

Part 1

Stochastic simulation

Part 2

Simulation $\{ \text{MCMC} \}$ Bayes, aims to compute $\pi(\theta|y)$
Laplace approx. $\{ \text{INLA}, \text{RTMB} \}$ Frequentist, aim is MLE $\hat{\theta} = \underset{\theta}{\operatorname{argmax}} \pi(y|\theta)$

Part 3

Estimate sampling distribution of $\hat{\theta}$ given θ (incl. bias and SD)

using a non-parametric or parametric model

(bootstrapping). Frequentist

E-M alg. (after easter)

Non-parametric bootstrap

- Simplest case: iid sample $x_1, x_2, \dots, x_n \sim F$
where F is an unknown cdf.

Find estimator of

$$\theta = t(F)$$

e.g.

$$\theta = E(X) = \underbrace{\int_{-\infty}^{\infty} x \, dF(x)}_{\text{Lebesgue-Stieltjes integral}} = \begin{cases} \int_{-\infty}^{\infty} x f(x) dx & \text{if } X \text{ is cont.} \\ \sum_x x f(x) & \text{if } X \text{ is discr.} \end{cases}$$

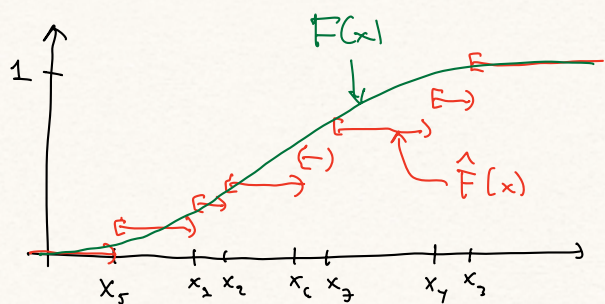
or

$$\theta = \operatorname{Var}(X) = \int (x - EX)^2 dF$$

- F can be estimated non-parametrically by the estimator

$$\hat{F}(\underline{x}) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}(\underline{x}_i \leq \underline{x}),$$

the empirical cdf. Note that x can be a vector



- The corresponding (plug-in) estimator of θ is

$$\hat{\theta} = t(\hat{F})$$

Thus, if $\theta = EX = \int x dF$

$$\hat{\theta} = \int_{-\infty}^{\infty} x d\hat{F} = \sum_x x \hat{f}(x) = \sum_{i=1}^n x_i \frac{1}{n} = \bar{x} = E_{\hat{F}}(X)$$

and if $\theta = \text{Var} X = \int (x - EX)^2 dF$

$$\hat{\theta} = \widehat{\text{Var}(X)} = \sum_x (x - EX)^2 \hat{f}(x)$$

$$= \sum_{i=1}^n (x_i - \bar{x})^2 \frac{1}{n}$$

$$= \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = \text{Var}_{\hat{F}}(X)$$

If $\theta = SD(X)$,

$$\hat{\theta} = \widehat{SD(X)} = SD_{\hat{F}}(X)$$

If $\theta = \text{Skew}(X) = E\left(\left(\frac{X-\mu}{\sigma}\right)^3\right)$ the plug-in estimator of θ is

$$\hat{\theta} = \text{Skew}_{\hat{F}}(X)$$

What is the sampling distribution of these plug-in estimators if we replace F by $\hat{F}$?

- Approximate method

1. Generate B bootstrap samples

$$X^{1*}, X^{2*}, \dots, X^{B*} \text{ where } X^{b*} = (X_1^{b*}, X_2^{b*}, \dots, X_n^{b*})$$

and $X_i^{b*} \stackrel{iid}{\sim} \hat{F}$ for $i=1, 2, \dots, n$, $b=1, 2, \dots, B$.

2. Compute corresponding bootstrap replicates of plug-in estimator $\hat{\theta}$ of θ ,

$$\hat{\theta}^{b*} = t(\hat{F}^{b*})$$

3. Estimate e.g. $E_{\hat{F}}(\hat{\theta})$ and $SD_{\hat{F}}(\hat{\theta})$ or, in general $E(h(\hat{\theta}))$, by Monte-Carlo integration, i.e.

$$\widehat{E(\hat{\theta})} = \frac{1}{B} \sum_{b=1}^B \hat{\theta}^{b*}$$

and

$$\widehat{SD(\hat{\theta})} = \sqrt{\frac{1}{B-1} \sum_{b=1}^B (\hat{\theta}^{b*} - \widehat{E(\hat{\theta})})^2}$$

- Exact method (applicable for small n)

If the original sample $x_1, x_2, \dots, x_n$ consist of distinct values the number of unordered outcomes $((x_4, x_4, x_1, x_2) \text{ the same as } (x_1, x_2, x_4, x_4))$ is

$$\binom{2n-1}{n-1} \underset{n=10}{\approx} \underset{= 92378}{(n\pi)^{-\frac{1}{2}} 2^{2n-1}}$$

and the probability of observing $m_1, m_2, \dots, m_n$ of each value $x_1, x_2, \dots, x_n$ is

$$\frac{n!}{m_1! m_2! m_3!} \cdot \underbrace{\left(\frac{1}{n}\right)^{m_1} \left(\frac{1}{n}\right)^{m_2} \dots \left(\frac{1}{n}\right)^{m_n}}_{n^{-n}}$$

since $m_1, \dots, m_n \sim \text{multinomial}(n, (\frac{1}{n}, \frac{1}{n}, \dots, \frac{1}{n}))$

Ex: $x = (5, 7, 10)$ $n = 3$

#5's	#7's	#10's	Ball-bar repr.	Prob
3	0	0	0 0 0	3^{-3}
2	1	0	0 0 0	3^{-2}
2	0	1	0 0 1 0	3^{-2}
1	2	0	0 0 0	$\vdots$
1	1	1	0 0 1 0	$\vdots$
1	0	2	0 1 0 0	
0	3	0	1 0 0 0	
0	2	1	1 0 0 1 0	
0	1	2	1 0 1 0 0	
0	0	3	1 1 0 0 0	

n balls

$n-1$ bars

has

$$\binom{2n-1}{n-1} \text{ combinations}$$

$$= \binom{5}{2} = 10$$

Original sample with $k \leq n$ distinct values can be treated similarly,

$$\binom{n+k-1}{k-1} \text{ distinct outcomes}$$

Parametric bootstrapping

Suppose original data $\underline{x} = (x_1, x_2, \dots, x_n) \sim F(\underline{x}; \underline{\theta})$ (no iid assumption)
where F is known and $\underline{\theta}$ is unknown.

Estimate $\underline{\theta}$ by $\hat{\underline{\theta}}$ (by given method, e.g. ML) and
 $F(\underline{x}; \underline{\theta})$ by $F(\underline{x}; \hat{\underline{\theta}})$. Algorithm

1. Simulate $\underline{x}_1^{1*}, \underline{x}_2^{2*}, \dots, \underline{x}^{B*} \stackrel{\text{iid}}{\sim} F(\underline{x}; \hat{\underline{\theta}})$
2. Compute corresponding bootstrap replicates of $\hat{\underline{\theta}}$, $\hat{\underline{\theta}}^{b*}$, $b = 1, 2, \dots, B$ (by same given method)
3. Estimate $E(h(\hat{\underline{\theta}}))$ by Monte-Carlo integration

Estimating and correcting for bias

Having estimated $E(\hat{\underline{\theta}})$ by $\widehat{E(\hat{\underline{\theta}})} = \frac{1}{B} \sum_{b=1}^B \hat{\underline{\theta}}^{b*}$
an estimator of

$$\text{Bias}(\hat{\underline{\theta}}) = E(\hat{\underline{\theta}}) - \underline{\theta}$$

is

$$\widehat{\text{Bias}(\hat{\underline{\theta}})} = \widehat{E(\hat{\underline{\theta}})} - \hat{\underline{\theta}}$$

Subtracting the estimated bias from $\hat{\underline{\theta}}$ we obtain
a bias-corrected estimator

$$\hat{\underline{\theta}}_c = \hat{\underline{\theta}} - \widehat{\text{Bias}(\hat{\underline{\theta}})}$$

Typically has less but not zero bias and
sometimes higher variance.

Ex. August 2024, problem 4

$X_1, \dots, X_n \stackrel{i.i.d.}{\sim} F$, $F(x) = 1 - e^{-\lambda x}$ (exponential distribution)

MLE of λ is $\hat{\lambda} = \frac{n}{\sum_{i=1}^n X_i}$.

Observed data gives $\hat{\lambda} = 2.0$.

Let $\hat{\lambda}^{bx}$ and $\hat{\lambda}^{bx}$, $b = 1, 2, \dots, B$ denote parametric bootstrap samples from $F(\hat{\lambda})$, and suppose

$$\frac{1}{B} \sum \hat{\lambda}^{bx} = 2.24.$$

a) An estimate of $E(\hat{\lambda})$ is then $\widehat{E(\hat{\lambda})} = \frac{1}{B} \sum \hat{\lambda}^{bx} = 2.24$
and an estimate of $\text{Bias}(\hat{\lambda}) = E(\hat{\lambda}) - \lambda$ is

$$\widehat{\text{Bias}(\hat{\lambda})} = \widehat{E(\hat{\lambda})} - \hat{\lambda} = 2.24 - 2 = 0.24.$$

A bias corrected estimate of λ is then

$$\begin{aligned} \hat{\lambda}_c &= \hat{\lambda} - \widehat{\text{Bias}(\hat{\lambda})} = \hat{\lambda} - \left(\frac{1}{B} \sum_{b=1}^B \hat{\lambda}^{bx} - \hat{\lambda} \right) = 2\hat{\lambda} - \frac{1}{B} \sum_{b=1}^B \hat{\lambda}^{bx} \\ &= 2 - 0.24 = 1.76 \end{aligned}$$

b) In this particular example, we know analytically that

$$E(\hat{\lambda}) = \frac{n}{n-1} \lambda \quad \stackrel{n=10}{=} \quad \frac{10}{9} \lambda = 1.11 \lambda$$

→ this implies that

$$E(\hat{\lambda}^{b*} | \hat{\lambda}) = \frac{n}{n-1} \hat{\lambda}$$

Thus,

$$\begin{aligned} E(\hat{\lambda}_c | \hat{\lambda}) &= E\left(2\hat{\lambda} - \frac{1}{B} \sum \hat{\lambda}^{b*} \mid \hat{\lambda}\right) \\ &= 2\hat{\lambda} - \frac{1}{B} \sum_{b=1}^B E(\hat{\lambda}^{b*} | \hat{\lambda}) \\ &= 2\hat{\lambda} - \frac{n}{n-1} \hat{\lambda} = \frac{2n-2-n}{n-1} \hat{\lambda} = \frac{n-2}{n-1} \hat{\lambda} \end{aligned}$$

and

$$\begin{aligned} E(\hat{\lambda}_c) &= E E(\hat{\lambda}_c | \hat{\lambda}) = E\left(\frac{n-2}{n-1} \hat{\lambda}\right) = \frac{n-2}{n-1} E(\hat{\lambda}) \\ &= \frac{n-2}{n-1} \cdot \frac{n}{n-1} \lambda = \frac{n(n-2)}{(n-1)^2} \lambda \quad \stackrel{n=10}{=} \quad \frac{80}{81} \lambda = 0.99 \lambda \end{aligned}$$

i.e. $\hat{\lambda}_c$ is only less biased than $\hat{\lambda}$.

Similarly, we can find $\text{Var}(\hat{\lambda}_c)$ (via law of total variance)
but this will depend on B (why?)

$$\text{Var}(\hat{\lambda}_c) = E \text{Var}(\hat{\lambda}_c | \hat{\lambda}) + \text{Var} E(\hat{\lambda}_c | \hat{\lambda})$$

Bootstrap confidence intervals

→ Percentile method

1. Generate $x^{1*}, x^{2*}, \dots, x^{B*} \stackrel{iid}{\sim} \hat{F}$

2. Compute $\hat{\theta}^{1*}, \hat{\theta}^{2*}, \dots, \hat{\theta}^{B*}$

3. Compute empirical $\alpha/2$ and $1-\alpha/2$ quantiles of $\hat{\theta}^{b*}$

Justification: Valid if a strictly increasing transformation $\phi(\theta)$ exist

such that $\phi(\hat{\theta}) - \phi(\theta)$ has cdf $H(z) = 1 - H(-z)$.

Then

$$P(h_{\alpha/2} \leq \phi(\hat{\theta}) - \phi(\theta) \leq h_{1-\alpha/2}) = 1 - \alpha \quad (1)$$

where $h_{\alpha} = H^{-1}(\alpha)$.

Given an existing ϕ bootstrap replicates of $\phi(\hat{\theta}) - \phi(\theta)$ are given by $\phi(\hat{\theta}^*) - \phi(\hat{\theta})$ from which $h_{\alpha/2}$ and $h_{1-\alpha/2}$ can be estimated by empirical quantiles, i.e.

$$P(h_{\alpha/2} \leq \phi(\hat{\theta}^*) - \phi(\hat{\theta}) \leq h_{1-\alpha/2}) \approx 1 - \alpha$$

and so

$$P\left(\underbrace{\phi^{-1}(h_{\alpha/2} + \phi(\hat{\theta}))}_{\hat{\theta}_{\alpha/2}^*} \leq \hat{\theta}^* \leq \underbrace{\phi^{-1}(h_{1-\alpha/2} + \phi(\hat{\theta}))}_{\hat{\theta}_{1-\alpha/2}^*}\right) \approx 1 - \alpha$$

The empirical quantiles $\hat{\theta}_{\alpha/2}^*$ and $\hat{\theta}_{1-\alpha/2}^*$ of $\hat{\theta}^*$ are explicitly known.

From (1) we have

$$P(-h_{\alpha/2} \geq \phi(\theta) - \phi(\hat{\theta}) \geq -h_{1-\alpha/2}) = 1 - \alpha$$

Since $H(z) = 1 - H(-z)$, $h_{\alpha/2} = -h_{1-\alpha/2}$ and so

$$P(h_{\alpha/2} \leq \phi(\theta) - \phi(\hat{\theta}) \leq h_{1-\alpha/2}) = 1 - \alpha$$

and

$$P\left(\underbrace{\phi^{-1}(\phi(\theta) - h_{\alpha/2})}_{=\hat{\theta}_{\alpha/2}^*} \leq \theta \leq \underbrace{\phi^{-1}(\phi(\hat{\theta}) - h_{1-\alpha/2})}_{=\hat{\theta}_{1-\alpha/2}^*}\right) = 1 - \alpha$$

Thus, $(\hat{\theta}_{\alpha/2}^*, \hat{\theta}_{1-\alpha/2}^*)$ is a conf. int. for θ . Note that ϕ and $h_{\alpha/2}, h_{1-\alpha/2}$ do not need to be known explicitly!

Ex. of failure of percentile method:

Suppose $X_1, \dots, X_n \stackrel{iid}{\sim} N(\mu, \sigma^2)$.

It follows that parametric bootstrap replicates of $\frac{\hat{\sigma}^2(n-1)}{\sigma^2}$, i.e. $\frac{\hat{\sigma}^{2*}(n-1)}{\hat{\sigma}^2} \sim \chi_{n-1}^2$ and bootstrap replicates

of $\hat{\sigma}^2$, i.e. $\hat{\sigma}^{2*} \sim \frac{\hat{\sigma}^2}{n-1} \chi_{n-1}^2$ so the percentil-method interval is (up to Monte-Carlo error)

$$\left(\frac{\hat{\sigma}^2 \chi_{\alpha/2, n-1}^2}{n-1}, \frac{\hat{\sigma}^2 \chi_{1-\alpha/2, n-1}^2}{n-1} \right)$$

But the classical exact interval is

$$\left(\frac{\hat{\sigma}^2(n-1)}{\chi_{1-\alpha/2, n-1}^2}, \frac{\hat{\sigma}^2(n-1)}{\chi_{\alpha/2, n-1}^2} \right)$$

Bootstrapping regression

12.2

Model

$$y_i = x_i^T \beta + \varepsilon_i$$

where the errors ε_i are iid zero mean constant variance.

In matrix notation

$$\underline{y} = X \underline{\beta} + \underline{\varepsilon}$$

Least-squares estimate of $\underline{\beta}$ is

$$\hat{\underline{\beta}} = \underset{\underline{\beta}}{\operatorname{argmin}} \left\| \underline{y} - X \underline{\beta} \right\|_2^2 = (X^T X)^{-1} X^T \underline{y}$$

Residuals defined as

$$\hat{\varepsilon}_i = y_i - x_i^T \hat{\underline{\beta}}$$

- Bootstrapping residuals

1. Generate B bootstrap replicates

$\hat{\varepsilon}_1^*, \hat{\varepsilon}_2^*, \dots, \hat{\varepsilon}_n^*$ by resampling with replacement from $\hat{\varepsilon}_1, \dots, \hat{\varepsilon}_n$.

2. Compute $y_i^* = x_i^T \hat{\underline{\beta}} + \hat{\varepsilon}_i^*$

3. Compute bootstrap replicates $\hat{\underline{\beta}}^*$ by regressing each $\underline{y}^*$ on X

More appropriate for experimental data

R demo

— Bootstrapping pairs

1. Generate B bootstrap samples $(y_i^*, x_i^*), i=1, 2, \dots, n$ by sampling $(y_i, x_i), i=1, \dots, n$ with replacement.
2. Compute bootstrap replicates $\hat{\beta}^*$ by regressing each y_i^* on each x_i^*

More appropriate for observational data

- Accelerated Bias-corrected Percentile Method BC_a

Assume strictly increasing $\phi(\cdot)$ and constants a and b exist such that

$$U = \frac{\phi(\hat{\theta}) - \phi(\theta)}{1 + a\phi(\theta)} + b \sim N(0, 1) \quad (1)$$

Bootstrap replicates of U from $\hat{F}$,

$$U^* = \frac{\phi(\hat{\theta}^*) - \phi(\hat{\theta})}{1 + a\phi(\hat{\theta})} + b \sim N(0, 1) \quad (2)$$

Thus, from (2),

$$\beta = P(U^* \leq z_\beta)$$

$$= P\left(\frac{\phi(\hat{\theta}^*) - \phi(\hat{\theta})}{1 + a\phi(\hat{\theta})} + b \leq z_\beta \right)$$

$$= P\left(\hat{\theta}^* \leq \phi^{-1}\left(\phi(\hat{\theta}) + (z_\beta - b)[1 + a\phi(\hat{\theta})] \right) \right) \quad (3)$$

$= \hat{\theta}_\beta^* = \text{observable empirical } \beta\text{-quantile of } \hat{\theta}^*$

Similarly, from (1), isolating θ rather than $\hat{\theta}$

$$1 - \alpha = P(U > z_\alpha)$$

$$= P\left(\frac{\phi(\hat{\theta}) - \phi(\theta)}{1 + a\phi(\theta)} + b > z_\alpha \right)$$

$$= P\left(\phi(\hat{\theta}) - \phi(\theta) > (z_\alpha - b)(1 + a\phi(\theta)) \right)$$

$$= P\left(\phi(\hat{\theta}) + b - z_\alpha > \left[1 - a(b - z_\alpha) \right] \phi(\theta) \right)$$

$$= P\left(\theta < \Phi^{-1}\left(\frac{\Phi(\hat{\theta}) + b - z_\alpha}{1 - a(b - z_\alpha)}\right)\right)$$

$$= P\left(\theta < \underbrace{\Phi^{-1}\left(\Phi(\hat{\theta}) + \frac{b - z_\alpha}{1 - a(b - z_\alpha)}\right)}_{*} \left(1 + a\Phi(\hat{\theta})\right)\right)$$

$$= \Phi(\hat{\theta}) + \frac{-\Phi(\hat{\theta})(1 - a(b - z_\alpha)) + \Phi(\hat{\theta}) + b - z_\alpha}{1 - a(b - z_\alpha)} = \Phi(\hat{\theta}) + \frac{a\Phi(\hat{\theta})(b - z_\alpha) + b - z_\alpha}{1 - a(b - z_\alpha)} =$$

i.e., (*) is the upper limit of a $(1 - \alpha)$ one sided conf. int. for θ , equal to $\hat{\theta}_\beta^*$ for

$$\frac{b - z_\alpha}{1 - a(b - z_\alpha)} = z_\beta - b$$

$$\beta = \Phi\left(b + \frac{b - z_\alpha}{1 - a(b - z_\alpha)}\right)$$

Two-sided confidence limits given by $(\hat{\theta}_{\beta_1}^*, \hat{\theta}_{\beta_2}^*)$ with α replaced by $\alpha/2$ and $1 - \alpha/2$, respectively.

Choosing a and b (Shao & Tu, 1996):

$$b = \Phi^{-1}\left(F_{\hat{\theta}^*}(\hat{\theta})\right), \text{ where } F_{\hat{\theta}^*} \text{ is the cdf of } \hat{\theta}^*$$

$$\alpha = \frac{\frac{1}{6} \sum_{i=1}^n \psi_i^3}{\left(\sum_{i=1}^n \psi_i^2\right)^{\frac{2}{3}}}, \text{ where } \psi_i = \hat{\theta}_{(i)} - \hat{\theta}_{(-i)},$$

$\hat{\theta}_{(-i)}$ is estimate based on omitting i th observation and

$$\hat{\theta}_{(i)} = \frac{1}{n} \sum_{j=1}^n \theta_{(-j)}$$

R demo

Bootstrap & confidence interval

Suppose an estimate $\hat{V}$ of the variance of $\hat{\theta}$ is "directly" available.

Ex.: Regression, $\hat{\theta} = f(\hat{\beta}_0, \hat{\beta}_1) = \frac{\hat{\beta}_1}{\hat{\beta}_0}$

$$\hat{V} = \widehat{\text{Var}}(\hat{\theta}) \approx \left(\frac{\partial f}{\partial \beta_0}\right)^2 \widehat{\text{Var}}(\hat{\beta}_0) + 2\left(\frac{\partial f}{\partial \beta_0}\right)\left(\frac{\partial f}{\partial \beta_1}\right) \widehat{\text{Cov}}(\hat{\beta}_0, \hat{\beta}_1) + \left(\frac{\partial f}{\partial \beta_1}\right)^2 \widehat{\text{Var}}(\hat{\beta}_1)$$

Then

$$R = \frac{\hat{\theta} - \theta}{\sqrt{\hat{V}}}$$

is approximately pivotal. The bootstrap replicates

$$R^* = \frac{\hat{\theta}^* - \hat{\theta}}{\sqrt{\hat{V}^*}}$$

should have approximately same distribution and observable quantiles $R_{\alpha/2}^*$ and $R_{1-\alpha/2}^*$. Thus

$$P(R_{\alpha/2}^* < \frac{\hat{\theta} - \theta}{\sqrt{\hat{V}}} < R_{1-\alpha/2}^*) \approx 1 - \alpha$$

and so

$$\left(\hat{\theta} - \sqrt{\hat{V}} R_{1-\alpha/2}^* < \theta < \hat{\theta} - \sqrt{\hat{V}} R_{\alpha/2}^* \right)$$

is an approximate $1-\alpha$ conf. int for θ .

R demo

Bootstrapping time-series models

Ex. $Y_t - \mu = \phi(Y_{t-1} - \mu) + \varepsilon_t$, $\varepsilon_t \sim \text{WN}(\sigma^2)$ (1)

For $|\phi| < 1$, (1) has a causal stationary solution with stationary mean μ . Taking variance of (1), the stationary variance $\gamma_0 = \text{Var}(Y_t) = \text{Var}(Y_{t-1})$ satisfies

$$\gamma_0 = \phi^2 \gamma_0 + \sigma^2$$

$$\Rightarrow \gamma_0 = \frac{\sigma^2}{1 - \phi^2}$$

MLEs of μ , ϕ and σ^2 (assuming $\varepsilon_t \stackrel{\text{iid}}{\sim} N(0, \sigma^2)$) obtained by maximising $\pi(y_1, y_2, \dots, y_n | \theta, \phi, \sigma^2)$.

Let $\hat{\varepsilon}_t = y_t - \widehat{E}(y_t | y_{t-1})$, for $t = 2, 3, \dots, n$

Bootstrapping residuals: For $b = 1, 2, \dots, B$

1. Create $n+m+1$ bootstrap innovations $\varepsilon_{-m}^*, \varepsilon_{-m+1}^*, \dots, \varepsilon_n^*$ by resampling $\hat{\varepsilon}_2, \hat{\varepsilon}_3, \dots, \hat{\varepsilon}_n$ with replacement. "burnin"

2. Compute $y_{-m}^* = \varepsilon_{-m}^* + \hat{\mu}$ and for $t = -m+1, \dots, n$
$$y_t^* = \hat{\mu} + \hat{\phi}(\hat{y}_{t-1}^* - \hat{\mu}) + \varepsilon_t^*$$

3. Refit model to $y_1^*, y_2^*, \dots, y_n^*$ to obtain bootstrap replicates $\hat{\mu}^{b*}, \hat{\phi}^{b*}, \hat{\sigma}^{2b*}$

Scale parameter bootstrap CI

Approximate pivotal

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

Analogous to

$$\frac{s^2(n-1)}{\sigma^2} \sim \chi^2_{n-1}$$

Bootstrap replicates

$$V^* = \frac{\hat{\theta}^*}{\hat{\theta}}$$

has known quantiles $V^*_{\alpha/2}, V^*_{1-\alpha/2}$. Then

$$P\left(V^*_{\alpha/2} \leq \frac{\hat{\theta}}{\theta} \leq V^*_{1-\alpha/2}\right) = 1 - \alpha$$

$$P\left(\frac{\hat{\theta}}{V^*_{1-\alpha/2}} < \theta < \frac{\hat{\theta}}{V^*_{\alpha/2}}\right) = 1 - \alpha$$

EM-algorithm

14.1

Motivation:

Likelihood $f_Y(y|\theta)$ of "complete" data Y has simple form but only $X = M(Y)$ is observed where M is a many-to-fewer mapping.

Aim: Maximise "difficult" likelihood $f_X(x|\theta)$ of X .

Algorithm: Let

$$Q(\theta|\theta^{(t)}) = E(\ln f_Y(Y|\theta) | X=x, \theta = \theta^{(t)})$$

$$= \int_{\{y: M(y)=x\}} \ln f_Y(y|\theta) f_{Y|X}(y|x, \theta^{(t)}) dy$$

1. E-step: Find $Q(\theta|\theta^{(t)})$ in closed form

2. M-step: Set $\theta^{(t+1)} = \arg \max_{\theta} Q(\theta|\theta^{(t)})$

Repeat until convergence.

- Example: ABO - bloodtypes (Coppellini et al. 1955)

Genotypes	Observable phenotypes	Y	X
AA AO	A	n_{AA} n_{AO}	$n_A = n_{AA} + n_{AO}$
BB BO	B	n_{BB} n_{BO}	$n_B = n_{BB} + n_{BO}$
AB	AB	n_{AB}	n_{AB}
OO	O	n_{OO}	n_O

Model: $Y = (n_{AA}, n_{AO}, \dots, n_{OO}) \sim \text{multinomial}(n, p)$

where

$$p = \begin{bmatrix} p_A^2 \\ 2p_A p_O \\ p_B^2 \\ 2p_B p_O \\ 2p_A p_B \\ p_O^2 \end{bmatrix} \quad \left. \vphantom{\begin{bmatrix} p_A^2 \\ 2p_A p_O \\ p_B^2 \\ 2p_B p_O \\ 2p_A p_B \\ p_O^2 \end{bmatrix}} \right\} \text{Hardy-Weinberg proportions}$$

$$\text{and } p_O = 1 - p_A - p_B$$

The complete data log-likelihood is

$$\begin{aligned} & \ln f(n_{AA}, n_{Ao}, n_{Bo}, n_{Bb}, n_{AB}, n_{oo}) \\ &= \ln \left[\binom{n}{n_{AA}, n_{Ao}, n_{Bo}, n_{Bb}, n_{AB}, n_{oo}} (p_A^2)^{n_{AA}} (2p_A p_o)^{n_{Ao}} (p_B^2)^{n_{Bo}} (2p_B p_o)^{n_{Bb}} (2p_A p_B)^{n_{AB}} (p_o^2)^{n_{oo}} \right] \\ &= \ln \binom{n}{n_{AA}, n_{Ao}, n_{Bo}, n_{Bb}, n_{AB}, n_{oo}} + (2n_{AA} + n_{Ao} + n_{AB}) \ln p_A \\ & \quad + (2n_{Bo} + n_{Bb} + n_{AB}) \ln p_B + (2n_{oo} + n_{Ao} + n_{Bb}) \ln p_o \\ & \quad + (n_{Ao} + n_{Bo} + n_{AB}) \ln 2 \end{aligned}$$

E-step: Given $(p_A^{(t)}, p_B^{(t)}, p_o^{(t)})$ and the observed data (n_A, n_B, n_{AB}, n_o) ,

$$\begin{aligned} Q(\underline{\theta} | \underline{\theta}^{(t)}) &= E \left(\ln f(n_{AA}, n_{Ao}, n_{Bo}, n_{Bb}, n_{AB}, n_{oo}) \mid n_A, n_B, n_{AB}, n_o, p_A^{(t)}, p_B^{(t)}, p_o^{(t)} \right) \\ &= C + (2n_{AA}^* + n_{Ao}^* + n_{AB}) \ln p_A + (2n_{Bo}^* + n_{Bb}^* + n_{AB}) \ln p_B + (2n_{oo}^* + n_{Ao}^* + n_{Bb}^*) \ln p_o \end{aligned}$$

where

$$n_{AA}^* = E(n_{AA} \mid n_A, p_A^{(t)}, p_B^{(t)}, p_o^{(t)})$$

$$= n_A \cdot \frac{p_A^{(t)2}}{p_A^{(t)2} + 2p_A^{(t)} p_o^{(t)}}$$

$$n_{Ao}^* = n_A \cdot \frac{2p_A^{(t)} p_o^{(t)}}{p_A^{(t)2} + 2p_A^{(t)} p_o^{(t)}}$$

etc.. since $n_{AA} \mid n_A \sim \text{bin}(n_A, \quad)$

and C is a constant w.r.t. p_A, p_B, p_o (but not

w.r.t. $p_A^{(t)}, p_B^{(t)}, p_o^{(t)}$ since the conditional expectation of the log of the multinomial coefficient etc. depends on $p_A^{(t)}, p_B^{(t)}, p_o^{(t)}$)

M-step: Maximising $Q(\theta | \theta^{(t)})$ w.r.t. θ yields

$$\theta^{(t+1)} = (p_A, p_B, p_o)^{(t+1)}$$

$$= \frac{1}{2n} \left(2n_{AA}^* + n_{AO}^* + n_{AB}, 2n_{BB}^* + n_{Bo}^* + n_{AB}, 2n_o + n_{Ao}^* + n_{Bo}^* \right)$$

= sample frequencies of A, B and o.

- Ex.: Gaussian mixture:

$$x_1, x_2, \dots, x_n \sim f(x|\theta), \quad \theta = (\pi_1, \mu_1, \mu_2)$$

where

$$f(x|\theta) = \sum_{k=1}^2 \pi_k \phi(x - \mu_k), \quad \pi_2 = 1 - \pi_1$$

Full likelihood is

$$\begin{aligned} L(\theta) &= \prod_{i=1}^n f(x_i|\theta) \\ &= \prod_{i=1}^n \sum_{k=1}^2 \pi_k \phi(x_i - \mu_k) \end{aligned}$$

and log-likelihood

$$l(\theta) = \sum_{i=1}^n \ln \sum_{k=1}^2 \pi_k \phi(x_i - \mu_k)$$

Can be maximised numerically to find MLEs $\hat{\pi}_1, \hat{\mu}_1, \hat{\mu}_2$

Via EM-alg.: Introduce "missing" "component memberships"

$$z_1, z_2, \dots, z_n \stackrel{iid}{\sim} f(z|\pi_1) = \pi_1 \mathbb{I}(z=1) + (1-\pi_1) \mathbb{I}(z=2)$$

Letting

$$f(x_i|z_i=k) = \phi(x_i - \mu_k)$$

$x_1, \dots, x_n$ has the correct marginal distribution.

Conditional on x_i , and $\theta^{(t)} = (\mu_1^{(t)}, \mu_2^{(t)}, \pi_1^{(t)})$

$$w_{i,k}^{(t)} = f(z_i=k | x_i, \theta^{(t)}) = \frac{f(x_i|z_i=k, \theta^{(t)}) f(z_i=k | \theta^{(t)})}{f(x_i|\theta^{(t)})}$$

i.e

$$w_{i,1}^{(t)} = \frac{\phi(x_i - \mu_1) \pi_1}{\phi(x_i - \mu_1) \pi_1 + \phi(x_i - \mu_2) (1 - \pi_1)}$$

and

$$w_{i,2}^{(t)} = 1 - w_{i,1}^{(t)}$$

The likelihood based on "complete" data is

$$L(\theta) = \prod_{i=1}^n f(x_i | z_i, \theta) f(z_i | \theta)$$

and

$$\ell(\theta) = \sum_{i=1}^n \ln f(x_i | z_i, \theta) + \ln f(z_i | \theta)$$

E-step:

$$Q(\theta | \theta^{(t)}) = \sum_{i=1}^n E \left(\ln f(x_i | z_i, \theta) + \ln f(z_i | \theta) \mid x_i, \theta^{(t)} \right)$$

$$= \sum_{i=1}^n \sum_{k=1}^2 w_{i,k}^{(t)} \left(\ln f(x_i | z_i = k, \theta) + \ln f(z_i = k | \theta) \right)$$

$$= \sum_{i=1}^n \sum_{k=1}^2 w_{i,k}^{(t)} \left(-\frac{1}{2} (x_i - \mu_k)^2 + \ln \pi_k \right)$$

M-step:

$$\frac{\partial Q}{\partial \mu_k} = \sum_{i=1}^n w_{i,k}^{(t)} (x_i - \mu_k) = 0 \Rightarrow \mu_k^{(t+1)} = \frac{\sum_{i=1}^n x_i w_{i,k}^{(t)}}{\sum_{i=1}^n w_{i,k}^{(t)}}$$

for $k=1, 2$.

$$\frac{\partial Q}{\partial \pi_1} = \sum_{i=1}^n \left(\frac{w_{i,1}^{(t)}}{\pi_1} - \frac{1 - w_{i,1}^{(t)}}{1 - \pi_1} \right)$$

$$= \frac{1}{\pi_1(1-\pi_1)} \sum_{i=1}^n \left[(1-\pi_1) w_{i,1}^{(t)} - \pi_1 (1 - w_{i,1}^{(t)}) \right] = 0$$

$$\sum w_{i,1} - n \pi_1 = 0$$

$$\pi_1^{(t+1)} = \frac{1}{n} \sum_{i=1}^n w_{i,1}^{(t)}$$

Convergence: EM-alg. only guaranteed to find local optima or saddle point of $f_X(x|\theta)$.

Proof that $f_X(x|\theta^{(t+1)}) \geq f_X(x|\theta^{(t)})$.

Let $(X, Z) = M(Y)$ be a one-to-one mapping between X, Z and Y

Since

$$f_{Z|X}(z|x, \theta) = \frac{f_{X,Z}(x, z|\theta)}{f_X(x|\theta)}$$

$$\ln f_X(x|\theta) = \ln f_{X,Z}(x, z|\theta) - \ln f_{Z|X}(z|x, \theta)$$

and

$$E(\ln f_X(x|\theta) | x, \theta^{(t)}) = E(\ln f_{X,Z}(x, Z|\theta) | x, \theta^{(t)}) - E(\ln f_{Z|X}(Z|x, \theta) | x, \theta^{(t)})$$

i.e.

$$\ln f_X(x|\theta) \stackrel{c}{=} Q(\theta | \theta^{(t)}) - H(\theta | \theta^{(t)})$$

It follows that

$$-H(\theta | \theta^{(t)}) + H(\theta^{(t)} | \theta^{(t)})$$

$$= E\left(-\ln f_{Z|X}(Z|x, \theta) + \ln f_{Z|X}(Z|x, \theta^{(t)}) \mid x, \theta^{(t)}\right)$$

$$= E\left(-\ln \frac{f_{Z|X}(Z|x, \theta)}{f_{Z|X}(Z|x, \theta^{(t)})} \mid x, \theta^{(t)}\right)$$

$Eg(W) \geq g(EW)$ when $g(\cdot) = -\ln(\cdot)$ is convex.

$$\geq -\ln \left[E\left(\frac{f_{Z|X}(Z|x, \theta)}{f_{Z|X}(Z|x, \theta^{(t)})} \mid x, \theta^{(t)} \right) \right] \quad (\text{Jensen's ineq.})$$

$$= -\ln \int \underbrace{\frac{f_{Z|X}(z|x, \theta)}{f_{Z|X}(z|x, \theta^{(t)})} f_{Z|X}(z|x, \theta^{(t)})}_{=1} dz$$

$$= 0$$

for any θ .

Thus, since $Q(\theta^{(t+1)} | \theta^{(t)}) \geq 0$,

$$\begin{aligned}\ln f_X(x | \theta^{(t+1)}) &= Q(\theta^{(t+1)} | \theta^{(t)}) - H(\theta^{(t+1)} | \theta^{(t)}) \\ &\geq Q(\theta^{(t)} | \theta^{(t)}) - H(\theta^{(t)} | \theta^{(t)}) \\ &= \ln f_X(x | \theta^{(t)})\end{aligned}$$

Summary, part 1

Project 1, problem 2, & demo

$$F(x, y) = P(X \leq x \text{ and } Y \leq y)$$

$$\hat{F}(x, y) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}(X_i \leq x \text{ and } Y_i \leq y)$$

Ex.: Given $\underline{X} \sim N_2(\underline{\mu}, \Sigma)$, use Monte-Carlo integration to estimate

$$\mu = E h(\underline{X}) = E \left(\ln(w e^{X_1} + (1-w) e^{X_2}) \right)$$

simulate $\underline{Z}$ via Box-Muller and transform

to

$$\underline{X} = \underline{\mu} + A \underline{Z}$$

where A is cholesky factor or given by eig. decomp.

Simulate $\underline{X}_1, \underline{X}_2, \dots, \underline{X}_m$ and estimate μ by

$$\hat{\mu}_{MC} = \frac{1}{m} \sum_{i=1}^m h(\underline{X}_i)$$

Antithetic sampling: Let $\underline{X}_i^* = \underline{\mu} - A \underline{Z}_i$ and use

$$\hat{\mu}_A = \frac{1}{m} \sum_{i=1}^m \underbrace{\left(\frac{1}{2} (h(\underline{X}_i) + h(\underline{X}_i^*)) \right)}_{\hat{\mu}_{Ai}}$$

Randomized
Quasi-Monte Carlo (see e.g. Lemieux 2009, Springer)

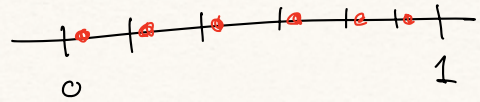
Univariate version:

For $i = 1, 2, \dots, m$

Generate V_i

For $j = 1, 2, \dots, n$,

Set $u_{ij} = \frac{j-1}{n} + V_i \bmod 1$



Set $x_{ij} = F^{-1}(u_{ij})$

$$\hat{\mu}_{\text{RQMC}} = \frac{1}{n} \sum_{i=1}^n \underbrace{\frac{1}{m} \sum_{j=1}^m h(F^{-1}(u_{ij}))}_{\hat{\mu}_{\text{QMC}, i}}$$

Multivariate version:

Generate $u_{i1}, u_{i2}, \dots, u_{in}$

using a Sobol-sequence

In R: `qrng::sobol( )`

R demo

Inversion method for truncated densities

Suppose that $X \sim F$ with quantile function F^{-1} .

$X | a < X \leq b$ can be simulated by

$$U \sim \text{unif}(F(a), F(b))$$

$$X \leftarrow F^{-1}(U)$$

Importance sampling, harmonic mean estimator of model evidence.

(the following also relates to project 1, problem 6)

Def.: If $\{X_n\}_{n=1}^{\infty}$ is a seq. of random variables s.t.

$$P\left(\lim_{n \rightarrow \infty} X_n = \mu\right) = 1$$

i.e.

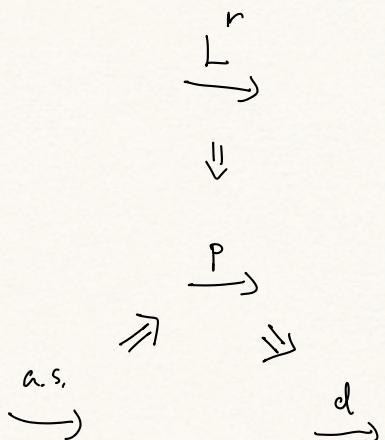
$$P\left(\left\{\omega : \lim_{n \rightarrow \infty} X_n(\omega) = \mu\right\}\right) = 1$$

then $X_n \xrightarrow{\text{a.s.}} \mu$ (almost surely)

Def.: If

$$\lim_{n \rightarrow \infty} E(|X_n - \mu|^2) = 0$$

then $X_n \xrightarrow{L^2} \mu$ (in mean square)



Importance sampling (self normalizing version)

Recall: If $X \sim f(x)$, then

$$\widehat{Eh(X)}_n = \frac{\sum_{i=1}^n h(x_i) \frac{f(x_i)}{g(x_i)}}{\sum_{i=1}^n \frac{f(x_i)}{g(x_i)}} \xrightarrow{\text{a.s.}} Eh(X) \quad \left(\begin{array}{l} \text{although not unbiased} \\ \text{(and not in mean sq.)} \end{array} \right)$$

where $x_1, \dots, x_n$ are iid samples from $g(x)$

Potential application:

Given prior $\pi(\theta)$, model $f(x|\theta)$ and

MCMC samples $\theta_1, \theta_2, \dots, \theta_n$ from

posterior $\pi(\theta|x) \propto f(x|\theta)\pi(\theta)$ how

can we estimate the model evidence

$$f(x) = \int f(x|\theta) \pi(\theta) d\theta \quad (\text{denominator in Bayes thm.})$$

$$= E \underbrace{f(x|\theta)}_{h(\theta)}, \quad \theta \sim \pi(\theta)$$

(important quantity in model selection)

Inefficient naive Monte Carlo estimate

Generate $\theta_1, \dots, \theta_n \sim \pi(\theta)$

Compute

$$\widehat{f(x)} = \frac{1}{n} \sum_{i=1}^n f(x|\theta_i)$$



Alternatively (self. norm. import. samp.):

Draw $\theta_1, \dots, \theta_n \sim \pi(\theta|x)$ (via MCMC)

Compute

$$\widehat{f(x)} = \frac{\sum_{i=1}^n \cancel{f(x|\theta_i)} \frac{\cancel{\pi(\theta_i)}}{\pi(\theta_i|x)}}{\sum_{i=1}^n \cancel{\pi(\theta_i)}} = \cancel{f(x|\theta_i)} \cancel{\pi(\theta_i)}$$

$$\prod_{i=1}^n \frac{f(x|\theta_i)}{\pi(\theta_i)} = f(x|\theta_i) \pi(\theta_i)$$

$$= \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{f(x|\theta_i)} \right)^{-1} \xrightarrow{\text{a.s.}} f(x),$$

the (strongly criticized) harmonic mean estimator (Newton & Raftery, 1994) of model evidence $f(x)$

Toy example:

$$x|\theta \sim N(\theta, 1) \quad (\text{likelihood})$$

$$\theta \sim N(0, 3) \quad (\text{prior})$$

Given $x=1$, the true model evidence, since the marginal density of $x \sim N(0, 4)$ is

$$f(1) = \frac{1}{\sqrt{2\pi \cdot 4}} e^{-\frac{(1-0)^2}{2 \cdot 4}} = \frac{1}{2\sqrt{\pi}} e^{-\frac{1}{32}}$$

The posterior

$$\theta|x=1 \sim N\left(\frac{0 \cdot \frac{1}{3} + 1 \cdot 1}{\frac{1}{3} + 1}, \frac{1}{\frac{1}{3} + 1}\right) = \left(\frac{3}{4}, \frac{3}{4}\right)$$

It follows that each term $\frac{1}{f(x|\theta_i)}$ in $\widehat{f(x)}$ has finite

$$\begin{aligned} E\left(\frac{1}{f(x|\theta_i)}\right) &= \int \frac{1}{f(x|\theta)} \frac{f(x|\theta) \pi(\theta)}{\int f(x|\theta) \pi(\theta) d\theta} d\theta \\ &= \frac{1}{\int f(x|\theta) \pi(\theta) d\theta} = \frac{1}{f(x|\theta)} \end{aligned}$$

but their second moments,

$$E\left[\left(\frac{1}{f(x|\theta)}\right)^2\right] = \int_{-\infty}^{\infty} \underbrace{\left(\frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(1-\theta)^2}\right)^2}_{f(x|\theta)} \underbrace{\frac{1}{\left(2\pi\frac{3}{4}\right)} e^{-\frac{(6-\frac{3}{4})^2}{2\frac{3}{4}}}}_{f(\theta|x)} d\theta$$

$$= \int_{-\infty}^{\infty} \dots e^{(6-1)^2 - \frac{2}{3}(6-\frac{3}{4})^2} d\theta = \infty$$

Thus, $\frac{1}{\widehat{f(x)}}$ does not converge in mean square to $\frac{1}{f(x)}$.

More generally, if using the prior $\theta \sim N(0, \sigma^2)$, the variance is infinite if

$$2 - \left(1 + \frac{1}{\sigma^2}\right)^{-1} > 0$$

$$2 - \frac{\sigma^2}{\sigma^2 + 1} > 0$$

$$2\sigma^2 + 2 - \sigma^2 > 0$$

$$\sigma^2 > 2$$

Constructing proposal from dominating terms of target

(as in project 2, problem 3)

$$x_1, \dots, x_n \stackrel{\text{iid}}{\sim} \exp(\lambda)$$

$$\pi(\lambda) \propto \frac{1}{\lambda} e^{-\frac{(\ln \lambda - \mu_0)^2}{2\sigma_0^2}} \quad (\text{non-conjugate log-normal prior})$$

Posterior

$$\pi(\lambda | x) \propto \underbrace{\lambda^{n-1} e^{-\lambda \sum x_i}}_{\propto Q(\lambda' | \lambda)} e^{-\frac{(\ln \lambda - \mu_0)^2}{2\sigma_0^2}}$$

Proposal

$$Q(\lambda' | \lambda) = \frac{(\sum x_i)^n}{\Gamma(n)} \lambda'^{n-1} e^{-\lambda' \sum x_i} \quad \left(\begin{array}{l} \text{does not depend} \\ \text{on } \lambda, \text{ thus we} \\ \text{have an independence} \\ \text{ sampler} \end{array} \right)$$

This leads to

$$\frac{\pi(\lambda' | x) Q(\lambda | \lambda')}{\pi(\lambda | x) Q(\lambda' | \lambda)} = e^{-\frac{1}{2\sigma_0^2} \underbrace{[(\ln \lambda' - \mu_0)^2 - (\ln \lambda - \mu_0)^2]}_{*}}$$

Should work well as long as $\pi(\lambda)$ is vague (* small)

MCMC, Constructing block proposals via conditioning.

(as in Project 2, problem 2c)

$$x_1, \dots, x_n \stackrel{i.i.d.}{\sim} f(x)$$

$$F(x) = 1 - e^{-ax^b} \quad (\text{Weibull}) \quad f(x) = abx^{b-1} e^{-ax^b}$$

$$\text{Prior} \quad \pi(a, b) \propto \frac{1}{a} \mathbb{I}(a > 0, b > 0)$$

$$\text{Posterior} \quad \pi(a, b | \underline{x}) \propto a^{n-1} b^n \left(\prod_{i=1}^n x_i \right)^b e^{-a \sum x_i^b}$$

Full conditionals

$$\pi(a | b, \underline{x}) = \frac{\left(\sum x_i^b \right)^n}{\Gamma(n)} a^{n-1} e^{-a \sum x_i^b}$$

$$\pi(b | a, \underline{x}) \propto b^n e^{-a \sum x_i^b}$$

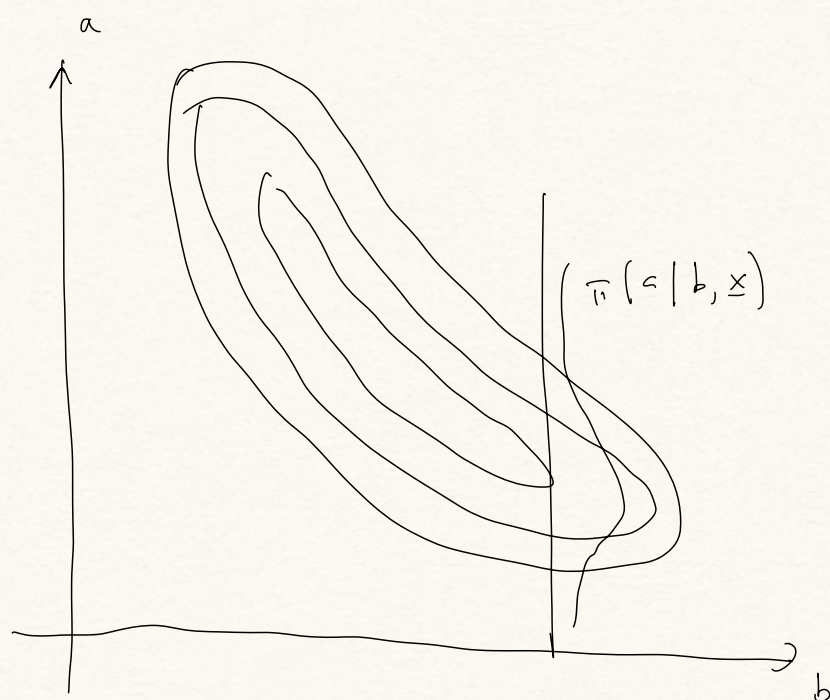
Block proposal:

$$Q(a', b' | a, b) = \underbrace{e^{-\frac{(b' - b)^2}{2\sigma^2}}}_{\text{random walk for } b'} \underbrace{\frac{\left(\sum x_i^{b'} \right)^n}{\Gamma(n)} a'^{n-1} e^{-a' \sum x_i^{b'}}}_{\pi(a' | b', \underline{x})}$$

This leads to

$$\frac{\pi(a', b' | \underline{x}) Q(a, b | a', b')}{\pi(a, b | \underline{x}) Q(a', b' | a, b)}$$

$$\begin{aligned}
&= \frac{a^{n-1} b'^n \left(\prod_{i=1}^n x_i \right)^b e^{-a' \sum x_i^{b'}} e^{-\frac{(b-b')^2}{2\sigma^2}} \frac{(\sum x_i^{b'})^n}{\Gamma(n)} a^{n-1} e^{-a \sum x_i^b}}{a^{n-1} b^n \left(\prod_{i=1}^n x_i \right)^b e^{-a \sum x_i^b} e^{-\frac{(b'-b)^2}{2\sigma^2}} \frac{(\sum x_i^{b'})^n}{\Gamma(n)} a^{n-1} e^{-a' \sum x_i^{b'}}} \\
&= \left(\frac{b' \sum x_i^b}{b \sum x_i^{b'}} \right)^n \left(\prod_{i=1}^n x_i \right)^{b'-b}
\end{aligned}$$



Similarly, in project 2 (and for latent Gaussian models in general),
a single-block proposal

$$Q(\underline{x}', \underline{\theta}' | \underline{x}, \underline{\theta}, \underline{y}) \propto \underbrace{\pi(\underline{x}' | \underline{\theta}', \underline{y})}_{\text{via Gaussian approximation around mode}} \underbrace{\pi(\underline{\theta}' | \underline{\theta})}_{\text{random-walk}}$$

$\hat{\underline{x}} = \arg \max_{\underline{x}} \pi(\underline{x} | \underline{\theta}', \underline{y})$

is feasible (see discussion section to Rue et. al. 2009, p. 356)

Summary, part 3

Example EM-als.: Suppose $X_1, \dots, X_n, Y_1, \dots, Y_n$ are indep. and $Y_i \sim \text{Pois}(\beta \tau_i)$, $X_i \sim \text{Pois}(\tau_i)$. (Casella & Berger)

Complete data likelihood

$$L(\beta, \underline{\tau}; \underline{x}, \underline{y}) = \prod_{i=1}^n \frac{e^{-\beta \tau_i} (\beta \tau_i)^{y_i}}{y_i!} \frac{e^{-\tau_i} \tau_i^{x_i}}{x_i!}$$

$$\ell(\beta, \underline{\tau}; \underline{x}, \underline{y}) = -(\beta + 1) \sum_{i=1}^n \tau_i + \sum (x_i + y_i) \ln \tau_i + (\sum y_i) \ln \beta$$

MLEs

$$\frac{\partial \ell}{\partial \beta} = -\sum \tau_i + \frac{\sum y_i}{\beta} = 0 \Rightarrow \beta = \frac{\sum y_i}{\sum \tau_i} = \frac{(\beta + 1) \sum y_i}{\sum (x_i + y_i)}$$

$$\frac{\partial \ell}{\partial \tau_i} = -(\beta + 1) + \frac{x_i + y_i}{\tau_i} = 0 \Rightarrow \hat{\tau}_i = \frac{x_i + y_i}{\beta + 1}$$

$$\beta \sum x_i + \cancel{\beta \sum y_i} = \cancel{\beta \sum y_i} + \sum y_i$$
$$\hat{\beta} = \frac{\sum y_i}{\sum x_i}$$

Suppose x_1 is missing.

E-step: Given $\theta^{(t)} = (\beta^{(t)}, \underline{\tau}^{(t)})$,

$$Q(\beta, \underline{\tau} | \beta^{(t)}, \underline{\tau}^{(t)}) = E(\ell(\beta, \underline{\tau}; \underline{X}, \underline{Y}) | \underline{X}_{-1} = \underline{x}_{-1}, \underline{Y} = \underline{y}, \beta^{(t)}, \underline{\tau}^{(t)})$$
$$= -(\beta + 1) \sum_{i=1}^n \tau_i + \left(x_1^* + \sum_{i=2}^n x_i + \sum_{i=1}^n y_i \right) \ln \tau_i + (\sum y_i) \ln \beta$$

where $x_1^* = E(X_1 | Y_1 = y_1, \beta^{(t)}, \tau_1^{(t)})$

$$= \beta^{(t)} \tau_1^{(t)}$$

M-step:

Maximising $Q(\quad)$ w.r.t. β and $\underline{\tau}$ yields

$$\beta^{(t+1)} = \frac{\sum_{i=1}^n \gamma_i}{x_1^* + \sum_{i=2}^n \gamma_i}$$

$$\tau_i^{(t+1)} = \begin{cases} \frac{x_i + \gamma_i}{\beta^{(t+1)} + 1} & \text{for } i \geq 2 \\ \frac{x_i^* + \gamma_i}{\beta^{(t+1)} + 1} & \text{for } i = 1 \end{cases}$$

Exercise:

Can we find the MLE via the observed data likelihood?

$$L(\beta, \underline{\tau}) = \frac{e^{-\beta \tau_1} (\beta \tau_1)^{\gamma_1}}{\gamma_1!} \prod_{i=2}^n \frac{e^{-\beta \tau_i} (\beta \tau_i)^{\gamma_i}}{\gamma_i!} \frac{e^{-\tau_i} x_i}{x_i!}$$

Show that the MLEs are

$$\hat{\beta} = \frac{\sum_{i=2}^n \gamma_i}{\sum_{i=2}^n x_i}, \quad \hat{\tau}_i = \frac{x_i + \gamma_i}{\beta + 1} \quad \text{for } i \geq 2$$

$$\tau_1 = \frac{\gamma_1}{\hat{\beta}}$$

Proof: Profile likelihood for $(\beta, \underline{\tau}_{-1})$? Log of first term (only term containing τ_1),

$$l_1(\beta, \tau_1) = -\beta \tau_1 + \gamma_1 (\ln \beta + \ln \tau_1),$$

is maximised when

$$\frac{\partial \ell_1}{\partial \alpha_1} = 0$$

$$-\beta + \frac{Y_1}{\alpha_1} = 0$$

$$\hat{\alpha}_1 = \frac{Y_1}{\beta}$$

Thus, the profile likelihood

$$L(\beta, \hat{\alpha}_1) = \frac{e^{-Y_1} Y_1^{Y_1}}{Y_1!} \prod_{i=2}^n \left(\frac{e^{-\beta} \beta^{Y_i}}{Y_i!} \right)$$

$\propto$ complete data likelihood
for observations $i = 2, 3, \dots, n$.