidden}}^y} P(x_{\text{obs}}^y, x_{\text{hidden}}^y) \cdot P_{\text{missing}}(a_z | x_{\text{obs}}^y, x_{\text{hidden}}^y)$$

$$= \sum_{x_{\text{hidden}}^y} P(x_{\text{obs}}^y, x_{\text{hidden}}^y) \cdot P_{\text{missing}}(o_z | x_{\text{obs}}^y)$$

$$= \underbrace{P_{\text{missing}}(x_2 | x_{\text{obs}}^y)}_{\text{Missingness}} \cdot \underbrace{p(x_{\text{obs}}^y)}_{\text{Observed}}$$

depends only on ψ depends only on θ

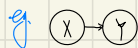
The Likelihood Function

Given a BIV G over a set of var $\mathcal{X}$. each instance has a different set of observed vars

$O(m)$ and $o(m)$ are observed vars and values of mth instance

$H(m)$ is the missing vars in mth instance

$$l(\theta; D) = \log L(\theta; D) = \sum_{m=1}^M \log P(o(m) | \theta)$$



fully observed: $L(\theta; D) = \theta_x^{ME(X)} (1-\theta_x)^{ME(X)} \theta_{yx}^{ME(X,Y)} (1-\theta_{yx})^{ME(X,Y)} \theta_{y|x}^{ME(X,Y)} (1-\theta_{y|x})^{ME(X,Y)}$

$L(\theta; D)$ is log-concave, has a close-form for global maxi.

$L(\theta; D)$ can be decomposed into function of each θ

partially observed: $D = \{ (x=?, y=?), (x=x(m), y=y(m)) \dots (x=x(M), y=y(M)) \}$

$$L(\theta; D[1]) = P(X=x, Y=y) + P(X=x, Y=y)$$

$$= (1-\theta_x) \theta_{yx} + \theta_x \theta_{yx}$$

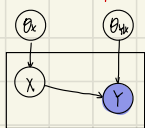
$$L(\theta; D[1:n]) = \theta_x^{ME(X)} (1-\theta_x)^{ME(X)} \theta_{yx}^{ME(X,Y)} (1-\theta_{yx})^{ME(X,Y)} \theta_{y|x}^{ME(X,Y)} (1-\theta_{y|x})^{ME(X,Y)}$$

$$l(\theta; D) = l(\theta; D[1]) + l(\theta; D[2:n])$$

$$= \log [(1-\theta_x) \theta_{yx} + \theta_x \theta_{yx}] + l(\theta; D[2:n])$$

can not be decomposed; and not necessarily log-concave

each partially observed data introduce a log-plus term.



when y is observed and x is not

$\theta_x \rightarrow x \rightarrow y \leftarrow \theta_{yx}$ is an active trail

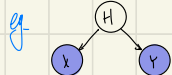
or $x \rightarrow \theta_{yx}$

$$P(\theta | y) = \sum_{x \in \text{nodes}} P(\theta | x, y) = \sum_x \frac{P(\theta) P(x, y | \theta)}{P(y)} = \frac{P(\theta)}{P(y)} \cdot \sum_x P(x, y | \theta) = \frac{P(\theta)}{P(y)} \cdot \underbrace{L(\theta, y)}$$

can not be decomposed as $f(\theta_x) \cdot f(\theta_{yx})$

lose the property of parameter independence

if θ_x is known, $P(\theta_{yx} | y) = C \cdot \log [(1-\theta_x) \theta_{yx} + \theta_x \theta_{yx}]$ the local decomposability is lost



x, y observed

$$P(x, y) = \sum_{h \in \text{nodes}} p(h) p(x|h) p(y|h)$$

$$L(\theta; D) = \prod_{(x,y) \in \text{nodes}} \left[\sum_h p(h) p(x|h) p(y|h) \right]^{ME(x,y)}$$

cannot be decomposed as $p(x|h)$ and $p(y|h)$ (as in mte)

In the presence of partially observed data, we lose all important properties of likelihood function

log-concavity, close-form solution, decomposability of $L(\theta; D)$

Identifiability

ability to identify uniquely a model from data

If randomly choose one from 2 biased coins, then toss

$(\theta_H, \theta_{H|H}, \theta_{H|T})$

$$P(\theta=0) = P(\theta')^{\text{true}} \cdot (1 - P(\theta'))^{\text{true}}$$

$$P(\theta') = \theta_H \cdot \theta_{H|H} + (1 - \theta_H) \cdot \theta_{H|T}$$

$P(\theta)$ is a weighted average of $\theta_{H|H}$ and $\theta_{H|T}$

different choices of $(\theta_H, \theta_{H|H}, \theta_{H|T})$ can give the same likelihood

Given the observation (no matter how many instances), we cannot hope to recover a unique set of parameters

Suppose we have a parametric model with $\theta \in \Theta$ that defines a distribution $P(x|\theta)$

A choice of θ is identifiable if there is no $\theta' \neq \theta$ s.t. $P(x|\theta) = P(x|\theta')$

A model is identifiable if θ is identifiable $\forall \theta \in \Theta$

Gradient Ascent

Let B be a Bayesian Network over $\mathcal{X}$ that induces $P(\mathcal{X})$

Let o be a tuple of observations for some variables

$$\frac{\partial P(o)}{\partial P(x|u)} = \frac{1}{P(x|u)} P(x|u, o) \quad \text{if } P(x|u) > 0, x \in \text{val}(X), u \in \text{val}(U \setminus X)$$

Consider a full assignment ξ .

$$P(\xi) = \prod_{X \in \mathcal{X}} P(X=x_i | \xi \setminus \{X=x_i\})$$

$$\frac{\partial P(\xi)}{\partial P(x|u)} = \begin{cases} \prod_{X \neq X} P(X=x_i | \xi \setminus \{X=x_i\}) = \frac{P(\xi; \theta)}{P(x|u, \theta)} & \text{if } \xi \setminus \{X=x_i\} = (x, u) \\ 0 & \text{otherwise} \end{cases}$$

Consider a partial assignment o

$$P(o) = \sum_{\xi \setminus o = o} P(\xi)$$

$$\frac{\partial P(o)}{\partial P(x|u)} = \sum_{\xi \setminus o = o} \frac{\partial P(\xi; \theta)}{\partial P(x|u, \theta)} = \sum_{\xi \setminus o = o, \xi \setminus \{X=x_i\} = (x, u)} \frac{1}{P(x|u, \theta)} P(\xi; \theta)$$

$$= \frac{1}{P(x|u)} P(o, x, u) \quad P(x|u) \neq 0$$

Let G be a Bayesian network over $\mathcal{X}$ $D = \{o_1, \dots, o_m\}$ be a partially observed dataset

$$\frac{\partial \ell(\theta; D)}{\partial P(x|u)} = \frac{1}{P(x|u)} \sum_{m=1}^M P(x, u | o_m; \theta)$$

$$\begin{aligned} \frac{\partial \ell(\theta; D)}{\partial P(x|u)} &= \frac{\partial}{\partial P(x|u)} \cdot \sum_{m=1}^M \log P(o_m; \theta) = \sum_{m=1}^M \frac{\partial}{\partial P(x|u)} \log P(o_m; \theta) \\ &= \sum_{m=1}^M \frac{1}{P(o_m; \theta)} \cdot \frac{\partial P(o_m; \theta)}{\partial P(x|u)} = \sum_{m=1}^M \frac{1}{P(o_m; \theta)} \cdot \frac{1}{P(x|u)} \cdot P(o_m, x, u; \theta) \\ &= \frac{1}{P(x|u)} \sum_{m=1}^M P(x, u | o_m; \theta) \end{aligned}$$

def compute_gradient(G, θ, D):

for each $\emptyset \rightarrow \emptyset \in G$:

for each $x_i, u_i \in \text{val}(X, U)$

$\bar{\mu}[x_i, u_i] \leftarrow$

} can use similar data structure as Factor

for $m = 1 \dots M$:

run clique tree calibration on $\langle G, \theta \rangle$ with evidence $\sigma[m]$ // $\sum_{m=1}^M P(x, u | \sigma[m]; \theta)$ for each x, u

for each $\emptyset \rightarrow \emptyset \in G$:

for each $x_i, u_i \in \text{val}(X, U)$

$\bar{\mu}[x_i, u_i] += P(x_i, u_i | \sigma[m])$

for each $\emptyset \rightarrow \emptyset \in G$:

for each $x_i, u_i \in \text{val}(X, U)$

$\nabla P(x_i | u_i) = \frac{1}{\theta_{x_i u_i}} \bar{\mu}[x_i, u_i]$

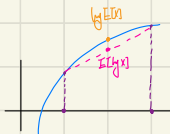
return $\{\nabla P(x_i | u_i)\}$

EM Algorithm

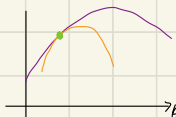


Y is observed.

$$\begin{aligned}
 \sum_m \log P(y[m], \theta) &= \sum_m \log \sum_{x \in \text{val}(X)} P(y[m], x[m]; \theta) \\
 &= \sum_m \log \sum_{x \in \text{val}(X)} \underbrace{Q_m(x[m])}_{\text{construct a probability } Q_m(x) \neq m=1, \dots, M} \frac{P(y[m], x[m]; \theta)}{Q_m(x[m])} \\
 &= \sum_m \log E_{x \sim Q_m} \left[\frac{P(y[m], x[m]; \theta)}{Q_m(x[m])} \right] \\
 &\geq \sum_m E_{x \sim Q_m} \left[\log \frac{P(y[m], x[m]; \theta)}{Q_m(x[m])} \right] \quad \log(\cdot) \text{ is concave, Jensen's Inequality} \\
 &= \sum_m \sum_{x \in \text{val}(X)} Q_m(x[m]) \log \frac{P(y[m], x[m]; \theta)}{Q_m(x[m])}
 \end{aligned}$$



$$l(\theta, D) \geq \sum_m \sum_{x \in \text{val}(X)} Q_m(x[m]) \log \frac{P(y[m], x[m]; \theta)}{Q_m(x[m])} \quad \text{a lower bound for } l(\theta, D)$$



$l(\theta, D)$ is concave when $P(x, y, \theta)$ is log-concave in θ .

$$\begin{aligned}
 \sum_m \sum_{x \in \text{val}(X)} Q_m(x[m]) \log \frac{P(y[m], x[m]; \theta)}{Q_m(x[m])} \\
 l(\theta, D) = \sum_m \log E_{x \sim Q_m} \left[\frac{P(y[m], x[m]; \theta)}{Q_m(x[m])} \right] = \sum_m E_{x \sim Q_m} \left[\log \frac{P(y[m], x[m]; \theta)}{Q_m(x[m])} \right]
 \end{aligned}$$

$$Q_m(x[m]) = \frac{1}{Z_m} P(y[m], x[m]; \theta)$$

$$\sum_{x \in \text{val}(X)} Q_m(x[m]) = \frac{1}{Z_m} \sum_{x \in \text{val}(X)} P(y[m], x[m]; \theta) = \frac{1}{Z_m} \cdot P(y[m]; \theta) = 1 \quad Z_m = P(y[m]; \theta)$$

$$Q_m(x[m]) = \frac{P(y[m], x[m]; \theta)}{P(y[m]; \theta)} = P(x[m] | y[m]; \theta) \neq x[m] \in \text{val}(x)$$

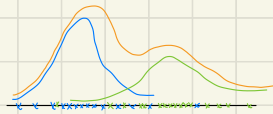
then maximize the lower bound.

$$\theta := \arg \max_{\theta} \sum_m \sum_{x \in \text{val}(X)} Q_m(x[m]) \log \frac{P(y[m], x[m]; \theta)}{Q_m(x[m])}$$

g. mixture of Gaussians

$X \sim \text{Bernoulli}(p_x)$

$(Y|X=j) \sim \mathcal{N}(u_j, \Sigma_j)$



$$\begin{aligned} \ell(\theta) &= \sum_{n=1}^M \log P(y^{(n)}; \theta) = \sum_{n=1}^M \log \sum_{j \in \text{values}} P(X^{(n)} | y^{(n)}; \theta) = \sum_{n=1}^M \log \sum_{j \in \text{values}} \alpha_n(X^{(n)}) \frac{P(X^{(n)}, y^{(n)}; \theta)}{\alpha_n(X^{(n)})} = \sum_{n=1}^M \log \mathbb{E}_{\alpha_n(X^{(n)})} \left[\frac{P(X^{(n)}, y^{(n)}; \theta)}{\alpha_n(X^{(n)})} \right] \\ &> \sum_{n=1}^M \mathbb{E}_{\alpha_n(X^{(n)})} \left[\log \frac{P(X^{(n)}, y^{(n)}; \theta)}{\alpha_n(X^{(n)})} \right] = \sum_{n=1}^M \sum_{j \in \text{values}} \alpha_n(X^{(n)}) \cdot \log \frac{P(X^{(n)}, y^{(n)}; \theta)}{\alpha_n(X^{(n)})} \end{aligned}$$

E-Step:

$$\alpha_n(X^{(n)}) = P(X^{(n)} | y^{(n)}; \theta) = \frac{P(X^{(n)}, y^{(n)}; \theta)}{P(y^{(n)}; \theta)}$$

$$\alpha_n(X^{(n)}) = \frac{1}{2} \left(p_x \cdot \mathcal{N}(y^{(n)}; u_1, \Sigma_1) + \frac{1}{(2\pi)^{d/2} |\Sigma_1|^{1/2}} \exp\left(-\frac{1}{2} (y^{(n)} - u_1)^T \Sigma_1^{-1} (y^{(n)} - u_1)\right) \right)$$

$$\alpha_n(X^{(n)}) = \frac{1}{2} \left((1-p_x) \cdot \mathcal{N}(y^{(n)}; u_0, \Sigma_0) + \frac{1-p_x}{(2\pi)^{d/2} |\Sigma_0|^{1/2}} \exp\left(-\frac{1}{2} (y^{(n)} - u_0)^T \Sigma_0^{-1} (y^{(n)} - u_0)\right) \right)$$

M-Step:

$$\begin{aligned} \max_{\theta} \sum_{n=1}^M \sum_{j \in \text{values}} \alpha_n(X^{(n)}) \cdot \log \frac{P(X^{(n)}, y^{(n)}; \theta)}{\alpha_n(X^{(n)})} \\ = \max_{\theta} \sum_{n=1}^M \sum_{j \in \text{values}} \alpha_n(X^{(n)}) \cdot \left[\log \frac{B_{\theta n}}{(2\pi)^{d/2} |\Sigma_n|^{1/2}} - \frac{1}{2} (y^{(n)} - u_{\theta n})^T \Sigma_n^{-1} (y^{(n)} - u_{\theta n}) \right] \end{aligned}$$

max u :

$$\begin{aligned} \nabla u_j &= \sum_{n=1}^M \alpha_n(j) \cdot \left[-\frac{1}{2} \cdot 2 \Sigma_j^{-1} (y^{(n)} - u_j) \right] \\ &= \Sigma_j^{-1} \left[\sum_{n=1}^M \alpha_n(j) y^{(n)} - \sum_{n=1}^M \alpha_n(j) u_j \right] \Rightarrow \\ u_j &= \frac{\sum_{n=1}^M \alpha_n(j) y^{(n)}}{\sum_{n=1}^M \alpha_n(j)} \end{aligned}$$

max Σ :

$$\begin{aligned} \max_{\Sigma_j} \sum_{n=1}^M \alpha_n(j) \left[-\frac{1}{2} \log |\Sigma_j| - \frac{1}{2} (y^{(n)} - u_j)^T \Sigma_j^{-1} (y^{(n)} - u_j) \right] \\ = \max_{\Sigma_j} \sum_{n=1}^M \alpha_n(j) \left[\log \det(\Sigma_j) - \text{tr}(\Sigma_j^{-1} (y^{(n)} - u_j)(y^{(n)} - u_j)^T) \right] \\ = \max_{\Sigma_j} \log \det(\Sigma_j) \sum_{n=1}^M \alpha_n(j) - \text{tr}(\Sigma_j^{-1} \sum_{n=1}^M \alpha_n(j) (y^{(n)} - u_j)(y^{(n)} - u_j)^T) \end{aligned}$$

$$\nabla \Sigma_j = \Sigma_j^{-1} \sum_{n=1}^M \alpha_n(j) - \sum_{n=1}^M \alpha_n(j) (y^{(n)} - u_j)(y^{(n)} - u_j)^T \Rightarrow$$

$$\Sigma_j = \Sigma_j^{-1} = \frac{\sum_{n=1}^M \alpha_n(j) (y^{(n)} - u_j)(y^{(n)} - u_j)^T}{\sum_{n=1}^M \alpha_n(j)}$$

max θ

$$\max_{\theta} \sum_{j \in \text{values}} \alpha_n(j) \log \frac{B_j \cdot \mathcal{N}(y^{(n)}; u_j, \Sigma_j)}{\alpha_n(j)}$$

$$\max_{\theta} \sum_{j \in \text{values}} \alpha_n(j) \log B_j$$

$$\text{st. } \sum_j B_j = 1 \quad \text{!}$$

$$\log \det(X + \alpha X) = \log \det [X^{1/2} (I + X^{-1/2} \alpha X^{1/2}) X^{1/2}]$$

$$= \log \det [X (I + X^{-1/2} \alpha X^{1/2})]$$

$$= \log \det(X) + \log \det(I + X^{-1/2} \alpha X^{1/2})$$

$$= \log \det(X) + \sum_i \log(1 + \lambda_i) \quad \lambda_i \text{ is the } i\text{th eigenvalue of } X^{-1/2} \alpha X^{1/2}$$

$$\approx \log \det(X) + \sum_i \lambda_i \quad \lim_{\lambda \rightarrow 0} \frac{\lambda}{\log(1+\lambda)} = 1$$

$$= \log \det(X) + \text{tr}(X^{-1/2} \alpha X^{1/2})$$

$$= \log \det(X) + \text{tr}(\alpha X, X^{-1})$$

$$\mathcal{L}(\theta, V) = -\sum_{m=1}^M \sum_j \log(\Omega_m(j)) \log \theta_j + V \left(\sum_j \theta_j - 1 \right)$$

$$\nabla_{\theta} \mathcal{L} = -\frac{1}{\theta_j} \sum_{m=1}^M \Omega_m(j) + V = 0$$

$$\theta_j = \frac{\sum_{m=1}^M \Omega_m(j)}{V}$$

$$g(V) = \inf_{\theta} \mathcal{L}(\theta, V) = -\sum_{m=1}^M \sum_j \log(\Omega_m(j)) \log \frac{\sum_{m=1}^M \Omega_m(j)}{V} + \sum_j \sum_{m=1}^M \Omega_m(j) - V$$

$$= \sum_{m=1}^M \sum_j \log(\Omega_m(j)) \log V - V + C$$

$$\nabla_V g = \frac{1}{V} \sum_{m=1}^M \sum_j \log(\Omega_m(j)) - 1 = 0$$

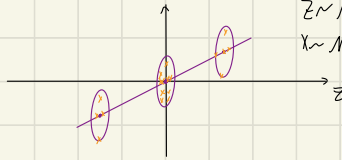
$$V^* = \sum_{m=1}^M \sum_j \Omega_m(j) = \sum_{m=1}^M 1 = M$$

$$\theta_j^* = \frac{\sum_{m=1}^M \Omega_m(j)}{M}$$

eg. Factor Analysis

$$z \sim \mathcal{N}(0, 1)$$

$$x|z \sim \mathcal{N}(\mu + \lambda z, \Psi)$$



$$z \sim \mathcal{N}(0, 1)$$

$$x \sim \mathcal{N}(\begin{bmatrix} \mu \\ \lambda \end{bmatrix} z + \begin{bmatrix} \mu \\ \lambda \end{bmatrix}; \begin{bmatrix} \Psi & 0 \\ 0 & \Sigma \end{bmatrix})$$

$$\mathcal{L}(\theta) = \sum_{m=1}^M \log p(x[m], z[m]; \theta)$$

$$= \sum_{m=1}^M \log \int_{-\infty}^{+\infty} p(x[m], z[m]; \theta) dz[m]$$

$$= \sum_{m=1}^M \log \int_{-\infty}^{+\infty} \Omega_m(z[m]) \frac{p(x[m], z[m]; \theta)}{\Omega_m(z[m])} dz[m] = \sum_{m=1}^M \log E_{z[m] \sim \Omega_m} \left[\frac{p(x[m], z[m]; \theta)}{\Omega_m(z[m])} \right]$$

$$\geq \sum_{m=1}^M E_{z[m] \sim \Omega_m} \left[\log \frac{p(x[m], z[m]; \theta)}{\Omega_m(z[m])} \right]$$

$$= \sum_{m=1}^M \int_{-\infty}^{+\infty} \Omega_m(z[m]) \cdot \log \frac{p(x[m], z[m]; \theta)}{\Omega_m(z[m])} dz[m]$$

$$\Omega_m(z[m]) = \frac{1}{\Sigma} p(x[m], z[m]; \theta) \quad \forall z[m] \in \text{val}(z)$$

$$\int_{-\infty}^{+\infty} \Omega_m(z[m]) dz[m] = \frac{1}{\Sigma} \int_{-\infty}^{+\infty} p(x[m], z[m]; \theta) dz[m] = \frac{1}{\Sigma} \cdot p(x[m]; \theta) = 1$$

$$\Omega_m(z[m]) = \frac{p(x[m], z[m]; \theta)}{p(x[m]; \theta)} = p(z[m] | x[m]; \theta)$$

E-Step:

$$p(z[m], x[m]; \theta) = p(z[m]) \cdot p(x[m] | z[m]; \theta)$$

$$= \frac{1}{(2\pi)^{\frac{M}{2}}} \exp\left(-\frac{1}{2} z^T z\right) \cdot \frac{1}{(2\pi)^{\frac{M}{2}} |\Psi|^{\frac{M}{2}}} \exp\left(-\frac{1}{2} [x[m] - (\mu + \lambda z)]^T \Psi^{-1} [x[m] - (\mu + \lambda z)]\right)$$

$$= \frac{1}{2} \exp\left\{-\frac{1}{2} [z^T z + x^T \Psi^{-1} x - 2x^T \Psi^{-1} (\mu + \lambda z) + (\mu + \lambda z)^T \Psi^{-1} (\mu + \lambda z)]\right\}$$

$$= \frac{1}{2} \exp\left\{-\frac{1}{2} [z^T z + x^T \Psi^{-1} x - 2x^T \Psi^{-1} \mu - 2x^T \Psi^{-1} \lambda z + \mu^T \Psi^{-1} \mu + 2\mu^T \Psi^{-1} \lambda z + \lambda^T \Psi^{-1} \lambda z]\right\}$$

$$= \frac{1}{2} \exp\left\{-\frac{1}{2} [z^T z, x^T \mu] \begin{bmatrix} I + \lambda^T \Psi^{-1} \lambda & -\lambda^T \Psi^{-1} \\ -\Psi^{-1} \lambda & \Psi^{-1} \end{bmatrix} \begin{bmatrix} z \\ x \mu \end{bmatrix}\right\}$$

$$\begin{bmatrix} z \\ x \mu \end{bmatrix} \sim \mathcal{N}\left(\begin{bmatrix} 0 \\ \mu \end{bmatrix}, \Sigma\right)$$

$$\Sigma = \begin{bmatrix} I + \lambda^T \Psi^{-1} \lambda & -\lambda^T \Psi^{-1} \\ -\Psi^{-1} \lambda & \Psi^{-1} \end{bmatrix} = \begin{bmatrix} I & \lambda^T \\ \lambda & \Psi + \lambda \lambda^T \end{bmatrix}$$

$$\begin{bmatrix} z_1 & z_2 \\ x_1 & x_2 \end{bmatrix}^{-1} = \begin{bmatrix} (z_1 - z_2 z_2^{-1} z_1^{-1})^{-1} & -(z_1 - z_2 z_2^{-1} z_1^{-1})^{-1} z_2 z_2^{-1} \\ -z_2^{-1} z_1 (z_1 - z_2 z_2^{-1} z_1^{-1})^{-1} & z_2^{-1} + z_2^{-1} z_1 (z_1 - z_2 z_2^{-1} z_1^{-1})^{-1} z_2 z_2^{-1} \end{bmatrix}$$

by block elimination

$$(z[m] | \lambda[m]) \sim \mathcal{N}(-\lambda^T (\Psi + \lambda \lambda^T)^{-1} (\lambda[m] - u), \tau^{-1} \lambda^T (\Psi + \lambda \lambda^T)^{-1} \lambda)$$

$$Q_m(z[m]) = P(z[m] | \lambda[m]) = \mathcal{N}(z[m]; \underbrace{-\lambda^T (\Psi + \lambda \lambda^T)^{-1} (\lambda[m] - u)}_{E[z[m] | \lambda[m]]}, \underbrace{\tau^{-1} \lambda^T (\Psi + \lambda \lambda^T)^{-1} \lambda}_{\text{var}(z[m] | \lambda[m])})$$

M-Step:

$$\begin{aligned} & \max_{\lambda, u, \Psi} \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \log \frac{P(\lambda[m], z[m]; \theta)}{Q_m(z[m])} dz[m] \\ &= \max_{\lambda, u, \Psi} \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) [\log P(z[m]) + \log P(\lambda[m] | z[m]; \theta) - \log Q_m(z[m])] dz[m] \\ &= \max_{\lambda, u, \Psi} \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \log P(\lambda[m] | z[m]; \theta) dz[m] \\ &= \max_{\lambda, u, \Psi} \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \left[\log \frac{1}{(2\pi)^2 |\Phi|} - \frac{1}{2} (\lambda[m] - (u + \lambda z[m]))^T \Phi^{-1} (\lambda[m] - (u + \lambda z[m])) \right] dz[m] \end{aligned}$$

max λ :

$$\begin{aligned} \nabla_{\lambda} &= \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \Psi^{-1} (\lambda[m] - u - \lambda z[m]) z[m]^T dz[m] \\ &= \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \Psi^{-1} (\lambda[m] - u) z[m]^T dz[m] - \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \Psi^{-1} \lambda z[m] z[m]^T dz[m] \\ &= \sum_{m=1}^M \Psi^{-1} (\lambda[m] - u) E_{z[m] \sim Q_m} [z[m]]^T - \sum_{m=1}^M \Psi^{-1} \lambda E_{z[m] \sim Q_m} [z[m] z[m]^T] \stackrel{!}{=} 0 \\ \lambda &= \left(\sum_{m=1}^M (\lambda[m] - u) E_{z[m] \sim Q_m} [z[m]]^T \right) \left(\sum_{m=1}^M E_{z[m] \sim Q_m} [z[m] z[m]^T] \right)^{-1} \end{aligned}$$

max u :

$$\begin{aligned} \nabla_u &= \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \Psi^{-1} (\lambda[m] - u - \lambda z[m]) dz[m] \\ &= \sum_{m=1}^M \Psi^{-1} \int_{z[m]}^{p+\tau} Q_m(z[m]) (\lambda[m] - u) dz[m] - \sum_{m=1}^M \Psi^{-1} \lambda \int_{z[m]}^{p+\tau} Q_m(z[m]) z[m] dz[m] \\ &= \Psi^{-1} \sum_{m=1}^M (\lambda[m] - u) - \Psi^{-1} \lambda \sum_{m=1}^M E_{z[m] \sim Q_m} [z[m]] \\ &= \Psi^{-1} \sum_{m=1}^M (\lambda[m] - u) - \Psi^{-1} \lambda \cdot \sum_{m=1}^M -\lambda^T (\Psi + \lambda \lambda^T)^{-1} (\lambda[m] - u) \\ &= (\Psi^{-1} + \Psi^{-1} \lambda \lambda^T (\Psi + \lambda \lambda^T)^{-1}) \sum_{m=1}^M (\lambda[m] - u) \stackrel{!}{=} 0 \end{aligned}$$

$$u = \frac{1}{M} \sum_{m=1}^M \lambda[m]$$

max Ψ :

$$\begin{aligned} & \max_{\Psi} \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \cdot \left\{ -\frac{1}{2} \log \det(\Psi) - \frac{1}{2} (\lambda[m] - (u + \lambda z[m]))^T \Psi^{-1} (\lambda[m] - (u + \lambda z[m])) \right\} dz[m] \\ &= \min_{\Psi} \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \cdot \left\{ \log \det(\Psi) + \text{tr} \left\langle \Psi^{-1}, (\lambda[m] - (u + \lambda z[m])) (\lambda[m] - (u + \lambda z[m]))^T \right\rangle \right\} dz[m] \\ \text{let } J &= \Psi^{-1} \\ &= \min_J \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \cdot \left\{ -\log \det(J) + \text{tr} \left\langle J, (\lambda[m] - (u + \lambda z[m])) (\lambda[m] - (u + \lambda z[m]))^T \right\rangle \right\} dz[m] \\ \nabla_J &= \sum_{m=1}^M \int_{z[m]}^{p+\tau} Q_m(z[m]) \cdot \left\{ -J^{-1} + (\lambda[m] - u - \lambda z[m]) (\lambda[m] - u - \lambda z[m])^T \right\} dz[m] \end{aligned}$$

$$= \sum_{m=1}^M -J^T + \sum_{m=1}^M \int_{z[m]}^{+\infty} Q_m(z[m]) \lambda z[m] z[m]^T \rightarrow \lambda z[m] (\lambda z[m] - u)^T + (\lambda z[m] - u)(\lambda z[m] - u)^T d(z[m])$$

$$= -MJ^T + \sum_{m=1}^M \lambda E_{z[m] \sim Q_m} [z[m] z[m]^T] \lambda^T \rightarrow \lambda E_{z[m] \sim Q_m} [z[m]] (\lambda z[m] - u)^T + (\lambda z[m] - u)(\lambda z[m] - u)^T \quad \text{if } \lambda \rightarrow 0$$

$$J = J^T = \frac{1}{M} \sum_{m=1}^M \lambda E_{z[m] \sim Q_m} [z[m] z[m]^T] \lambda^T \rightarrow \lambda E_{z[m] \sim Q_m} [z[m]] (\lambda z[m] - u)^T + (\lambda z[m] - u)(\lambda z[m] - u)^T$$

EM algorithm for Bayesian Network

let $o[m]$ be observed vars of mth data and $z[m]$ be hidden vars of mth data

$$\begin{aligned} \ell(\theta) &= \sum_{m=1}^M \log P(o[m]; \theta) \\ &= \sum_{m=1}^M \log \sum_{z[m] \in \mathcal{Z}(o)} P(o[m], z[m]; \theta) \\ &= \sum_{m=1}^M \log \sum_{z[m]} Q_m(z[m]) \cdot \frac{P(o[m], z[m]; \theta)}{Q_m(z[m])} \\ \textcircled{1} &\geq \sum_{m=1}^M \sum_{z[m]} Q_m(z[m]) \log \frac{P(o[m], z[m]; \theta)}{Q_m(z[m])} \\ &= \sum_{m=1}^M \sum_{z[m]} Q_m(z[m]) \cdot \left[\sum_{(o, z) \in \mathcal{O} \times \mathcal{Z}} \log P(o, z) \mathbb{1}_{(o, z) < x} - \log(Q_m(z)) \right] \end{aligned}$$

E-Step:

to achieve equality at $\textcircled{1}$ $\frac{P(o[m], z[m]; \theta)}{Q_m(z[m])}$ is constant

$$Q_m(z[m]) = \frac{1}{Z} P(o[m], z[m]; \theta)$$

$$\sum_{z[m]} Q_m(z[m]) = \frac{1}{Z} \sum_{z[m]} P(o[m], z[m]; \theta) = \frac{1}{Z} P(o[m]; \theta) = 1$$

$$Q_m(z[m]) = \frac{P(o[m], z[m]; \theta)}{P(o[m]; \theta)} = P(z[m] | o[m]; \theta)$$

M-Step:

Find $\theta_{\text{new}} \forall x \in \mathcal{V}(X)$

$$\text{min.} \quad -\sum_{m=1}^M \sum_{z[m] \in \mathcal{Z}(o)} Q_m(z[m]) \log P(o[m], z[m] < x | o[m], z[m] < u)$$

$$\text{s.t.} \quad \sum_{x \in \mathcal{V}(X)} \theta_{\text{new}} = 1$$

$$\begin{aligned} \mathcal{L}(\theta_{\text{new}}, V) &= -\sum_{m=1}^M \sum_{z[m]} Q_m(z[m]) \log P(o[m], z[m] < x | o[m], z[m] < u) + V \left(\sum_{x \in \mathcal{V}(X)} \theta_{\text{new}} - 1 \right) \\ &= -\sum_{x \in \mathcal{V}(X)} \log \theta_{\text{new}} \sum_{m=1}^M \sum_{z[m] \in \mathcal{Z}(o)} Q_m(z[m]) \cdot \mathbb{1}\{o[m], z[m] < x, u\} + V \left(\sum_{x \in \mathcal{V}(X)} \theta_{\text{new}} - 1 \right) \end{aligned}$$

$$\nabla_{\theta_{\text{new}}} \mathcal{L} = -\frac{1}{\theta_{\text{new}}} \cdot \sum_{m=1}^M \sum_{z[m] \in \mathcal{Z}(o)} Q_m(z[m]) \mathbb{1}\{o[m], z[m] < x, u\} + V \quad \text{if } \theta_{\text{new}} = 0$$

$$\theta_{\text{new}}(V) = \frac{1}{V} \cdot \sum_{m=1}^M \sum_{z[m] \in \mathcal{Z}(o)} Q_m(z[m]) \mathbb{1}\{o[m], z[m] < x, u\}$$

$$\text{let } h(x, u) = \sum_{m=1}^M \sum_{z[m]} Q_m(z[m]) \mathbb{1}\{o[m], z[m] < x, u\} \quad \text{is constant } \forall x, u$$

$$b_{xu}(v) = \frac{1}{v} h(xu)$$

$$g(u) = \inf_{b_{xu}} L(b_{xu}, v)$$

$$= - \sum_x \log\left(\frac{1}{v} h(xu)\right) \cdot h(xu) + v \left(\sum_x \frac{1}{v} h(xu) - 1 \right)$$

$$= - \sum_x (\log h(xu) - \log v) \cdot h(xu) + \sum_x h(xu) - v$$

$$\nabla_v g = \frac{1}{v} \cdot \sum_x h(xu) - 1 \stackrel{!}{=} 0$$

$$v^* = \sum_x h(xu)$$

$$\begin{aligned} b_{xu}^* &= \frac{h(xu)}{\sum_x h(xu)} = \frac{\sum_{m=1}^M \sum_{z(m)} \alpha_m(z(m)) \cdot \mathbb{1}\{(0(m), z(m)) \downarrow x, u\} = (xu)\}}{\sum_x \sum_{m=1}^M \sum_{z(m)} \alpha_m(z(m)) \cdot \mathbb{1}\{(0(m), z(m)) \downarrow x, u\} = (xu)\}} \\ &= \frac{\sum_{m=1}^M \sum_{z(m)} p(z(m)|o(m); \theta) \cdot \mathbb{1}\{(0(m), z(m)) \downarrow x, u\} = (xu)\}}{\sum_x \sum_{m=1}^M \sum_{z(m)} p(z(m)|o(m); \theta) \cdot \mathbb{1}\{(0(m), z(m)) \downarrow x, u\} = (xu)\}} \\ &= \frac{\sum_{m=1}^M p(xu|o(m); \theta)}{\sum_x \sum_{m=1}^M p(xu|o(m); \theta)} \end{aligned}$$

def compute_marginal_Q(G, \theta, D: partially observed dataset):

for each \textcircled{D} \rightarrow \textcircled{X} \in G

for each (xu) \in \text{val}(x, u):

$$Q(xu) = 0$$

for m=1..M:

run inference with evidence o(m)

for each \textcircled{D} \rightarrow \textcircled{X} \in G

for each (xu) \in \text{val}(x, u):

$$Q(xu) += p(xu|o(m))$$

return \sum Q(xu) \forall \textcircled{D} \rightarrow \textcircled{X} \in G, \forall (xu) \in \text{val}(x, u)

def EM-Bayesian_Net(G, \theta^0, D):

for t=0,1,2,... until convergence

$$\{Q(xu)\} = \text{compute_marginal_Q}(G, \theta^t, D)$$

for each \textcircled{D} \rightarrow \textcircled{X} \in G

for each (xu) \in \text{val}(x, u):

$$\theta_{xu} = \frac{Q(xu)}{\sum_x Q(xu)}$$

return \theta^t