

# Bayesian Inference

Introduction and Ideas

**Abhinavan, Anant, Prateek, Mukesh**

Indian Statistical Institute

May 2026

## What is Bayesian Inference?

---

There are two main schools of thought regarding the approach to statistical inference (also called learning). We are familiar to an extent with frequentist methods. The key difference lies in the interpretation of “probability” in the two schools of thought.

- **Frequentist Inference.** Probabilities are interpreted as long-run frequencies. The goal is to create procedures with long-run frequency guarantees.
- **Bayesian Inference.** Probabilities are interpreted as subjective degrees of belief. The goal is to state and analyze one’s beliefs.

Essentially, the Bayesian approach requires us to drop the idea of probability as a long-run frequency and instead think of it in a subjective sense, conditioned on known facts.

Following from this philosophy, the parameter  $\theta$  of any probability distribution is taken to be *fixed* in the frequentist approach. On the contrary, in the Bayesian approach  $\theta$  is treated as a *random variable* with its own distribution — chosen subjectively based on prior knowledge — called the **prior**.

## The Prior and the Posterior

---

The prior density is represented by  $\pi(\theta)$ .

Given a sample  $X_1, X_2, \dots, X_n$ , we wish to find a new density for  $\theta$  that incorporates the information contained in the observations. This new density, conditioned on the data, is known as the **posterior** and is denoted

$$p(\theta \mid X_1, X_2, \dots, X_n).$$

## Computing the Posterior via Bayes’ Theorem

We recall Bayes’ theorem for probability densities:

$$f_X(x \mid y) = \frac{f_Y(y \mid x) f_X(x)}{f_Y(y)} = \frac{f_Y(y \mid x) f_X(x)}{\int_{-\infty}^{\infty} f_Y(y \mid x) f_X(x) dx}. \quad (1)$$

Setting  $x = \theta$  and  $y = \mathbf{x} = (x_1, \dots, x_n)$ , we obtain:

$$f_{\theta|\mathbf{X}}(\theta \mid \mathbf{x}) = \frac{f_{\mathbf{X}|\theta}(\mathbf{x} \mid \theta) f_{\theta}(\theta)}{f_{\mathbf{X}}(\mathbf{x})}. \quad (2)$$

Renaming each factor by its role in the Bayesian framework:

$$\underbrace{p(\theta \mid \mathbf{X})}_{\text{posterior}} = \frac{\underbrace{p(\mathbf{X} \mid \theta)}_{\text{likelihood}} \underbrace{\pi(\theta)}_{\text{prior}}}{\underbrace{p(\mathbf{X})}_{\text{evidence}}}. \quad (3)$$

This can equivalently be written as:

$$p(\theta \mid X_1, X_2, \dots, X_n) = \frac{p(X_1, X_2, \dots, X_n \mid \theta) \pi(\theta)}{p(X_1, X_2, \dots, X_n)} = \frac{L_n(\theta) \pi(\theta)}{c_n} \propto L_n(\theta) \pi(\theta), \quad (4)$$

where  $L_n(\theta)$  is the familiar likelihood function and the denominator  $c_n = p(X_1, \dots, X_n)$  is called the **evidence**. Since the evidence does not depend on  $\theta$ , the posterior is proportional to the product of the likelihood and the prior.

## Inference Using the Posterior

---

The posterior  $p(\theta \mid X_1, \dots, X_n)$  is the central object used for inference in the Bayesian approach.

### Posterior Mean

A natural point estimate of  $\theta$  is the posterior mean:

$$\bar{\theta} = \int \theta p(\theta \mid X_1, X_2, \dots, X_n) d\theta = \frac{\int \theta L_n(\theta) \pi(\theta) d\theta}{\int L_n(\theta) \pi(\theta) d\theta}. \quad (5)$$

## Interval Estimation: Frequentist vs. Bayesian

---

To illustrate the difference between the two approaches, consider interval estimation. Suppose

$$X_1, X_2, \dots, X_n \sim N(\theta, 1),$$

and we wish to construct an interval estimate for  $\theta$ .

### Frequentist Confidence Interval

The classical  $(1 - \alpha)$  confidence interval is:

$$C_1 = \left[ \bar{X}_n - \frac{\tau_{\alpha/2}}{\sqrt{n}}, \bar{X}_n + \frac{\tau_{\alpha/2}}{\sqrt{n}} \right], \quad (6)$$

where  $\tau_\alpha$  denotes the upper  $\alpha$ -quantile of the  $N(0, 1)$  distribution.

### Bayesian Credible Interval

In the Bayesian approach, we seek an interval  $C_2 = [a, b]$  such that:

$$\int_{C_2} p(\theta \mid X_1, X_2, \dots, X_n) d\theta = 1 - \alpha, \quad (7)$$

which is equivalently specified by choosing  $a$  and  $b$  satisfying:

$$\int_{-\infty}^a p(\theta \mid X_1, \dots, X_n) d\theta = \int_b^{\infty} p(\theta \mid X_1, \dots, X_n) d\theta = \frac{\alpha}{2}. \quad (8)$$

This interval  $C_2$  is known as an **equal-tailed credible interval** (or Bayesian credible set). Unlike a frequentist confidence interval, a credible interval carries the direct probabilistic interpretation: given the observed data,  $\theta$  lies in  $C_2$  with posterior probability  $1 - \alpha$ .

## Conjugate Priors: Definition and Mechanics

---

A family of prior distributions  $\mathcal{P}$  is said to be **conjugate** to a family of sampling distributions  $\mathcal{F}$  if, for every prior  $\pi \in \mathcal{P}$ , the resulting posterior  $p(\cdot | x)$  also belongs to  $\mathcal{P}$ . Essentially, the parametric family of the distribution remains closed under sampling.

This property allows for the instant calculation of posterior distributions. Instead of performing complex integrations, one simply updates the hyperparameters of the distribution.

## Conjugacy in the Exponential Family

---

Most standard conjugate pairs arise from the exponential family. Suppose the sampling model takes the form:

$$p(x | \theta) = \exp(\theta^T x - A(\theta)). \quad (1)$$

The conjugate prior for this model is defined as:

$$\pi_{x_0, n_0}(\theta) \propto \exp(n_0 x_0^T \theta - n_0 A(\theta)). \quad (2)$$

Here,  $n_0$  can be interpreted as the number of “virtual observations” and  $x_0$  as their average. After observing  $n$  real data points with mean  $\bar{X}$ , the posterior hyperparameters update as:

- $n'_0 = n + n_0$
- $x'_0 = \frac{n\bar{X} + n_0 x_0}{n + n_0}$

### Example: Poisson–Gamma Hyperparameter Update

Consider a Poisson model  $X \sim \text{Poisson}(\lambda)$  with a  $\text{Gamma}(\alpha, \beta)$  prior. Suppose our prior belief is based on  $n_0 = 5$  virtual observations with a virtual average  $x_0 = 2$ .

1. We observe  $n = 10$  new data points with sample mean  $\bar{X} = 4$ .
2. The updated virtual sample size is:  $n'_0 = 10 + 5 = 15$ .
3. The updated mean is:  $x'_0 = \frac{(10 \times 4) + (5 \times 2)}{15} = \frac{50}{15} \approx 3.33$ .

The posterior is then a Gamma distribution defined by these updated parameters.

### Example: The Poisson–Gamma Conjugate Pair

Suppose we are counting events, so each observation follows  $X_i \sim \text{Poisson}(\lambda)$ .

**The Likelihood:** For  $n$  observations, the joint likelihood is proportional to:

$$\mathcal{L}(\lambda) \propto \lambda^{\sum X_i} e^{-n\lambda}. \quad (3)$$

**The Conjugate Prior:** We choose the Gamma distribution for  $\lambda$ . With  $\lambda \sim \text{Gamma}(\alpha, \beta)$ :

$$\pi(\lambda) \propto \lambda^{\alpha-1} e^{-\beta\lambda}. \quad (4)$$

Observe that the likelihood and the prior share the exact same algebraic structure:  $\lambda^{(\cdot)} \times e^{(\cdot)\lambda}$ . Multiplying them together:

$$p(\lambda \mid X_1, \dots, X_n) \propto (\lambda^{\sum X_i} e^{-n\lambda}) \times (\lambda^{\alpha-1} e^{-\beta\lambda}). \quad (5)$$

Combining the exponents:

$$p(\lambda \mid X_1, \dots, X_n) \propto \lambda^{(\sum X_i + \alpha) - 1} e^{-(\beta + n)\lambda}. \quad (6)$$

This is exactly the kernel of a new Gamma distribution — no complex integration required:

$$\lambda \mid X_1, \dots, X_n \sim \text{Gamma}(\alpha + n\bar{X}, \beta + n). \quad (7)$$

## Technical Foundations and Proofs

---

### The Linearity of Posterior Means

A significant theoretical result by Diaconis (1979) states that in the continuous case, the posterior mean of the gradient of the log-normalising constant  $\nabla A(\theta)$  is linear if and only if the prior is conjugate. Specifically:

$$\mathbb{E}(\nabla A(\theta) \mid X) = aX + b. \quad (8)$$

This linearity ensures that the posterior mean is always a weighted average of the prior information and the observed data.

### Non-Exponential Family: Uniform/Pareto

While conjugacy is most common in the exponential family, it exists elsewhere. For the model  $X \sim \text{Uniform}(0, \theta)$ , the conjugate prior is the Pareto distribution. The posterior updates based on the maximum observed value  $X_{(n)}$ :

$$\theta \mid X_1, \dots, X_n \sim \text{Pareto}(n + k, \max\{X_{(n)}, \nu_0\}). \quad (9)$$

## Summary of Key Conjugate Pairs

---

The following table summarises the most important conjugate relationships used in statistical machine learning.

| Sampling Model             | Conjugate Prior               | Posterior Distribution                              |
|----------------------------|-------------------------------|-----------------------------------------------------|
| Bernoulli( $\theta$ )      | Beta( $\alpha, \beta$ )       | Beta( $\alpha + n\bar{X}, \beta + n(1 - \bar{X})$ ) |
| Poisson( $\lambda$ )       | Gamma( $\alpha, \beta$ )      | Gamma( $\alpha + n\bar{X}, \beta + n$ )             |
| Normal (known $\sigma^2$ ) | Normal( $\mu_0, \sigma_0^2$ ) | Normal (weighted average)                           |
| Normal (known $\mu$ )      | Inverse-Gamma                 | Inverse-Gamma                                       |
| Multinomial( $\theta$ )    | Dirichlet( $\alpha$ )         | Dirichlet( $\alpha + n\bar{X}$ )                    |
| Uniform( $0, \theta$ )     | Pareto( $\nu_0, k$ )          | Pareto( $\max\{\nu_0, X_{(n)}\}, n + k$ )           |

## Conjugate Priors in Machine Learning

---

Conjugate priors are not merely mathematical conveniences; they underpin classical machine learning algorithms.

- **Text Classification (Spam Filters):** Naïve Bayes classifiers often use the Dirichlet–Multinomial conjugate pair. The model counts word occurrences (Multinomial); the Dirichlet prior acts as Laplace smoothing, adding virtual word counts so the model does not break on unseen vocabulary. The posterior update is fast enough to allow real-time learning.
- **Topic Modelling (LDA):** Latent Dirichlet Allocation relies on the same Dirichlet–Multinomial conjugacy to infer document-topic and topic-word distributions at scale.
- In Big Data settings, running full MCMC simulations for every new data point is computationally infeasible. Conjugate updates provide the closed-form arithmetic required for real-time inference.

## The Limitations of Conjugacy

---

If conjugate priors are so easy to compute with, why do we not always use them?

- **Inflexibility:** The prior must be chosen for mathematical convenience rather than to faithfully reflect genuine subjective beliefs.
- **Complex Models:** In modern hierarchical models or deep learning, conjugate priors simply do not exist for the intricate relationships we wish to model.
- **The Solution:** When we step outside the bounds of conjugacy, the posterior integral becomes analytically intractable. This is precisely the motivation for computational simulation methods such as MCMC.

## Conclusion

---

Conjugate priors serve as a mathematically convenient tool that bypasses the need for high-dimensional integration. By providing closed-form posterior updates, they power modern applications such as Topic Modelling (LDA) and Naïve Bayes classifiers, where computational efficiency is paramount. Understanding their theoretical foundations — and their limitations — is essential for knowing when to use them and when to turn to simulation-based methods instead.

## Example 1: Bernoulli/Beta Conjugate Model

---

In the Bernoulli/Beta conjugate model, we observe  $n$  Bernoulli trials with  $S_n$  successes. Taking a Uniform prior on  $\theta$ , the posterior is  $\text{Beta}(S_n + 1, n - S_n + 1)$ . The following R code constructs and plots the prior, scaled likelihood, and posterior for  $n = 15$ ,  $S_n = 7$ .

```
1 n <- 15
2 S_n <- 7
3 alpha <- S_n + 1 # 8
4 beta <- n - S_n + 1 # 9
5
6 # Create theta values (grid from 0 to 1)
7 theta <- seq(0, 1, length.out = 1000)
8
9 # 1. Prior (Uniform)
10 prior <- dunif(theta, min = 0, max = 1)
11
12 # 2. Likelihood (scaled to maximum = 1)
13 likelihood <- theta^S_n * (1 - theta)^(n - S_n)
14 likelihood_scaled <- likelihood / max(likelihood)
15
16 # 3. Posterior (Beta(8, 9))
17 posterior <- dbeta(theta, shape1 = alpha, shape2 = beta)
18
19 # Create the plot
20 plot(theta, posterior, type = "l",
21       col = "red", lwd = 3,
22       xlab = expression(theta),
23       ylab = "Density / Scaled Value",
24       main = paste0("Prior, Likelihood (Scaled), and Posterior for n = ",
25                     n,
26                     ", S_n = ", S_n, "\nPosterior: Beta(", alpha, ", ",
27                     beta, ")"),
28       ylim = c(0, max(posterior, prior, likelihood_scaled) * 1.1))
29
30 # Add prior (uniform)
31 lines(theta, prior, col = "blue", lwd = 3, lty = 2)
32
33 # Add scaled likelihood
34 lines(theta, likelihood_scaled, col = "green", lwd = 3, lty = 3)
35
36 # Add legend
37 legend("topright",
38       legend = c("Prior (Uniform)",
39                 "Likelihood (scaled)",
40                 paste0("Posterior: Beta(", alpha, ", ", beta, ")")),
41       col = c("blue", "green", "red"),
42       lty = c(2, 3, 1),
43       lwd = 3,
44       bty = "n",
45       bg = "white")
```

## Example 2: Normal/Normal Conjugate Model

---

In the Normal/Normal conjugate model, the prior is  $\theta \sim N(a, b^2)$  and the likelihood comes from  $X_i | \theta \sim N(\theta, \sigma^2)$ . The posterior is also Normal, with mean a weighted combination

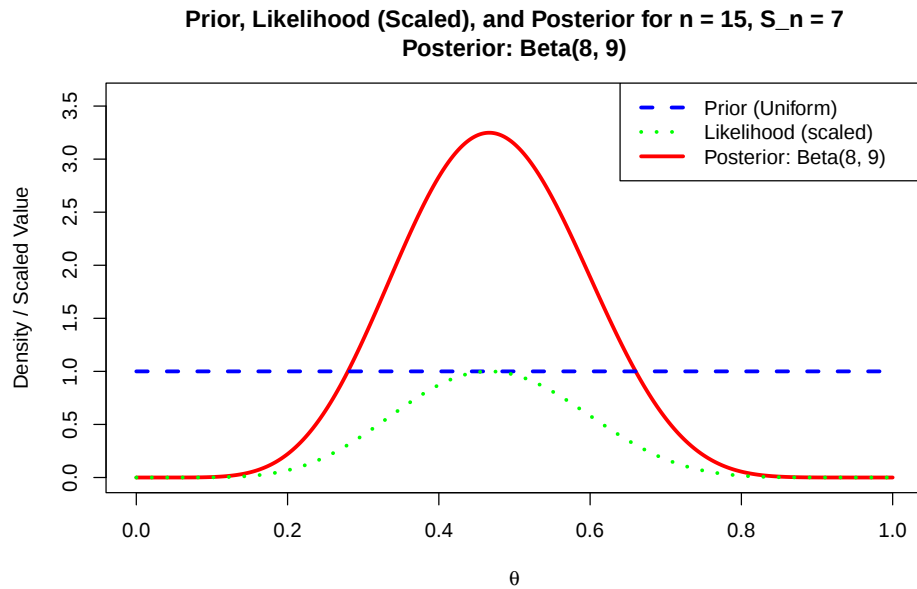


Figure 1: Bernoulli/Beta Conjugate Model: Prior (Uniform), Likelihood (scaled), and Posterior (Beta(8, 9)).

of the prior mean and the sample mean, and a variance that shrinks with increasing sample size. The following R code computes and plots the three densities for  $n = 10$ ,  $\bar{X} = 0.5$ ,  $\sigma = 1$ ,  $a = 0$ ,  $b = 1$ .

```

1 n <- 10
2 Xbar <- 0.5
3 sigma <- 1
4 a <- 0
5 b <- 1
6 w <- (n / sigma^2) / (n / sigma^2 + 1 / b^2)
7 theta_bar <- w * Xbar + (1 - w) * a
8 tau <- sqrt(1 / (n / sigma^2 + 1 / b^2))
9 cat("Posterior mean (theta_bar):", theta_bar, "\n")
10 cat("Posterior standard deviation (tau):", tau, "\n")
11 cat("Weight (w):", w, "\n")
12
13 theta <- seq(-3, 3, length.out = 1000)
14
15 # 1. Prior: N(a, b^2)
16 prior <- dnorm(theta, mean = a, sd = b)
17
18 # 2. Likelihood: proportional to N(Xbar, sigma^2/n)
19 likelihood <- dnorm(theta, mean = Xbar, sd = sigma / sqrt(n))
20 likelihood_scaled <- likelihood / max(likelihood)
21
22 # 3. Posterior: N(theta_bar, tau^2)
23 posterior <- dnorm(theta, mean = theta_bar, sd = tau)
24
25 # Create the plot
26 plot(theta, posterior, type = "l",
27       col = "red", lwd = 3,
28       xlab = expression(theta),
29       ylab = "Density / Scaled Value",

```

```

30 main = paste0("Prior, Likelihood (Scaled), and Posterior\n",
31               "Normal/Normal Model (n = ", n, ", Xbar = ", Xbar,
32               ", sigma = ", sigma, ")"),
33   ylim = c(0, max(posterior, prior, likelihood_scaled) * 1.1))
34 lines(theta, prior, col = "blue", lwd = 3, lty = 2)
35 lines(theta, likelihood_scaled, col = "green", lwd = 3, lty = 3)
36 legend("topright",
37       legend = c(paste0("Prior: N(", a, ", ", b^2, ")"),
38                 "Likelihood (scaled)",
39                 paste0("Posterior: N(", round(theta_bar, 3), ", ",
40                 round(tau^2, 3), ")")),
41       col = c("blue", "green", "red"),
42       lty = c(2, 3, 1),
43       lwd = 3,
44       bty = "o",
45       bg = "white")

```

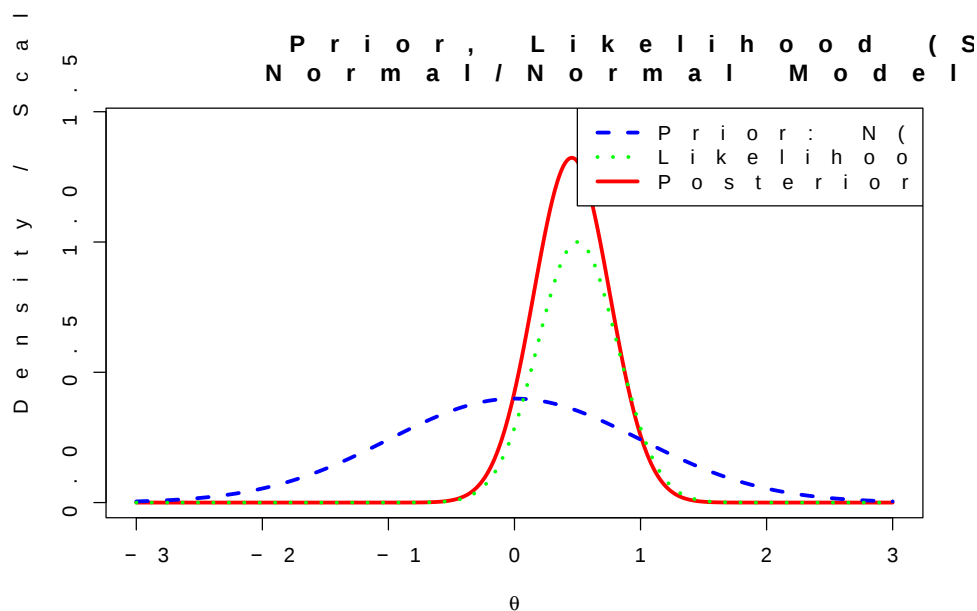


Figure 2: Normal/Normal Conjugate Model: Prior ( $\theta \sim N(a, b^2)$ ), Likelihood (scaled), and Posterior (Normal distribution).

## Inference about Functions of Parameters

### Definition

A common inferential goal is to reason about a derived quantity  $\tau = g(\theta)$ , where  $g$  is a known (possibly nonlinear) function. The posterior distribution of  $\tau$  is characterised by its CDF:

$$H(t \mid D_n) = \int_A p(\theta \mid D_n) d\theta, \quad A = \{\theta : g(\theta) \leq t\}.$$

When  $g$  is monotone, the posterior density of  $\tau$  follows by the change-of-variables formula:

$$p(\tau \mid D_n) = p(g^{-1}(\tau) \mid D_n) \cdot \left| \frac{d\theta}{d\tau} \right|.$$

## Simulation Approach

Analytical computation of  $H(t \mid D_n)$  is often intractable. A universal simulation-based algorithm proceeds as follows:

1. Draw  $\theta_1, \dots, \theta_B \sim p(\theta \mid D_n)$ .
2. Transform:  $\tau_b = g(\theta_b)$  for  $b = 1, \dots, B$ .
3. The empirical distribution of  $\{\tau_b\}$  approximates  $p(\tau \mid D_n)$ .

This approach works for *any* function  $g$ ; credible intervals are obtained directly as empirical quantiles of  $\{\tau_b\}$ , and it extends naturally to multiparameter settings.

## Multiparameter Problems

---

### General Framework

Let  $\theta = (\theta_1, \dots, \theta_d)^T$  with prior  $\pi(\theta)$ . The joint posterior satisfies

$$p(\theta \mid D_n) \propto L_n(\theta) \pi(\theta).$$

Interest often centres on a single component, say  $\theta_1$ . Its marginal posterior requires integrating out all other parameters:

$$p(\theta_1 \mid D_n) = \int \cdots \int p(\theta_1, \dots, \theta_d \mid D_n) d\theta_2 \cdots d\theta_d,$$

which is generally intractable in closed form. The simulation solution avoids this entirely:

1. Draw  $\theta^{(1)}, \dots, \theta^{(B)} \sim p(\theta \mid D_n)$ .
2. Collect the first components  $\theta_1^{(1)}, \dots, \theta_1^{(B)}$ .
3. These form a sample from  $p(\theta_1 \mid D_n)$  — no numerical integration required.

### Example: Treatment Effect $\tau = \theta_2 - \theta_1$

Consider two independent Bernoulli experiments: Group 1 (control) with  $n_1$  trials and  $X_1$  successes, and Group 2 (treatment) with  $n_2$  trials and  $X_2$  successes. Under independent Uniform priors, the posteriors are:

$$\theta_1 \mid X_1 \sim \text{Beta}(X_1 + 1, n_1 - X_1 + 1), \quad \theta_2 \mid X_2 \sim \text{Beta}(X_2 + 1, n_2 - X_2 + 1).$$

The treatment effect  $\tau = \theta_2 - \theta_1$  has no simple closed-form posterior, so we proceed by simulation:

1. Draw  $\theta_1^{(b)} \sim \text{Beta}(X_1 + 1, n_1 - X_1 + 1)$  for  $b = 1, \dots, B$ .
2. Draw  $\theta_2^{(b)} \sim \text{Beta}(X_2 + 1, n_2 - X_2 + 1)$  independently.
3. Compute  $\tau^{(b)} = \theta_2^{(b)} - \theta_1^{(b)}$ .
4.  $\{\tau^{(1)}, \dots, \tau^{(B)}\}$  is a sample from  $p(\tau \mid X_1, X_2)$ .

## R Implementation

```
1 # Group 1: Control group
2 n1 <- 40
3 X1 <- 25
4 # Group 2: Treatment group
5 n2 <- 40
6 X2 <- 30
7 B <- 10000
8
9 theta1 <- rbeta(B, shape1 = X1 + 1, shape2 = n1 - X1 + 1)
10 theta2 <- rbeta(B, shape1 = X2 + 1, shape2 = n2 - X2 + 1)
11 tau <- theta2 - theta1
12
13 post_mean <- mean(tau)
14 cred_interval <- quantile(tau, probs = c(0.025, 0.975))
15 post_median <- median(tau)
16 post_sd <- sd(tau)
17
18 cat("=====\n")
19 cat("Treatment Effect (tau = theta2 - theta1)\n")
20 cat("=====\n")
21 cat("Group 1 (Control):", n1, "trials,", X1, "successes\n")
22 cat("Group 2 (Treatment):", n2, "trials,", X2, "successes\n")
23 cat("-----\n")
24 cat("Posterior Mean      :", round(post_mean, 4), "\n")
25 cat("Posterior Median    :", round(post_median, 4), "\n")
26 cat("Posterior SD        :", round(post_sd, 4), "\n")
27 cat("95% Credible Interval: [", round(cred_interval[1], 4), ",",
28     round(cred_interval[2], 4), "]\n")
29 cat("-----\n")
30 if (cred_interval[1] < 0 && cred_interval[2] > 0) {
31   cat("Conclusion: 0 is INSIDE the 95% credible interval\n")
32   cat("No significant treatment effect (at 95% level)\n")
33 } else if (cred_interval[2] < 0) {
34   cat("Conclusion: 0 is BELOW the 95% credible interval\n")
35   cat("Treatment has significant NEGATIVE effect\n")
36 } else {
37   cat("Conclusion: 0 is ABOVE the 95% credible interval\n")
38   cat("Treatment has significant POSITIVE effect\n")
39 }
40 cat("=====\n\n")
41
42 par(mfrow = c(2, 2))
43
44 # Plot 1: Histogram of tau with 95% CI
45 hist(tau,
46      breaks = 50,
47      col = "lightblue",
48      main = "Posterior of Treatment Effect\n(tau = theta2 - theta1)",
49      xlab = expression(tau),
50      ylab = "Frequency",
51      border = "white",
52      probability = TRUE)
53 lines(density(tau), col = "darkblue", lwd = 2)
54 abline(v = 0, col = "red", lwd = 3, lty = 2)
55 abline(v = cred_interval[1], col = "darkgreen", lwd = 2, lty = 3)
56 abline(v = cred_interval[2], col = "darkgreen", lwd = 2, lty = 3)
57 abline(v = post_mean, col = "purple", lwd = 2, lty = 4)
```

```

58 legend("topright",
59       legend = c("Posterior Density", "0 (No effect)", "95% CI", "
        Posterior Mean"),
60       col = c("darkblue", "red", "darkgreen", "purple"),
61       lty = c(1, 2, 3, 4),
62       lwd = 2,
63       bty = "n",
64       cex = 0.8)
65
66 # Plot 2: Density of tau with shaded 95% CI
67 plot(density(tau),
68      main = "Posterior Density with 95% Credible Interval",
69      xlab = expression(tau),
70      ylab = "Density",
71      col = "darkblue",
72      lwd = 2)
73 dens <- density(tau)
74 x_shade <- dens$x[dens$x >= cred_interval[1] & dens$x <= cred_interval
75 [2]]
76 y_shade <- dens$y[dens$x >= cred_interval[1] & dens$x <= cred_interval
77 [2]]
78 polygon(c(x_shade[1], x_shade, x_shade[length(x_shade)]),
79        c(0, y_shade, 0),
80        col = rgb(0, 0, 1, 0.3), border = NA)
81 abline(v = 0, col = "red", lwd = 2, lty = 2)
82 abline(v = post_mean, col = "purple", lwd = 2, lty = 4)
83
84 # Plot 3: Trace plots of theta1 and theta2 (first 1000 samples)
85 plot(theta1[1:1000], type = "l",
86      col = "blue",
87      main = "Trace Plot (First 1000 samples)",
88      xlab = "Iteration",
89      ylab = expression(theta),
90      ylim = range(c(theta1[1:1000], theta2[1:1000])))
91 lines(theta2[1:1000], col = "red")
92 legend("topright",
93       legend = c(expression(theta[1]), expression(theta[2])),
94       col = c("blue", "red"),
95       lty = 1,
96       lwd = 2,
97       bty = "n")
98
99 # Plot 4: Scatter plot of theta2 vs theta1
100 plot(theta1, theta2,
101      col = rgb(0, 0, 0, 0.1),
102      pch = 16,
103      main = "Joint Posterior Distribution",
104      xlab = expression(theta[1] * " (Control)"),
105      ylab = expression(theta[2] * " (Treatment)"))
106 abline(0, 1, col = "red", lwd = 2, lty = 2)
107 cor_val <- cor(theta1, theta2)
108 text(max(theta1) * 0.7, max(theta2) * 0.9,
109      paste("Correlation =", round(cor_val, 3)),
110      cex = 0.9)
111
112 par(mfrow = c(1, 1))

```

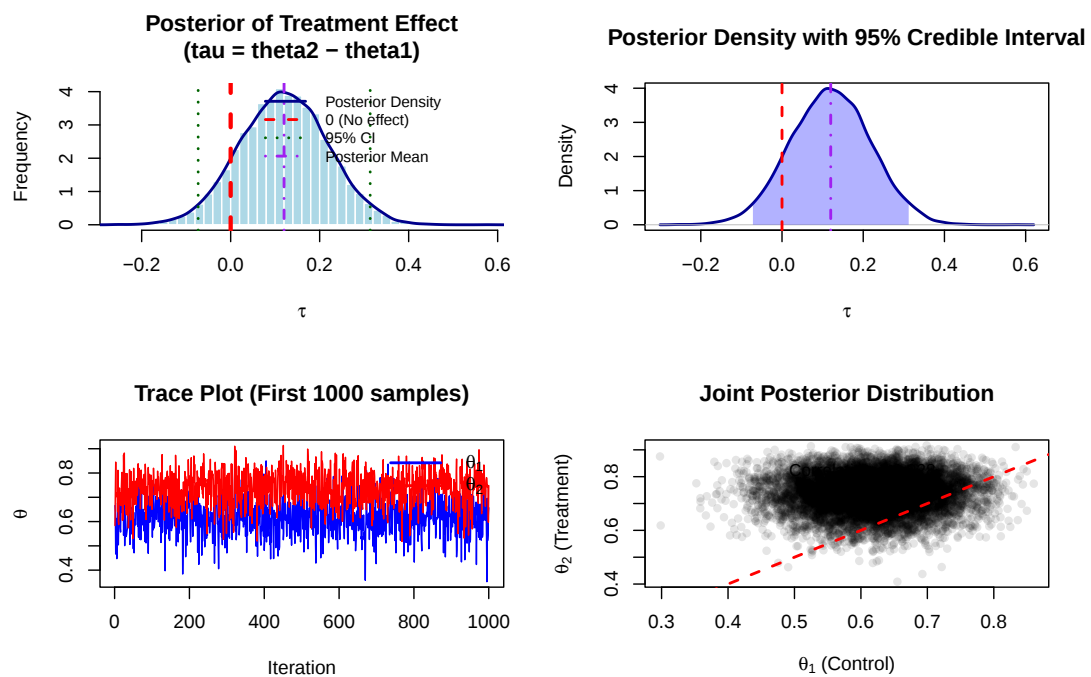


Figure 3: Multiparameter Example: Posterior analysis of the treatment effect  $\tau = \theta_2 - \theta_1$ , showing the posterior histogram with 95% credible interval, shaded density, trace plots, and joint posterior distribution.

## Frequentist vs. Bayesian Testing

Frequentist inference and Bayesian inference are two fundamentally different processes. The core difference lies in how the data are used to perform inference.

| Frequentist Testing                                                | Bayesian Testing                                           |
|--------------------------------------------------------------------|------------------------------------------------------------|
| Probability based on repeated sampling                             | Probability interpreted as degree of belief                |
| “How unusual is the observed data under $H_0$ ?”                   | “How probable is the hypothesis after observing the data?” |
| Probability assigned to procedures ( $p$ -value, rejection region) | Probability assigned directly to hypotheses                |

**Our Goal:** Study Bayesian testing, Bayes factors, and the BIC.

## Development of Bayesian Testing

We test  $H_0 : \theta = \theta_0$  vs.  $H_1 : \theta \neq \theta_0$  with prior belief  $P(H_0) = p_0$ . Under  $H_1$ ,  $\theta$  has a prior density  $\pi(\theta)$ . After observing data  $D_n$ , we apply Bayes' Theorem:

$$P(H_0 | D_n) = \frac{p(D_n | H_0) P(H_0)}{p(D_n | H_0) P(H_0) + p(D_n | H_1) P(H_1)}.$$

Assuming equal prior beliefs  $P(H_0) = P(H_1) = 1/2$ , this simplifies to:

$$P(H_0 | D_n) = \frac{p(D_n | \theta_0)}{p(D_n | \theta_0) + \int p(D_n | \theta) \pi(\theta) d\theta}. \quad (1)$$

This yields a direct probability statement about the truth of  $H_0$ .

## Interpretation and Marginal Likelihood

The numerator  $p(D_n | \theta_0)$  measures how well the null hypothesis explains the observed data. The integral in the denominator,

$$\int p(D_n | \theta) \pi(\theta) d\theta,$$

represents the *average* explanatory power of all parameter values under the alternative hypothesis. This leads to the central question of Bayesian testing:

“How well does the null explain the data compared to the entire alternative model?”

The quantity  $P(D_n | M) = \int L(\theta) \pi(\theta) d\theta$  is called the **Marginal Likelihood**.

## Balance: Fit vs. Complexity

---

The Bayesian Information Criterion (BIC) approximates twice the log marginal likelihood and embodies a trade-off between two competing forces:

$$\text{BIC} = \underbrace{\log L(\hat{\theta})}_{\text{Model Fit}} - \underbrace{\frac{d}{2} \log n}_{\text{Complexity Penalty}}.$$

- **Model Fit:** Measures how well the model explains the data.
- **Complexity Penalty:** Penalizes the model for each additional parameter. Larger models have larger parameter spaces and greater “uncertainty volume.”

Between two models, the one with the *higher* BIC is preferred.

## Summary and Limitations

---

- **Computational Problem:** Exact marginal likelihoods are difficult to calculate in high dimensions.
- **Asymptotics:** BIC is only accurate as  $n \rightarrow \infty$ . In small samples, the constant terms ignored in the BIC derivation may matter.
- **Prior Sensitivity:** Bayesian testing results can change significantly based on the choice of  $\pi(\theta)$ .