

Linear & Generalized Linear Models/ Advanced Regression for Independent Data

Andrew J. Spieker, Ph.D.

Associate Professor of Biostatistics
Vanderbilt University

Set 18 (Supplement 4): Bayesian methods for linear regression

Version: 07/03/2026

TABLE OF CONTENTS

- 1 Conjugate priors for linear regression
- 2 Metropolis-Hastings for linear regression
- 3 Example: Metropolis-Hastings from scratch

Linear regression:

- Suppose we're in the setting of linear regression.
- Model: $E[y|\mathbf{X}] = \mathbf{X}\boldsymbol{\beta} + \epsilon$, with $E[\epsilon] = \mathbf{0}$ and $\text{Cov}[\epsilon] = \sigma^2\mathbf{I}$.
- $\boldsymbol{\beta}$ is our target of inference, but σ^2 is a parameter of the data generating mechanism that needs attention. In the example that follows, it is easier to parameterize the model in terms of $\tau = 1/\sigma^2$.
- Note: In these slides, we're thinking of the models parametrically.

CONJUGATE PRIORS FOR LINEAR REGRESSION

Example 18.1: Conjugate prior

- The following is a Bayesian linear regression model:

$$\begin{aligned} \mathbf{y} | \boldsymbol{\beta}, \tau &\sim \mathcal{N}(\mathbf{X}\boldsymbol{\beta}, (1/\tau)\mathbf{I}), \\ \boldsymbol{\beta} | \tau &\sim \mathcal{N}(\boldsymbol{\phi}, (1/\tau)\mathbf{V}), \\ \tau &\sim \text{Gamma}(\alpha, \delta). \end{aligned}$$

- The first specification is a likelihood, a parametric model that might be correctly or incorrectly specified.
- The second two are prior specifications; conceptualizing “correct specification” doesn’t work quite the same way. The family needs to be flexible enough such that you can choose parameters that adequately represent your prior belief about $\boldsymbol{\beta}$ and τ .
- We assume, though we don’t *need* to, that $\boldsymbol{\phi}$, $\mathbf{V}$, α , and δ are fixed.

Prior specification:

- Choosing $\mathbf{V}$ to be a diagonal matrix with very large diagonal elements and a value of δ very close to zero would be an example of a weakly informative prior (on the scale of $\boldsymbol{\beta}$).
 - ▶ Choosing a prior with the goal of being “non-informative” may not always be a good idea, particularly when taking scale into account.

Example 18.2: Assay data

- Goal: Evaluate the relationship between an immunofluorescence assay (X) and a thermal assay (Y) to evaluate an (vaccine-induced) immunologic response to SARS-CoV-2.
- File: `assay.csv`.

CONJUGATE PRIORS FOR LINEAR REGRESSION

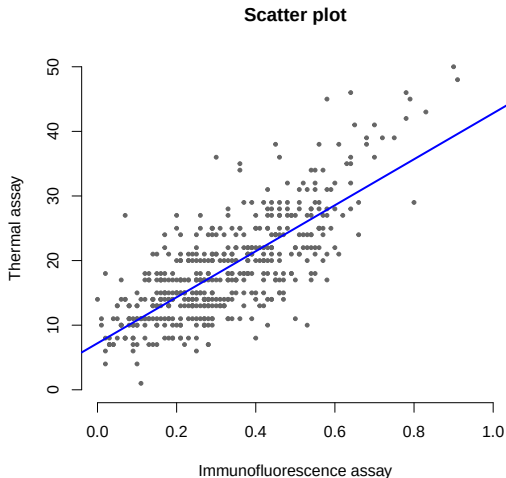


Figure S4.1: Assay data (linearity is approximated).



Example 18.2: Assay data

- Let's try a very basic prior that does not truly incorporate much of our prior knowledge about the data:
 - ▶ $\boldsymbol{\phi} = (20, 0)$.
 - ▶ $\mathbf{V} = 10\mathbf{I}$.
 - ▶ $\alpha = 2$.
 - ▶ $\delta = 100$.
- It is hard to understand just by looking at specifications for a prior what the data might look like under this combination of parameters. It is a good idea to check draws from the prior predictive distribution in order to make that assessment.
- Suppose we know in advance that the thermal assay is bounded between 0 and 50. We may want to check that our prior does not provide draws outside of this range too often (some such draws are okay).

CONJUGATE PRIORS FOR LINEAR REGRESSION

Example 18.2: Prior predictive (computational)

```
1 ## Read in data
2 dat <- read.csv("assay.csv")
3 x <- dat$immuno
4 y <- dat$therm
5 n <- dim(dat)[1]
6
7 ## Specify prior
8 phi <- c(20, 0)
9 V <- matrix(c(10, 0, 0, 10), nrow = 2)
10 alpha <- 2
11 delta <- 100
12
13 ## Set seed for reproducibility
14 set.seed(7345)
15
16 ## Prior predictive
17 taus <- matrix(0, nrow = n, ncol = 1)
18 betas <- matrix(0, nrow = n, ncol = 2)
19 for (m in 1:n)
20 {
21   tau.sample <- rgamma(1, shape = alpha, rate = delta)
22   beta.sample <- rmvnorm(1, mean = phi, sigma = (1/tau.sample)*V)
23   betas[m,] <- beta.sample
24   taus[m] <- tau.sample
25 }
```

Example 18.2: Prior predictive (computational)

```
1 ## Generate random draws
2 y.tildes <- rnorm(n, mean = betas[,1] + betas[,2]*x, sd = 1/sqrt(taus[,1]))
3
4 ## Plot densities repeatedly
5 plot(density(y.tildes), col = "gray90", ylim = c(0, 0.06), frame.plot = FALSE,
6      xlab = expression(tilde(y)), ylab = "Density", main = "Prior predictive: Poor choice",
7      xlim = c(-100, 100))
8 for (j in 1:500)
9 {
10   y.tildes <- rnorm(n, mean = betas[,1] + betas[,2]*x, sd = 1/sqrt(taus[,1]))
11   lines(density(y.tildes), col = "gray90")
12 }
13
14 ## Real data overlaid
15 lines(density(y), col = "blue", lwd = 2)
16 ## But we'd generally do this before we have the actual data
```

CONJUGATE PRIORS FOR LINEAR REGRESSION

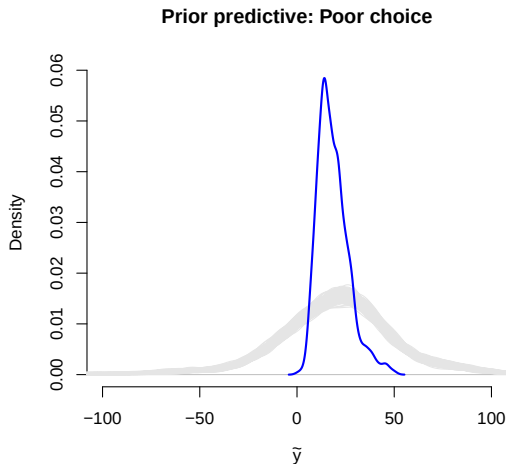


Figure S4.2: Prior predictive draws (poor choice).



Example 18.2: Assay data

- This is an opportunity to better understand prior specification.
- Here may be some knowledge that we want to exploit:
 - ▶ Mean thermal assay ≈ 0 when immunofluorescence = 0.
 - ▶ Mean thermal assay ≈ 50 when immunofluorescence = 1.
 - ▶ There should be a strictly positive association.
 - ▶ If $\beta_0 > 0$, β_1 “less positive,” and vice versa.
 - ▶ Conditional SD of Y (given X) between 5-10 thermal units.
- Let's update out prior to reflect this.
 - ▶ $\boldsymbol{\phi} = (0, 50)^\top$.
 - ▶ $\mathbf{V} = \begin{bmatrix} 0.6 & -0.05 \\ -0.05 & 0.1 \end{bmatrix}$.
 - ▶ $\alpha = 5$.
 - ▶ $\delta = 500$.

CONJUGATE PRIORS FOR LINEAR REGRESSION

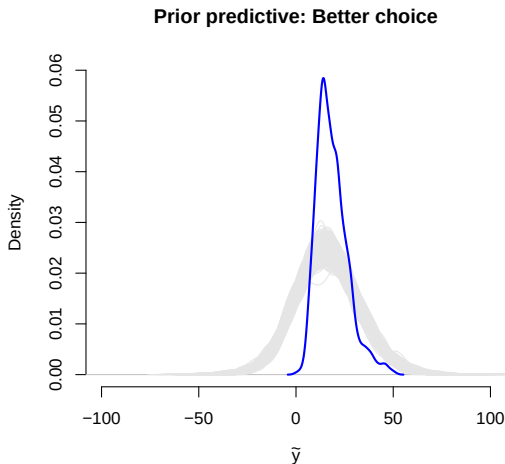


Figure S4.3: Prior predictive draws (better choice).



CONJUGATE PRIORS FOR LINEAR REGRESSION

Posterior: Conjugate relationship

- It can be shown (not-so-light algebra) that $\pi(\boldsymbol{\beta}, \tau | \mathbf{y})$, is of the same family (normal/Gamma) as the joint prior, $\pi(\boldsymbol{\beta}, \tau)$, which is to say the prior is conjugate.
- The *marginal* posterior $\pi(\boldsymbol{\beta} | \mathbf{y})$, takes the form (not-so-light algebra) of a multivariate t -distribution, $t_{N+2\alpha}(\boldsymbol{\phi}_*, \boldsymbol{\Sigma}_*)$, where:

$$\begin{aligned}\boldsymbol{\phi}_* &= (\mathbf{V}^{-1} + \mathbf{X}^T \mathbf{X})^{-1}(\mathbf{V}^{-1} \boldsymbol{\phi} + \mathbf{X}^T \mathbf{y}), \\ \boldsymbol{\Sigma}_* &= \left[\frac{(\mathbf{y} - \mathbf{X} \boldsymbol{\phi})^T (\mathbf{I} + \mathbf{X} \mathbf{V} \mathbf{X}^T)^{-1} (\mathbf{y} - \mathbf{X} \boldsymbol{\phi}) + 2\delta}{N + 2\alpha} \right] (\mathbf{V}^{-1} + \mathbf{X}^T \mathbf{X})^{-1}\end{aligned}$$

- $E[\boldsymbol{\beta} | \mathbf{y}] = \boldsymbol{\phi}_*$ and $\text{Cov}[\boldsymbol{\beta} | \mathbf{y}] = [(N + 2\alpha)/(N + 2\alpha - 2)] \boldsymbol{\Sigma}_*$.
- Let's think about these formulas in the context of the last statement of the previous slide.

Point estimation:

- In our simple example, ϕ_* marks the posterior mean, posterior median, and posterior mode. This will not always be the case.

Posterior: Conjugate relationship

- The *marginal* posterior $\pi(\tau|\mathbf{y})$, takes the form (not-so-light algebra) of a $\text{Gamma}(\alpha_*, \delta_*)$ distribution, where:

$$\alpha_* = \alpha + N/2, \text{ and}$$

$$\delta_* = (-\boldsymbol{\phi}_*^T(\mathbf{V}^{-1} + \mathbf{X}^T\mathbf{X})\boldsymbol{\phi}_* + \boldsymbol{\phi}^T\mathbf{V}^{-1}\boldsymbol{\phi} + \mathbf{y}^T\mathbf{y} + 2\delta)/2.$$

- Intuitive explanation a bit more challenging.
- This knowledge makes it easier to determine posterior predictive distribution.

Interval estimation:

- Conveniently, it turns out in our conjugate prior example that

$$\frac{\mathbf{c}^\top \boldsymbol{\phi}_* - \mathbf{c}^\top \boldsymbol{\beta}}{\sqrt{\mathbf{c}^\top \boldsymbol{\Sigma}_* \mathbf{c}}} \sim t_{N+2\alpha}.$$

- More generally, a suitably centered and scaled a linear combination of a multivariate t -distributed random vector follows a univariate t -distribution.
- Choosing $\mathbf{c} = (0, \dots, 0, 1, 0, \dots, 0)^\top$, zeros everywhere except the k^{th} position, informs us about the marginal posterior for $\beta_k | \mathbf{y}$:

$$\frac{\beta_k - \phi_{*,k}}{\sigma_{*,ii}} \sim t_{N+2\alpha}.$$

- Sensible 95% credible interval: $\phi_{*,k} \pm t_{0.975, N+2\alpha} \sigma_{*,ii}$.

Posterior: Posterior mean and variance for β_1

```
1 posterior <- function(phi0, phi1, v1, v2, v3, alpha, delta, x, Y)
2 {
3   N <- length(Y)
4   X <- cbind(1, x)
5   phi <- matrix(c(phi0, phi1), nrow = 2)
6   V <- matrix(c(v1, v2, v2, v3), nrow = 2)
7   post.mean <- as.numeric(solve(solve(V) + t(X) %*% X) %*% (solve(V) %*% phi + t(X) %*% Y))
8   post.var <- matrix(as.numeric(((t(Y - X %*% phi) %*%
9                                     solve((diag(N) + X %*% V %*% t(X))) %*%
10                                      (Y - X %*% phi) + 2*delta)/(N + 2*alpha - 2))) *
11                     solve((solve(V) + t(X) %*% X)), nrow = 2)
12   out <- list(mean = post.mean, var = post.var)
13   return(out)
14 }
```

Example 18.2: 95% credible interval for β_1

```
1 ## Posterior
2 pst <- posterior(0, 50, 0.6, -0.05, 0.1, 5, 500, x, y)
3
4 ## Parameters of marginal posterior
5 phistar <- as.numeric(pst$mean)
6 Sigmastar <- pst$var
7
8 ## Interested in beta1
9 c <- matrix(c(0,1), ncol = 1)
10
11 ## Define sigmall
12 sigmall <- sqrt((t(c) %*% Sigmastar %*% c))
13
14 ## Point estimate
15 > t(c) %*% phistar
16      [,1]
17 [1,] 41.442
18
19 ## 95% credible interval
20 crt <- qt(0.975, df = n + 2*alpha)
21 > c(t(c) %*% phistar - crt * sigmall,
22    t(c) %*% phistar + crt * sigmall)
23 [1] 39.259 43.626
```

CONJUGATE PRIORS FOR LINEAR REGRESSION

Example 18.2: 95% confidence interval for β_1

```
1 ## OLS
2 zz <- lm(y ~ x)
3
4 ## Coefficients
5 > summary(zz)
6
7 Call:
8 lm(formula = y ~ x)
9
10 Residuals:
11      Min       1Q   Median       3Q      Max
12 -16.074  -3.173  -0.037   3.123  18.115
13
14 Coefficients:
15             Estimate Std. Error t value Pr(>|t|)
16 (Intercept)    7.205     0.489    14.7  <2e-16 ***
17 x             35.601     1.350    26.4  <2e-16 ***
18 ---
19 ...
20 Residual standard error: 5.16 on 498 degrees of freedom
21 Multiple R-squared:  0.583, Adjusted R-squared:  0.582
22 F-statistic: 696 on 1 and 498 DF, p-value: <2e-16
23
24 ## 95% confidence interval
25 > confint(zz)
26             2.5 %   97.5 %
27 (Intercept)  6.244  8.1655
28 x           32.949 38.2536
```

Posterior predictive:

- In this example, we have a nice result for the posterior predictive distribution:

$$\tilde{\mathbf{y}}|\mathbf{X} \sim t_{N+2\alpha}\left(\mathbf{X}\boldsymbol{\phi}_*, \frac{\delta_*}{\alpha_*}(\mathbf{I}_N + \mathbf{X}(\mathbf{V}^{-1} + \mathbf{X}^\top\mathbf{X})^{-1}\mathbf{X}^\top)\right).$$

- Each draw, $\tilde{y}$, will follow a (shifted/scaled) univariate t -distribution.
- Suppose that we didn't know this form for the posterior predictive distribution. Could we generalize our previous trick to generate draws computationally?
 - ▶ You bet!

CONJUGATE PRIORS FOR LINEAR REGRESSION

Example 18.2: Posterior predictive

```
1 ## Design matrix and outcome
2 X <- cbind(1,x)
3 y <- cbind(y)
4
5 ## Posterior hyperparameters
6 phistar <- solve(solve(V) + t(X) %*% X) %*% (solve(V) %*% phi + t(X) %*% y)
7 alphastar <- alpha + n/2
8 deltastar <- (-t(phistar) %*% (solve(V) + t(X) %*% X) %*% phistar
9             + t(phi) %*% solve(V) %*% phi + t(y) %*% y + 2*delta)/2
10
11 ## Storage for posterior draws
12 taus <- matrix(0, nrow = n, ncol = 1)
13 sigmas <- matrix(0, nrow = n, ncol = 1)
14 betas <- matrix(0, nrow = n, ncol = 2)
15
16 for (m in 1:n)
17 {
18   tau.sample <- rgamma(1, alphastar, deltastar)
19   beta.sample <- rmvnorm(1, mean = phistar,
20                        sigma = (1/tau.sample)*solve(solve(V) + t(X) %*% X))
21   taus[m,1] <- tau.sample
22   betas[m,] <- beta.sample
23 }
24
25 ## Draw from posterior predictive
26 y.tildes <- rnorm(n, mean = betas[,1] + x*betas[,2], sd = 1/sqrt(taus[,1]))
```

CONJUGATE PRIORS FOR LINEAR REGRESSION

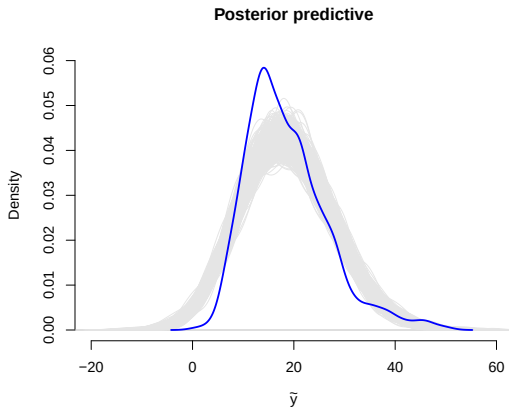


Figure S4.4: Posterior predictive draws.



CONJUGATE PRIORS FOR LINEAR REGRESSION

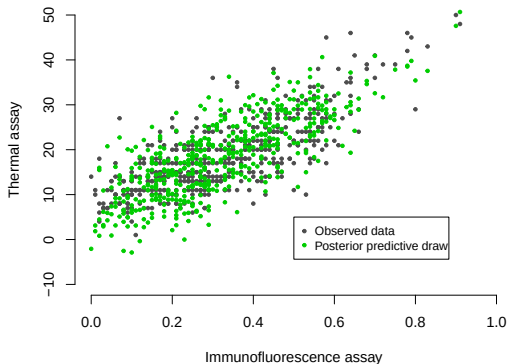


Figure S4.5: Posterior predictive draws (conditional on X).



TABLE OF CONTENTS

- 1 Conjugate priors for linear regression
- 2 Metropolis-Hastings for linear regression
- 3 Example: Metropolis-Hastings from scratch

Bayesian linear regression model:

- The following is a Bayesian linear regression model:

$$\begin{aligned} \mathbf{y} | \boldsymbol{\beta}, \sigma &\sim \mathcal{N}(\mathbf{X}\boldsymbol{\beta}, \sigma^2 \mathbf{I}), \\ \boldsymbol{\beta} &\sim \mathcal{N}(\boldsymbol{\phi}, \mathbf{V}), \\ \sigma &\sim \text{Gamma}(\alpha, \delta). \end{aligned}$$

- This example differs from the previous example in that:
 - ▶ We parameterize in terms of σ in place of $\tau = 1/\sigma^2$.
 - ▶ $\boldsymbol{\beta}$ and σ are independent.
- Again, assume $\boldsymbol{\phi}$, $\mathbf{V}$, α , and δ are fixed.
- We seek a computational tool to sample draws from $\pi(\boldsymbol{\beta}, \sigma | \mathbf{y})$.

METROPOLIS-HASTINGS FOR LINEAR REGRESSION

Algorithm: Metropolis-Hastings (Big picture)

- Once you're satisfied with your model, the general concept behind the Metropolis-Hastings algorithm is as follows:
 - 1 Initialize a posterior draw as $\theta^{(0)}$ (it doesn't have to be a great guess, but it does have to be "not awful").
 - 2 Starting from a posterior draw, $\theta^{(j)}$, take a random (but structured, in a way we will discuss) step away to a "nearby" value of θ , call it $\theta_{\text{prop}}^{(j+1)}$.
 - 3 With stochasticity (we will discuss), evaluate the plausibility of $\theta_{\text{prop}}^{(j+1)}$ as a candidate to include as a valid posterior draw.
 - 4 If algorithm says "yes" in Step 3, that new draw can be added to your collection of posterior draws **and is your new basis to go back to Step 2 and repeat the process** (that is $\theta^{(j+1)} = \theta_{\text{prop}}^{(j+1)}$). If "no," go back to Step 2 and try again (i.e., from $\theta^{(j)}$, the last accepted draw).
 - 5 Iterate Steps 2-4 until there are enough samples from the apparent stationary distribution.
- Under some assumptions, the stationary distribution is $\pi(\theta|\mathbf{y})$.

Metropolis-Hastings: Elaboration on Step 2 in context of example

- Example:

$$\begin{aligned}\boldsymbol{\beta} &\sim \mathcal{N}(\boldsymbol{\phi}, \mathbf{V}), \\ \sigma &\sim \text{Gamma}(\alpha, \delta).\end{aligned}$$

- Starting from $\boldsymbol{\theta}^{(j)}$, consider taking a random draw, $\boldsymbol{\theta}_{\text{prop}}^{(j+1)}$ from the following proposal distribution:

$$\begin{aligned}\boldsymbol{\beta}_{\text{prop}}^{(j+1)} &\sim \mathcal{N}(\boldsymbol{\beta}^{(j)}, \omega_1 \mathbf{I}), \\ \sigma_{\text{prop}}^{(j+1)} &\sim \text{Gamma}(\sigma^{(j)} / \omega_2, 1 / \omega_2).\end{aligned}$$

- The distribution for the draw is centered at the previous accepted proposal, $\boldsymbol{\theta}^{(j)}$; ω_1 and ω_2 control the search window (variance).
- Why are my proposals of the same family as my prior? Although they don't have to be, there is a convenience aspect in this example.

Metropolis-Hastings: Elaboration on Step 2 in context of example

- The values of ω_1 and ω_2 control how close to $\theta^{(j)}$ you'd like your subsequent proposal to be, inviting the following trade-off:
 - ▶ Larger window: search for new draws in unexplored areas. Too large, and the algorithm may fail to converge to a stationary distribution.
 - ★ In our example, higher values of ω_1 and ω_2 cast a wider net as compared to smaller values.
 - ▶ Narrower window: refine search closer to the areas where good candidates have been identified. Too narrow, and the algorithm may not find the optimum in a reasonable amount of time.
- See the third previously possibility listed for “what can go wrong.”

METROPOLIS-HASTINGS FOR LINEAR REGRESSION

Metropolis-Hastings: Elaboration on Step 3 in context of example

- Once we have our candidate, $\boldsymbol{\theta}_{\text{prop}}^{(j+1)}$, we compare it to our previously accepted step, $\boldsymbol{\theta}^{(j)}$, in the following way:

$$Q = \frac{f(\mathbf{y}|\mathbf{X}; \boldsymbol{\theta}_{\text{prop}}^{(j+1)})}{f(\mathbf{y}|\mathbf{X}; \boldsymbol{\theta}^{(j)})} \cdot \frac{\pi(\boldsymbol{\theta}_{\text{prop}}^{(j+1)})}{\pi(\boldsymbol{\theta}^{(j)})} \cdot \frac{g_{\boldsymbol{\omega}}(\boldsymbol{\theta}^{(j)}|\boldsymbol{\theta}_{\text{prop}}^{(j+1)})}{g_{\boldsymbol{\omega}}(\boldsymbol{\theta}_{\text{prop}}^{(j+1)}|\boldsymbol{\theta}^{(j)})},$$

where, for instance, $g_{\boldsymbol{\omega}}(\boldsymbol{\theta}_{\text{prop}}^{(j+1)}|\boldsymbol{\theta}^{(j)})$ evaluates the proposal density of $\boldsymbol{\theta}_{\text{prop}}^{(j+1)}$ assuming parameters $\boldsymbol{\theta}^{(j)}$.

- The value of $\boldsymbol{\omega} = (\omega_1, \omega_2)$ does not update across iterations.
- Let $U \sim \text{Uniform}(0, 1)$ denote an iteration-specific random draw. If $Q > U$, we accept the candidate, so that $\boldsymbol{\theta}^{(j+1)} = \boldsymbol{\theta}_{\text{prop}}^{(j+1)}$. Otherwise, go back to Step 2 and try generating another proposal.
- If the proposal distributions are symmetric, the last fraction should be one (Metropolis-Hastings $\rightarrow$ Metropolis).

Metropolis-Hastings: That's it!

- That's the Metropolis-Hastings algorithm!
- It is not in any way “obvious” that this algorithm would have the property of converging to the posterior as a stationary distribution (again, such technical details are not in the scope of this course).
- However, we can gain some intuition for the value of Q .
 - ▶ First note that—all else being equal—higher values of Q are more likely to be accepted as compared to lower values. Q is higher when...
 - ★ $\theta_{\text{prop}}^{(j+1)}$ is more consistent with the likelihood as compared to $\theta^{(j)}$.
 - ★ $\theta_{\text{prop}}^{(j+1)}$ is more consistent with the prior as compared to $\theta^{(j)}$.
 - ★ $\theta_{\text{prop}}^{(j+1)}$ is *less* consistent with the proposal as compared to $\theta^{(j)}$.
- Intuitively, this procedure encourages the acceptance of draws that are more consistent with your model and prior belief, but can also be designed to encourage exploration of under-explored areas.

More notes: Assessment of stationarity

- To assess whether a stationary distribution has been achieved, you'd be nervous about *autocorrelation* (tendency for nearby iterations to be closer to one another as compared to other further-away iterations).
 - ▶ Of course, because this is a Markov chain, we're not expecting to see the random draws completely uncorrelated.
 - ▶ High autocorrelation reduces the effective number of iterations.
 - ▶ When there is high autocorrelation, we sometimes benefit from reparameterizing the model.
- Thinning is the process by which certain draws are discarded to reduce autocorrelation. This is generally not considered a good practice.
- If chains appear not to be approximately stationary after a moderate number of iterations, a good strategy—other than searching for and fixing bugs in the code—may be to tweak the model (as opposed to just running more iterations and hoping for the best).

More notes: Strategies

- One typically runs MCMC with different initial parameters to evaluate whether they converge to the same stationary distribution.
 - ▶ When this appears to be the case, we say there is *good mixing*.
 - ▶ *Slow mixing* signifies challenges.
- If after many iterations there does not appear to be convergence to a stationary distribution, there are many things that could have gone wrong. Possibilities include:
 - ▶ A prior that is insufficiently informative.
 - ▶ A poor initializer.
 - ▶ A poor choice of a proposal “window” (I will elaborate on this).
 - ▶ Parameters that are too tightly correlated.

TABLE OF CONTENTS

- 1 Conjugate priors for linear regression
- 2 Metropolis-Hastings for linear regression
- 3 Example: Metropolis-Hastings from scratch

Example 18.3: Toy data

- Let's return to the immune assay data set and implement a Bayesian analysis using Metropolis-Hastings from scratch.

Example 18.3: Setting the prior

- Prior for $\boldsymbol{\beta}$:

$$\boldsymbol{\beta} \sim \mathcal{N}(\boldsymbol{\phi}, \mathbf{V}).$$

- This time, let's try the following prior hyperparameters:
 - ▶ $\boldsymbol{\phi} = (0, 50)^\top$.
 - ▶ $\mathbf{V} = \begin{bmatrix} 50 & 0 \\ 0 & 200 \end{bmatrix}$.

Example 18.3: Setting the prior

- Recall the form of the prior for σ :

$$\sigma \sim \text{Gamma}(\alpha, \delta).$$

- This time, let's try the following prior hyperparameters:
 - ▶ $\alpha = 3$ and $\delta = 3/5$.

Example 18.3: Coding the prior

```
1 ## Prior
2 pr <- function(theta, alpha = 3, delta = 3/5, v1 = 50, v2 = 200, phi = c(0,50)) {
3   d.beta <- dmvnorm(theta[1:2], mean = phi, sigma = matrix(c(v1, 0, 0, v2), nrow = 2))
4   d.sigma <- dgamma(theta[3], shape = alpha, rate = delta)
5   out <- log(d.beta) + log(d.sigma)
6   return(out)
7 }
```

Please note that use of the log-density has better computational stability.

Example 18.3: Setting the proposal

- Now, starting from $\theta^{(j)}$, we want to have a suitable window in which to propose our next draw:

$$\begin{aligned}\beta_{\text{prop}}^{(j+1)} &\sim \mathcal{N}(\beta^{(j)}, \omega_1 \mathbf{I}), \\ \sigma_{\text{prop}}^{(j+1)} &\sim \text{Gamma}(\sigma^{(j)} / \omega_2, 1/\omega_2).\end{aligned}$$

- Consider $\omega_1 = 0.5$ and $\omega_2 = 1/100$.
 - ▶ Modifying window will affect the properties of the algorithm.
 - ▶ Consider an exploration of this sort to be a good use of time.
- The most essential part is to center the draws to have mean $\theta^{(j)}$.
- That ω_1 and ω_2 are directly proportional to the variances is nice, but not always straightforward to achieve (nor is it essential). However, you must understand how the windows “control” the variance.
 - ▶ Example: Upcoming homework problem with a Weibull prior.

METROPOLIS-HASTINGS FROM SCRATCH

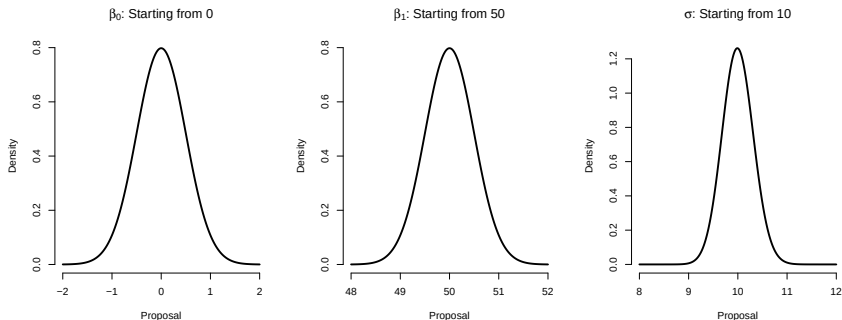


Figure S4.6: Proposal densities.



Example 18.3: Coding the proposal

```
1 ## Proposal
2 proposal <- function(theta, omega1, omega2) {
3   prop <- rep(0,3)
4   prop[1:2] <- rmvnorm(1, mean = theta[1:2], sigma = diag(2)*omega1)
5   prop[3] <- rgamma(1, shape = theta[3]/omega2, rate = 1/omega2)
6   return(prop)
7 }
```

Example 18.3: Coding the log-likelihood

```
1 ## Log-likelihood
2 loglik <- function(theta, x, y) {
3   out <- sum(log(dnorm(y, mean=theta[1] + theta[2]*x, sd=theta[3])))
4   return(out)
5 }
```

Example 18.3: Coding the proposal density ratio

```
1 ## Proposal ratio
2 prop.ratio <- function(theta.j, theta.proposed, omegal, omega2) {
3   d1 <- dmvnorm(theta.j[1:2], mean = theta.proposed[1:2], sigma = diag(2)*omegal)
4   d2 <- dgamma(theta.j[3], shape = theta.proposed[3]/omega2, rate = 1/omega2)
5   d3 <- dmvnorm(theta.proposed[1:2], mean = theta.j[1:2], sigma = diag(2)*omega1)
6   d4 <- dgamma(theta.proposed[3], shape = theta.j[3]/omega2, rate = 1/omega2)
7   out <- log(d1) + log(d2) - log(d3) - log(d4)
8   return(out)
9 }
```

METROPOLIS-HASTINGS FROM SCRATCH

Example 18.3: Setting up MCMC

```
1 ## Search window width
2 omega1 <- 0.5
3 omega2 <- 1/100
4
5 ## Track number accepted/rejected
6 accepted <- 0
7 rejected <- 0
8
9 ## Initialize based on a random draw from prior
10 theta.j <- as.numeric(c(rnorm(1, mean = 0, sd = sqrt(50)),
11                          rnorm(1, mean = 50, sd = sqrt(200)),
12                          rgamma(1, shape = 3, rate = 3/5)))
13
14 ## Number of MCMC iterations
15 M = 500000
16
17 ## Store accepted thetas
18 accepted.thetas <- matrix(0, nrow = M, ncol = 3)
```

Example 18.3: Running MCMC

```
1 ## Set seed for reproducibility
2 set.seed(7345)
3
4 ## Run MCMC
5 for (j in 1:M) {
6   theta.proposed <- proposal(theta.j, omega1, omega2)
7   ll.th.proposed <- loglik(theta.proposed, x, y)
8   ll.th.j <- loglik(theta.j, x, y)
9   prior.proposed <- pr(theta.proposed)
10  prior.theta.j <- pr(theta.j)
11  p.ratio <- prop.ratio(theta.j, theta.proposed, omega1, omega2)
12  logQ <- ll.th.proposed - ll.th.j + prior.proposed - prior.theta.j + p.ratio
13  U <- runif(1)
14  if (exp(logQ) > U) {
15    accepted <- accepted + 1
16    accepted.thetas[accepted,] <- theta.proposed
17    theta.j <- theta.proposed
18  }
19  if (exp(logQ) <= U) {rejected <- rejected + 1}
20  if (round(j/5000) == (j/5000)) {
21    print(paste(j, "iterations complete!"))
22  }
23 }
```

Exercise: Go through code to understand how it synthesizes the functions.

Example 18.3:

```
1 ## Report accepted/rejected draws
2 > accepted
3 [1] 23745
4
5 > rejected
6 [1] 76255
7
8 ## Don't forget that you have more storage than necessary
9 accepted.thetas <- accepted.thetas[1:accepted,]
```

More samples are rejected than accepted.

METROPOLIS-HASTINGS FROM SCRATCH

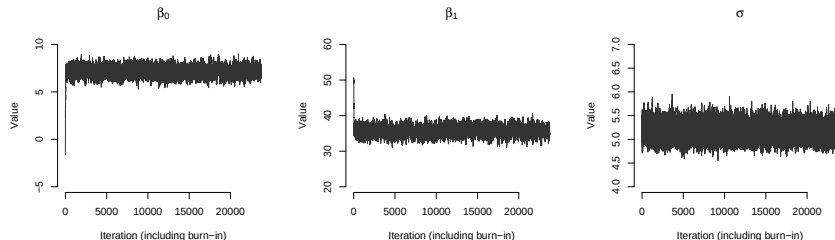


Figure S4.7: Trace plot (depicting burn-in).



METROPOLIS-HASTINGS FROM SCRATCH

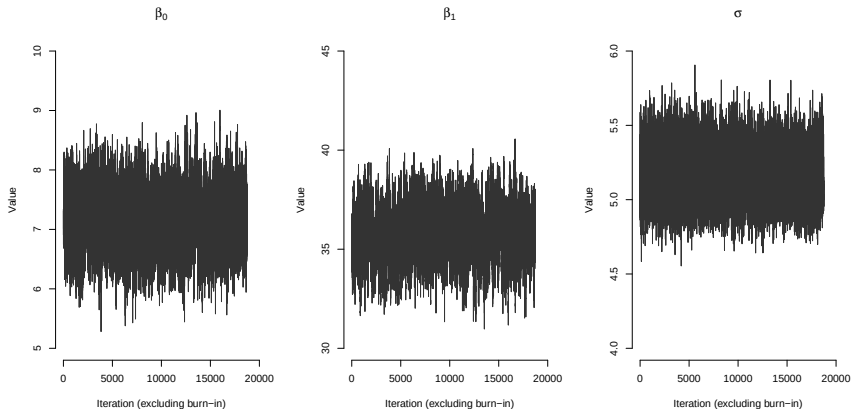


Figure S4.8: Trace plot (excluding 1:5000).



METROPOLIS-HASTINGS FROM SCRATCH

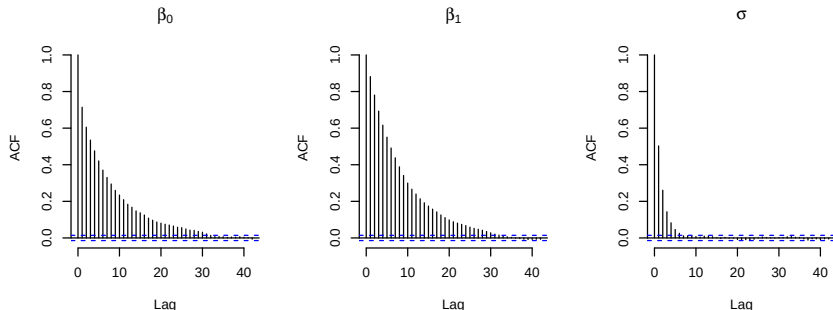


Figure S4.9: Assessing autocorrelation (`acf()` function in R).



METROPOLIS-HASTINGS FROM SCRATCH

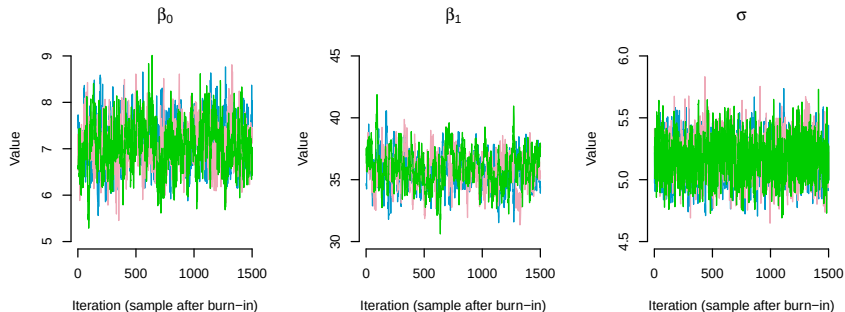


Figure S4.10: Multiple chains (done in practice; code not shown).



Example 18.3: Draws represent estimate of joint posterior

```
1 ## Combine three chains (code for three chains run but not shown):
2 final.thetas <- rbind(accepted.thetas, accepted.thetas2, accepted.thetas3)
3
4 > head(final.thetas)
5      [,1]      [,2]      [,3]
6 [1,] 6.747033 35.87393 5.264335
7 [2,] 7.342367 34.90827 5.588177
8 [3,] 7.083618 35.64153 5.348211
9 [4,] 6.907126 36.34663 5.192740
10 [5,] 7.067542 35.27238 5.159027
11 [6,] 7.309792 36.18202 5.129287
12
13 ## (etc.)
14
15 ## Column represents draws from marginal posterior.
16
17 ## Draws are correlated within an iteration
18 ## (not apparent from looking at marginal posteriors).
```

METROPOLIS-HASTINGS FROM SCRATCH

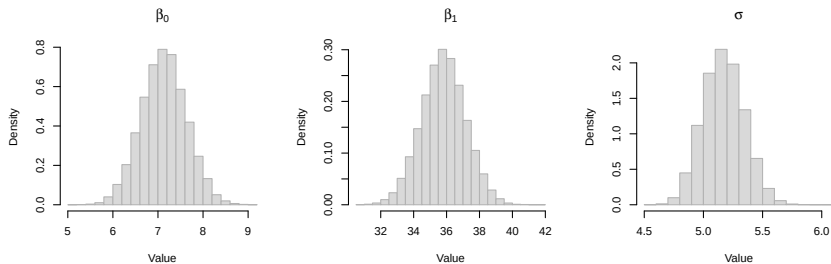


Figure S4.11: Marginal posteriors.



METROPOLIS-HASTINGS FROM SCRATCH

Example 18.3: Draws represent estimate of joint posterior

```
1 ## Posterior means
2 > colMeans(final.thetas)
3 [1]  7.131045 35.809624  5.170375
4
5 ## Posterior covariance
6 > cov(final.thetas)
7           [,1]      [,2]      [,3]
8 [1,]  2.495971e-01 -5.695972e-01 -8.547916e-05
9 [2,] -5.695972e-01  1.808847e+00  6.711686e-05
10 [3,] -8.547916e-05  6.711686e-05  2.984033e-02
11
12 ## OLS (comparison)
13 zz <- lm(y ~ x)
14
15 > vcov(zz)
16           (Intercept)          x
17 (Intercept)    0.23914 -0.58209
18 x             -0.58209  1.82210
```

Example 18.3: Results

```
1 ## Posterior mean
2 > as.numeric(mean(final.thetas[,2]))
3 [1] 35.80962
4
5 ## 2.5th and 97.5th percentiles
6 > as.numeric(quantile(final.thetas[,2], c(0.025, 0.975)))
7 [1] 33.14443 38.44808
8
9 ## Recall HPDint function
10 > as.numeric(HPDint(final.thetas[,2]))
11 [1] 33.21951 38.51218
12
13 ## 95% CI from OLS
14 > as.numeric(confint(zz)[2,])
15 [1] 32.949 38.254
```

Example 18.3: Discussion

- Indeed, in *this* particularly nice case, the interval estimates are quite concordant. But...
 - ▶ The sample size was not particularly small.
 - ▶ The normal likelihood was satisfied.
 - ▶ No interim analyses on the data (affects frequentist properties).
- Play around with these, and you'll start to see differences emerge.
 - ▶ As this is an introduction, we're starting with simple examples and they will behave accordingly.

Example 18.3: Posterior predictive

```
1 ## Number of final accepted iterations
2 n.mc <- dim(final.thetas)[1]
3
4 ## Sample draws (equal to sample size from data)
5 samp <- sample(c(1:n.mc), size = n, replace = TRUE)
6
7 ## Posterior predictive draw
8 y.tildes <- rnorm(n, mean = final.thetas[samp,1] + final.thetas[samp,2]*x,
9                   sd=final.thetas[samp,3])
```


METROPOLIS-HASTINGS FROM SCRATCH

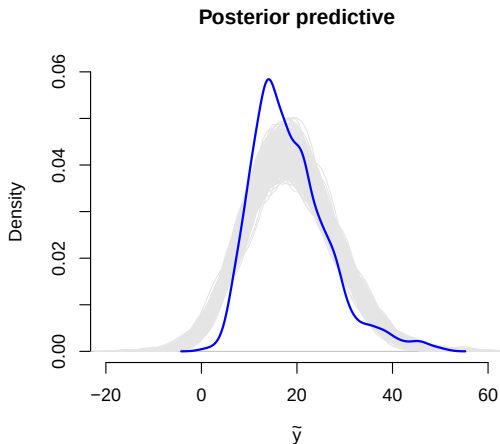


Figure S4.12: Posterior predictive.



METROPOLIS-HASTINGS FROM SCRATCH

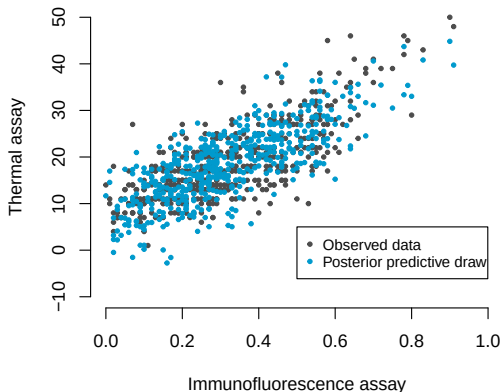


Figure S4.13: Posterior predictive draws (conditional on X).



METROPOLIS-HASTINGS FROM SCRATCH

Example 18.3: Sub-populations

- Having draws from the joint posterior distribution is quite helpful.
- Forming a credible interval for a subpopulation-specific mean (say, $X = 0.1$) is incomprehensibly easy.

```
1 ## Credible interval for E[Y|X = 0.1]
2 submean.draws <- final.thetas[,1] + 0.1*final.thetas[,2]
3
4 > as.numeric(quantile(submean.draws, c(0.025, 0.975)))
5 [1] 9.948057 11.478941
6
7 ## Analogous CI from OLS
8 C <- matrix(c(1, 0.1), nrow = 1)
9 submean.est <- C %*% coef(zz)
10 submean.se <- sqrt(C %*% vcov(zz) %*% t(C))
11
12 > c(submean.est - qnorm(0.975)*submean.se,
13     submean.est + qnorm(0.975)*submean.se)
14 [1] 10.029 11.501
```

- This generalizes in the way you might expect. For instance, in logistic regression, you can get a posterior for an odds ratio by exponentiating the posterior draws for a difference in log-odds.