

Jo Eidsvik, May 2012

Exam TMA 4300

## COMPUTATIONAL STATISTICS

1. a)  $U = F(x) = 1 - e^{-\left(\frac{x}{\lambda}\right)^k}$

$$-\ln(1-u) = \left(\frac{x}{\lambda}\right)^k, \quad \underline{x} = \lambda^{1/k} \sqrt[k]{-\ln(1-u)}$$

$k=2$ :  $\underline{Y} = \lambda \sqrt{-\ln(1-u)}$  Exponential

$k=1$ :  $\underline{X} = -\lambda \ln(1-u)$  Rayleigh

$$\underline{Y} = \sqrt{2\underline{X}} = \sqrt{\lambda \cdot 2 \cdot -\ln(1-u)} = \lambda \sqrt{-\ln(1-u)}$$

b)  $T^b = X_1^b + \dots + X_{10}^b$ , where  $X_j^b \sim \text{Rayleigh}(\lambda)$

$b=1, \dots, B.$

Sort  $T^1, \dots, T^B$  from smallest to largest.

$$T^{(1)} < T^{(2)} < \dots < T^{(B)}$$

$$\hat{\text{median}} = \begin{cases} T^{(B/2)} + T^{(B/2+1)} & B \text{ even} \\ T^{(B/2+1)} & B \text{ odd} \end{cases}$$

c)  $T^b = X_1^b + \dots + X_{10}^b$  like in b).

$$\hat{p} = \frac{\sum_{b=1}^B I(T^b > \tau)}{B}$$

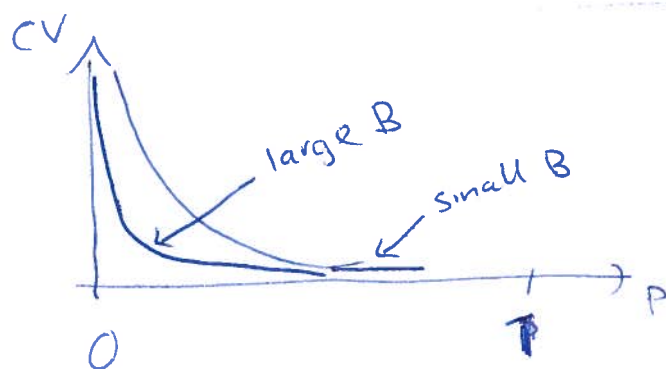
$B \cdot \hat{p}$  is binomial

$\sum_{b=1}^B \mathbb{I}(T^b > \tau) \sim \text{Bin}(B, p)$ , This means:

$$E(\hat{p}) = p$$

$$\text{Var}(\hat{p}) = \frac{p(1-p)}{B}$$

$$CV = \frac{\sqrt{\frac{p(1-p)}{B}}}{p} = \sqrt{\frac{1-p}{pB}}$$



For moderate CV for small  $p$ 's, the number of samples  $B$  needs to be very large.

d)  $\pi(x) \sim \text{Rayleigh}$   $\pi(x) = \left. \frac{dF}{dx} \right|_{x=2} = 2 \cdot \frac{x}{\lambda} e^{-\left(\frac{x}{\lambda}\right)^2}$

$q(x) \sim \text{Exponential}$   $q(x) = \left. \frac{dF}{dx} \right|_{x=1} = \frac{1}{\lambda} e^{-\frac{x}{\lambda}}$

Draw  $x_1^b, \dots, x_{10}^b$  from  $q(x)$   $T^b = x_1^b + \dots + x_{10}^b$

Evaluate  $W(x_1^b, \dots, x_{10}^b) = \prod_{i=1}^{10} \frac{\pi(x_i^b)}{q(x_i^b)}$

$$\hat{p}_{15} = \sum_{b=1}^B I(T^b > \tau) \cdot W(x_{1,}^b, \dots, x_{10,}^b)$$

which is unbiased  $E(\hat{p}_{15}) = p$ , since

$$p = E(I(T(x_{1,}, \dots, x_{10,}) > \tau))$$

$$= \int \int_{x_1, \dots, x_{10}} I(x_1, \dots, x_{10} > \tau) \cdot \pi(x_1) \cdots \pi(x_{10}) dx_1 \cdots dx_{10}$$

$$= \int \int_{x_1, \dots, x_{10}} I(x_1, \dots, x_{10} > \tau) \cdot W(x_1, \dots, x_{10}) \cdot q(x_1) \cdots q(x_{10}) dx_1 \cdots dx_{10}$$

2. a)  $f(\beta | y, q, \epsilon) \propto f(y | \beta, q, \epsilon) \cdot f(\beta)$

$$\propto \exp\left(-\frac{1}{2}(y - X\beta)^t R^{-1}(y - X\beta) - \beta^t S^{-1}\beta\right)$$

$$\propto \exp\left(\frac{1}{2}\beta^t (qX^t R^{-1}X + S^{-1})\beta + \beta^t (qX^t R^{-1}y)\right)$$

This is a Gaussian with covariance  
 $\Sigma_{\beta|y,q,\epsilon} = (qX^t R^{-1}X + S^{-1})^{-1}$  and mean

$$\mu_{\beta|y,q,\epsilon} = \Sigma_{\beta|y,q,\epsilon} [qX^t R^{-1}y]$$

Since Gaussian is  $y \sim N(\mu, \Sigma)$

$$\exp\left(-\frac{1}{2}(y - \mu)^t \Sigma^{-1}(y - \mu)\right)$$

$$\propto \exp\left(-\frac{1}{2}y^t \Sigma^{-1}y + y^t \Sigma^{-1}\mu\right)$$

iii)

$$a) f(\tau | \beta, y, \epsilon) \propto f(y | \beta, \tau, \epsilon) \cdot f(\tau)$$

$$\propto \frac{\tau^{n/2}}{|R|^{1/2}} \exp\left(-\frac{\tau}{2}(y - X\beta)^T R^{-1}(y - X\beta)\right) \cdot \tau^{a-1} e^{-b\tau}$$

$$\propto \tau^{\frac{n}{2}+a-1} \exp\left[-\tau\left(b + \frac{1}{2}(y - X\beta)^T R^{-1}(y - X\beta)\right)\right]$$

This is Gamma( $\frac{n}{2}+a, \frac{1}{2}(y - X\beta)^T R^{-1}(y - X\beta) + b$ )

b) A Gibbs sampler iterates between

Initiate  $\{\beta^0, \tau^0\}$  and  $\tau$  updates:  
 while  $b = 1, \dots, B$   
 $\{$   
 $\beta^b \sim f(\beta | y, \tau^{b-1}, \epsilon)$   
 $\tau^b \sim f(\tau | y, \beta^b, \epsilon)$   
 $\}$

This defines a Markov chain of  $(\beta^0, \tau^0), (\beta^1, \tau^1), \dots, (\beta^B, \tau^B)$  values that converges to  $f(\beta, \tau | y, \epsilon)$ .

c)

Proposed  $g(\phi|\phi(b)) = \text{Beta}(c, d)$ .

Accept rate:

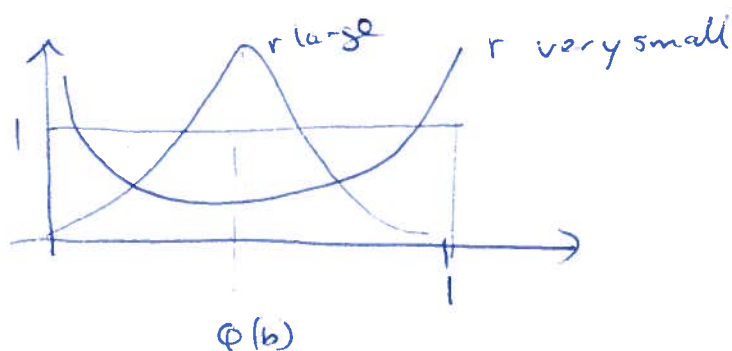
$$\alpha = \min\left(1, \frac{f(y|\phi, B, g) f(\phi) \cdot g(\phi(b)|\phi)}{f(y|\phi(b), B, g) f(\phi(b)) g(\phi|\phi(b))}\right)$$

$B, g$  fixed, so  $f(B) \cdot f(g)$  cancels.

$f(\phi) = f(\phi(b)) = 1$  also cancels.

$$\alpha = \min\left(1, \frac{|R(\phi)|^{-1/2} e^{-\frac{1}{2}(y-\mathbf{X}\beta)^T R(\phi)^{-1}(y-\mathbf{X}\beta)} \cdot \phi(b)^{c(\phi)-1} (1-\phi(b))^{d(\phi)-1}}{|R(\phi(b))|^{-1/2} e^{-\frac{1}{2}(y-\mathbf{X}\beta)^T R(\phi(b))^{-1}(y-\mathbf{X}\beta)} \cdot \phi^{c(\phi(b))-1} (1-\phi)^{d(\phi(b))-1}}\right)$$

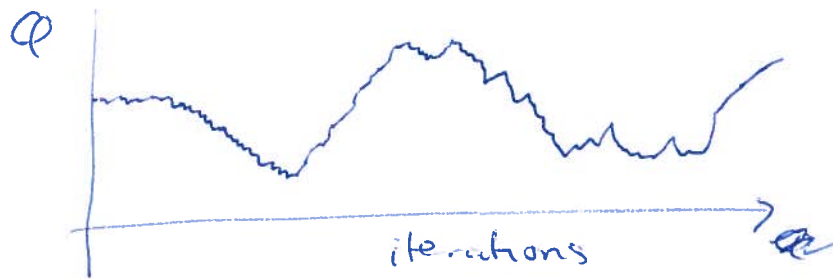
The proposal is beta-distributed, centered at the current value. If  $r$  is very small, the variance in the proposal ~~increases~~.



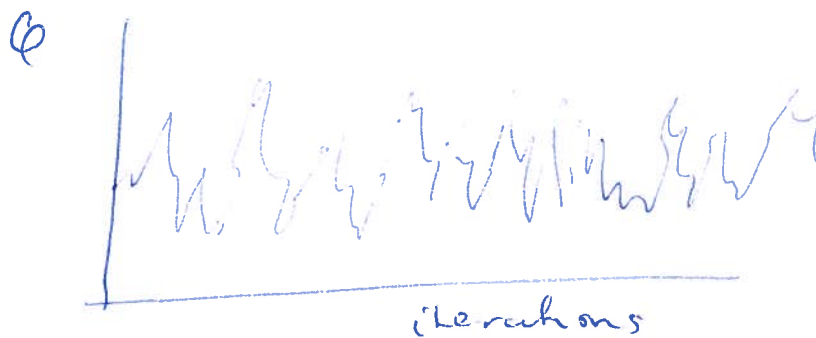
$r$  large gives smaller changes  $|\phi - \phi(b)|$ ,  
and the proposal is almost always accepted.  
 $r$  small gives larger changes  $|\phi - \phi(b)|$ ,  
and the proposal is often rejected.

v)

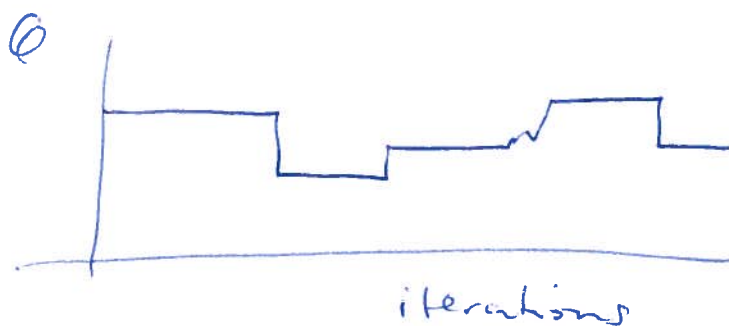
For intermediate  $r$  values, the mean acceptance rate is  $\sim 0.25-0.5$ , which tends to give reasonable mixing.



$r$  large



$r$  intermediate



$r$  small

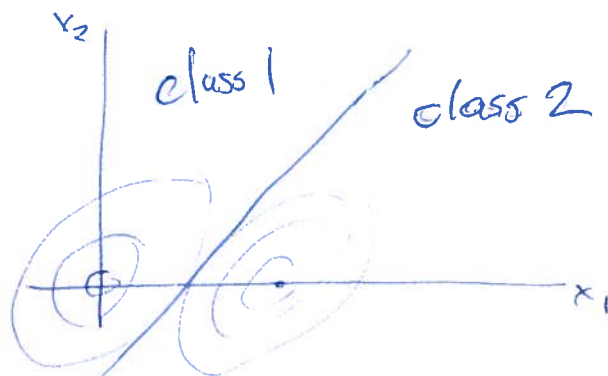
Illustration (traceplot) of  $\phi^1, \dots, \phi^B$  from a MH sampler with varying  $r$ -values.

3.

a) LDA picks  $k$ -class with largest density. When  $\pi_i = \frac{1}{k}$ :

$$k = \underset{l}{\operatorname{argmax}} f(x|l), \quad f(x|l) = N(x; \mu_l, \Sigma_l)$$

For  $\Sigma_1 = \Sigma_2 = \dots = \Sigma_k$ , this gives linear boundaries for classification



$$\Sigma^{-1} = \frac{1}{1-0.9^2} \begin{bmatrix} 1 & -0.9 \\ -0.9 & 1 \end{bmatrix} = \begin{bmatrix} 5.26 & -4.74 \\ -4.74 & 5.26 \end{bmatrix}$$

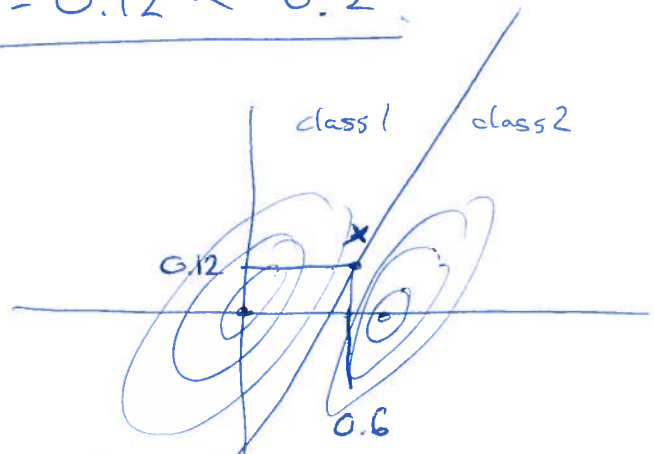
Boundary (since  $\pi_1 = \pi_2 = 0.5$ ):

$$\begin{aligned} \underbrace{x^t \Sigma^{-1} \mu_1}_0 - \underbrace{\frac{1}{2} \mu_1^t \Sigma^{-1} \mu_1}_0 &= x^t \Sigma^{-1} \mu_2 - \frac{1}{2} \mu_2^t \Sigma^{-1} \mu_2 \\ &= (x_1, x_2) \begin{bmatrix} 5.26 & -4.74 \\ -4.74 & 5.26 \end{bmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} - \frac{1}{2} (1, 0) \begin{bmatrix} 5.26 & -4.74 \\ -4.74 & 5.26 \end{bmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} \\ &= 5.26 \cdot x_1 - 4.74 x_2 - \frac{5.26}{2} \end{aligned}$$

$$\underline{x_2 = \frac{5.26}{4.74} x_1 - \frac{5.26}{4.74} \cdot \frac{1}{2} = 1.1 x_1 - 0.55}$$

$$x = (0.6, 0.2)$$

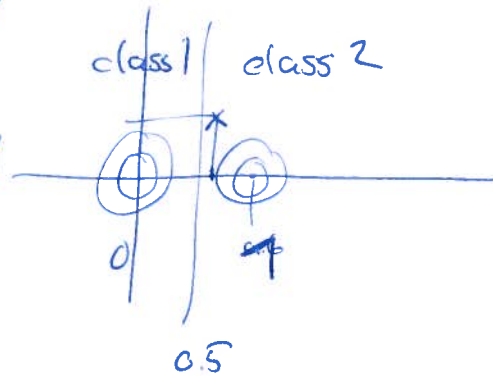
$$1.1 \cdot 0.6 - 0.55 = 0.12 < 0.2$$



Response is classified in class 1.

$$\text{If } \Sigma = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix},$$

the separation is made only on the first component, where  $\mu_1 \neq \mu_2$ .



Then  $x$  is classified

in class 2, since  $0.6 > 0.5$ .

3b)  $\hat{\mu}_1, \hat{\mu}_2, \hat{\Sigma}, \hat{\pi}_1$  are MLE

A parametric bootstrap uses the model (based on MLEs) to generate "new" data. From these "new" data, we compute the MLE's. As the number of sample increases, these are representative of the uncertainty of the  $\hat{\mu}_1, \hat{\mu}_2, \hat{\Sigma}, \hat{\pi}_1$ .

Idea

for  $b=1, \dots, B$

{ Simulate  $x^{*b} \sim \sum_{k=1}^K \frac{1}{\pi_k} N(x; \hat{\mu}_k, \hat{\Sigma})$ .

Compute

$\hat{\mu}_1^{*b}, \hat{\mu}_2^{*b}, \hat{\pi}_1^{*b}, \hat{\Sigma}^{*b}$  from  $x^{*b}$

The MLE's from data  $x^{*b}$

}

$\hat{\mu}_1^{*1}, \dots, \hat{\mu}_1^{*B}$  represent uncertainty in  $\hat{\mu}_1$   
 $\hat{\mu}_2^{*1}, \dots, \hat{\mu}_2^{*B}$  ————  $\hat{\mu}_2$   
 $\hat{\pi}_1^{*1}, \dots, \hat{\pi}_1^{*B}$  ————  $\hat{\pi}_1$   
 $\hat{\Sigma}^{*1}, \dots, \hat{\Sigma}^{*B}$  ————  $\hat{\Sigma}$ .