

Bayesian Statistics for Computer intensive methods sf2955

Timo Koski

April 29, 2009



- (a) Statistical Inference
- (b) Parametric statistical model
- (c) Bayes rule, Posterior Distributions
- (d) Bayes' Billiard Balls
- (e) Predictive Distributions
- (f) Asymptotic form of the posterior density



Learning/inference from data

By learning/inference from data one often means the process of inferring a general law or principle from the observations of particular instances. The general law is a piece of knowledge about the mechanism of nature that generates the data.



The intended learning is done by use of 'MODELS', which serve as the language in which the constraints predicated on the data can be described. We deal here with parametric statistical models.



Parametric statistical model

x is an observation of a random variable (X).

$$x \in \mathcal{X}$$

$f(x|\theta)$ is a probability density on $\mathcal{X}$. $f(x|\theta)$ is a known function of x and θ .

θ is an unknown parameter.



Parametric statistical model

x can be of the form $x^{(n)} = (x_1, \dots, x_n)$. x can be continuous or discrete variate or a mix thereof.

θ is an unknown parameter $\in \Theta =$ a vector space of finite dimension.
Hence we exclude, e.g., non-parametric statistics.



$$x \in f(x|\theta)$$

- x is distributed according to f ,
- x is an observation from the distribution f .

An outcome x of a random variable (r.v.) X .

Parametric statistical model: Examples; Normal distribution

$$\theta = (\mu, \sigma^2) \in \Theta = \mathbb{R} \times (0, \infty) .$$

$$f(x|\theta) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2\sigma^2}(x-\mu)^2}, -\infty < x < \infty.$$

We say that x is an observation from the normal distribution $N(\mu, \sigma^2)$.



Parametric statistical model: Examples; Bernoulli Distribution

Consider r.v. X with values $0, 1$, $0 < \theta < 1$ and

$$f(x | \theta) = \begin{matrix} x = 1 & x = 0 \\ \theta & 1 - \theta \end{matrix}$$

then we say x is distributed according to the Bernoulli distribution with the parameter θ .

$$X = x \in \text{Be}(\theta),$$



Parametric statistical model

$f(x|\theta)$ is a probabilistic mechanism of generating data, characterizes the behaviour of future observations conditional on θ .



Learning/inference from data: Inversion

Retrieve the parameters of the probabilistic generating mechanism using x .
 $f(x|\theta)$ is a probabilistic generating mechanism of data, characterizes the behaviour of future observations conditional on θ , but in inference the roles of x and θ are inverted.



Inversion and Bayes' Rule

Since

$$p(A | E) \cdot p(E) = p(E | A) \cdot p(A)$$

we have in a formal way Bayes' Rule of inversion

$$p(A | E) = \frac{p(E | A) \cdot p(A)}{p(E)}.$$

$$p(E) = p(E | A)p(A) + p(E | A^c)p(A^c).$$

A^c is the complement set of A .



Bayes' rule extended to continuous random variables:

$$g(y|x) = \frac{f(x|y) \cdot g(y)}{\int f(x|y) \cdot g(y) dy},$$

Due to the standardization $g(y|x)$ is a probability density; $g(y|x) \geq 0$, $\int g(y|x) dy = 1$.

Bayes' rule: parametric model

$$\pi(\theta|x) = \frac{f(x|\theta) \cdot \pi(\theta)}{\int_{\Theta} f(x|\theta) \cdot \pi(\theta) d\theta}$$

Terminology for Bayes' Rule:

- $\pi(\theta)$: **prior distribution** on Θ .
- $\pi(\theta|x)$: **posterior distribution** on Θ .
- $m(x) = \int_{\Theta} f(x|\theta) \cdot \pi(\theta) d\theta$: **marginal distribution** of x .



Uncertainty about the unknown θ is modeled by a probability distribution $\pi(\theta)$, and $\pi(\theta|x)$ expresses the uncertainty about the unknown θ after the observation of x .

We use probability as tool for all parts of our analysis. This is called coherence.

Mathematically: the unknown θ becomes a random variable. (x, θ) will have a joint distribution.

$$\pi(\theta|x) = \frac{f(x|\theta) \cdot \pi(\theta)}{m(x)},$$

Terminology:

- $m(x) = \int_{\Theta} f(x|\theta) \cdot \pi(\theta) d\theta$
- $\phi(x, \theta)$: **joint distribution** of (x, θ) .
-

$$\pi(\theta|x) m(x) = \phi(x, \theta) = f(x|\theta) \pi(\theta)$$

The notation

$$\int_{\Theta} f(x | \theta) \cdot \pi(\theta) d\theta$$

is imprecise by intent, as it can mean both a single integral and a multiple integral.

Bayesian Parametric Statistical Model

A Bayesian parametric statistical model consists of

- a parametric model

$$x \in f(x|\theta)$$

- a prior distribution

$$\theta \in \pi(\theta)$$

The quantity of interest

$$\theta|x \in \pi(\theta|x)$$



Any function $\pi(\cdot)$ such that

$$\pi(\theta) \geq 0,$$

and

$$\int_{\Theta} \pi(\theta) d\theta = 1,$$

can serve as a prior distribution.

Improper Prior Densities

But even functions with the properties

$$\pi(\theta) \geq 0,$$

and

$$\int_{\Theta} \pi(\theta) d\theta = \infty,$$

are also invoked as priors, and are called improper priors.



An Example

$X_i \mid M = m \in \mathbb{N} (m, \sigma_0^2)$, $M \in \mathbb{N} (\mu, s^2)$. $x^{(n)} = (x_1, \dots, x_n)$ a sample of i.i.d. X_i , $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$.

$$M \mid (X_1, \dots, X_n) \in \mathbb{N} \left(\frac{n\bar{x}/\sigma_0^2 + \mu/s^2}{n/\sigma_0^2 + 1/s^2}, \frac{1}{n/\sigma_0^2 + 1/s^2} \right)$$

i.e., $\pi(m|x^{(n)})$ is the density of this normal distribution. Here μ and s^2 are *hyperparameters*.

- $P(a(x) \leq \theta \leq b(x)) = \int_{a(x)}^{b(x)} \pi(\theta|x) d\theta$

- $\underbrace{P(a(x) \leq \theta \leq b(x))}$

This is a probability, not a degree of confidence

Confidence interval: Example

Take the example above

$$N\left(\frac{n\bar{x}/\sigma_0^2 + \mu/s^2}{n/\sigma_0^2 + 1/s^2}, \frac{1}{n/\sigma_0^2 + 1/s^2}\right)$$

Let $s \rightarrow \infty$ (the prior becomes improper). Then

$$N\left(\frac{n\bar{x}/\sigma_0^2 + \mu/s^2}{n/\sigma_0^2 + 1/s^2}, \frac{1}{n/\sigma_0^2 + 1/s^2}\right) \rightarrow N\left(\bar{x}, \frac{\sigma_0^2}{n}\right)$$

But then the Bayesian confidence interval $P(a(x) \leq \theta \leq b(x)) = 0.95$ becomes the familiar $\bar{x} \pm \lambda_{0.025} \frac{\sigma_0}{\sqrt{n}}$, and the common mistaken but natural interpretation of the confidence interval is correct !



General Aspects of Bayesian Inference

- inference is based on the observed x , not on an unobserved sample space.
- $\pi(\theta|x)$ is the only quantity evaluated for inference about θ .

On the other hand , evaluation of $\pi(\theta|x)$ is not in general possible by explicit means of integral calculus. This is where statistical inference needs Markov chain Monte Carlo (MCMC).



The distribution $f(x | \theta)$ regarded as a function of θ is known as the *likelihood function*

$$l(x; \theta) = f(x | \theta).$$

The likelihood function $l(\theta|x)$ thus compares the plausibilities of different parameter values for given x .

$$l(x; \theta) = -\log l(x; \theta).$$

is called the log likelihood function.

Likelihood Principle

The information brought by x about θ is entirely contained in the likelihood function $l(\theta | x)$.

If x_1 and x_2 are two observations depending on the same parameter θ such that there exists a constant c such that

$$l_1(x_1; \theta) = c \cdot l_2(x_2; \theta)$$

for every θ , then they bring the same information about θ and must lead to identical inferences.



Likelihood Principle

Conditions for the likelihood principle:

- inference should be about the same parameter set
- θ should include every unknown factor in the model.



Bayes' rule & likelihood

$$\begin{aligned}\pi(\theta|x) &= \frac{f(x|\theta) \cdot \pi(\theta)}{\int_{\Theta} f(x|\theta) \cdot \pi(\theta) d\theta'} \\ &= \frac{l(x;\theta) \cdot \pi(\theta)}{\int_{\Theta} l(x|\theta) \cdot \pi(\theta) d\theta}\end{aligned}$$

Hence Likelihood Principle is satisfied by Bayesian inference. There are ways of implementing the likelihood principle: MLE and MAP $\Rightarrow$

The Maximum Likelihood Estimate (MLE)

The **maximum likelihood estimate** MLE, $\hat{\theta}_{\text{ML}}$ of θ , is defined by

$$\begin{aligned}\hat{\theta}_{\text{ML}} &= \operatorname{argmax}_{\theta \in \Theta} f(x | \theta) \\ &= \operatorname{argmin}_{\theta \in \Theta} \ell(x; \theta)\end{aligned}$$

MLE is a parameter value that gives the observed x the highest possible probability.

The Maximum A Posterior Estimate (MAP)

The **maximum a posterior estimate** MAP $\hat{\theta}_{\text{MAP}}$ of θ is defined by

$$\hat{\theta}_{\text{MAP}} = \operatorname{argmax}_{\theta \in \Theta} \pi(\theta | x)$$

Definition: Conjugate Family of Priors

A family $\mathcal{F}$ of probability distributions on Θ is said to be **conjugate** or **closed under sampling** for a likelihood function

$$l(x; \theta) = f(x | \theta).$$

if for every $\pi \in \mathcal{F}$, the posterior distribution $\pi(\theta|x)$ also belongs to $\pi \in \mathcal{F}$. □

Above we have seen an example with normal prior density on the mean.



Definition: Conjugate Family of Priors

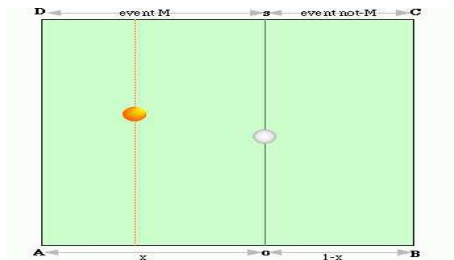
An intuitive way of understanding conjugate priors is that with conjugate priors the prior knowledge can be translated into equivalent sample information. See, e.g.,

$$N\left(\frac{n\bar{x}/\sigma_0^2 + \mu/s^2}{n/\sigma_0^2 + 1/s^2}, \frac{1}{n/\sigma_0^2 + 1/s^2}\right).$$

Next we reconsider another problem with a conjugate family of priors.



Bayes' Billiard Ball



The square table is made in such a way that, if the White or Orange ball be thrown on it, the probability that it rests on any part of the plane is the same. First, the White ball is thrown, and suppose it rests on line **os**; then the Orange ball is thrown n times. We define event M as any event the Orange ball rests between **A** and **o**, and not- M as its resting between **o** and **B**. What this means is this:

The first throw of the White ball determines the value of probability x (i.e. the probability of an unknown event) from a uniform distribution between 0 and 1; and then a series of trials, with probability x of success (i.e., M) is generated (this provides the data on which to infer the correct value of x).

Then, thanks to the geometrical representation of the problem in the Figure, we can obtain the solution to the initial problem, by calculating integration. Although we have omitted mathematical formulas, the preceding is the central idea.

Bayes' Billiard Ball

A billiard ball W is rolled on a line of length one, with a uniform probability of stopping anywhere. It stops at p , not disclosed to us. A second ball O is rolled n times under the same assumptions and X denotes the number of times O stops to the left of W . Given $X = x$, what inference can we make on p ?
(In the figure above $x \leftrightarrow p$.)



Modeling and Learning for Bayes' Billiard Ball

We let $\mathbf{P}$ be a random variable, whose values are denoted by p , $0 \leq p \leq 1$.
Parametric statistical model for rolls of Bayes' Billiard Ball O :
Conditional on $\mathbf{P} = p$, the rolls are outcomes of I.I.D $\text{Be}(p)$ R.V's.



Hence for $x = 0, 1, 2, \dots, n$,

$$\begin{aligned} f(x|p) &= P(X = x \mid \mathbf{P} = p) \\ &= \binom{n}{x} p^x \cdot (1 - p)^{n-x}, \end{aligned}$$

(the Binomial distribution)

The Posterior Density

Bayes' rule

$$\pi(p | x) = \frac{f(x | p) \cdot \pi(p)}{\int_0^1 f(x | p) \cdot \pi(p) dp}, 0 \leq p \leq 1$$

and zero elsewhere. The marginal distribution of x is

$$m(x) = \int_0^1 f(x | p) \cdot \pi(p) dp.$$



The Posterior Density

The posterior $\pi(p \mid x)$ expresses our updated uncertainty of the 'true' position of W given the data $X = x$.



The Posterior Density

One way to get further from here is to use an explicit expression for $\pi(p)$. There are many possible choices (some more systematic choices outlined below), some have straightforward analytical advantages. Laplace assumed that $p \in \mathbb{R}(0, 1)$. i.e.,

$$\pi(p) = \begin{cases} 1 & 0 \leq p \leq 1 \\ 0 & \text{elsewhere,} \end{cases}$$



The marginal distribution of x : uniform prior

$$\begin{aligned} m(x) &= \int_0^1 f(x | p) \cdot \pi(p) dp \\ &= \binom{n}{x} \int_0^1 p^x \cdot (1-p)^{n-x} dp. \end{aligned}$$

We use the Beta integral:

The Beta Integral

$$\int_0^1 p^{\alpha-1} (1-p)^{\beta-1} dp = \frac{\Gamma(\alpha)\Gamma(\beta)}{\Gamma(\alpha+\beta)}.$$

Recall also that $\Gamma(x+1) = x!$, if x is a positive integer. $\alpha = \beta = 1$ gives the distribution $R(0,1)$. We set

$$B(\alpha, \beta) := \frac{\Gamma(\alpha)\Gamma(\beta)}{\Gamma(\alpha+\beta)}.$$



The Beta Density

$$\pi(p) = \begin{cases} \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha)\Gamma(\beta)} p^{\alpha-1} (1-p)^{\beta-1} & 0 < p < 1 \\ 0 & \text{elsewhere.} \end{cases}$$

is a probability density $\text{Beta}(\alpha, \beta)$. $\alpha > 0$ and $\beta > 0$ are hyperparameters.



The marginal distribution of x : uniform prior

$$\begin{aligned} m(x) &= \int_0^1 f(x | p) \cdot \pi(p) dp \\ &= \binom{n}{x} \frac{x!(n-x)!}{(n+1)!} \end{aligned}$$

The marginal distribution of x , $p \in U(0,1)$

$$\begin{aligned} m(x) &= \int_0^1 f(x | p) \cdot dp = \binom{n}{x} \frac{x!(n-x)!}{(n+1)!} \\ &= \frac{n!}{x!(n-x)!} \frac{x!(n-x)!}{(n+1)!} = \frac{1}{(n+1)} \end{aligned}$$

There is an interpretation of Bayes' work claiming that the problem really attacked and solved by Bayes was: What should $\pi(p)$ be so that

$$\int_0^1 f(x | p) \cdot \pi(p) dp = \frac{1}{(n+1)}$$

holds for the Billiard Balls.



The Posterior Density for n rolls of Bayes' Orange Ball

$$\begin{aligned}\pi(p \mid x) &= \frac{\binom{n}{x} p^x \cdot (1-p)^{n-x}}{m(x)} \\ &= \begin{cases} \frac{(n+1)!}{x!(n-x)!} \cdot p^x (1-p)^{n-x} & 0 \leq p \leq 1 \\ 0 & \text{elsewhere.} \end{cases}\end{aligned}$$

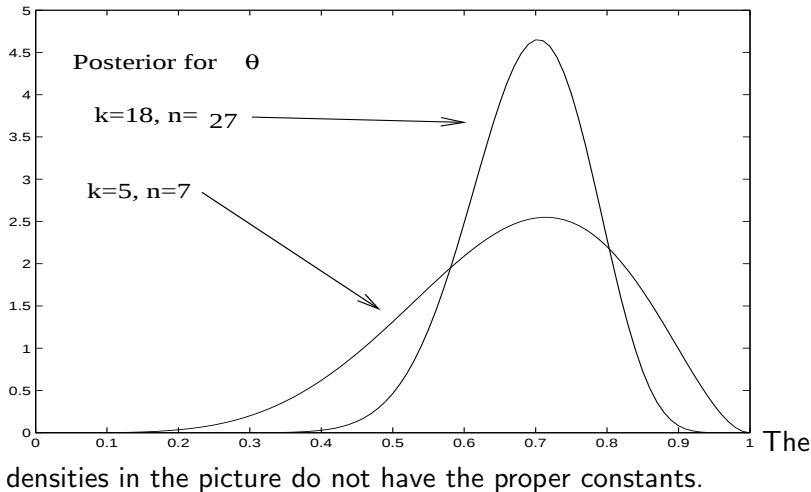
This is again a Beta density, i.e., we have used a conjugate family of priors.

The Posterior Density for n rolls of Bayes' Orange Ball

$$\frac{(n+1)!}{x!(n-x)!} = \frac{\Gamma(n+2)}{\Gamma(x+1)\Gamma(n-x+1)} = \frac{1}{B(x+1, n-x+1)}.$$



The Posterior Density for n rolls of Bayes' Orange Ball



The Posterior Density with n rolls of Bayes Ball, $p \in \text{Beta}(\alpha, \beta)$

$$\pi(p \mid x) = \begin{cases} \frac{1}{B(x+\alpha, n-x+\beta)} \cdot p^{x+\alpha-1} (1-p)^{\beta+n-x-1} & 0 \leq p \leq 1 \\ 0 & \text{elsewhere.} \end{cases}$$

This is the Beta density $\text{Beta}(\alpha + x, \beta + n - x)$.



End of Example

The example of Bayes' Billiard Balls W and O is discontinued for the moment, but will be reconsidered later.



The probability distribution $\phi(x, y)$ is a joint distribution of (x, y) .
We assume a parametric model $(x, y, \theta) \in \phi(x, y, \theta)$ so that:

$$\phi(x, y) = \int \phi(x, y, \theta) d\theta.$$



Predictive Distribution (1) (General)

The predictive distribution $g(y|x)$ of y conditional on x , denoted by $g(y|x)$ is

$$\begin{aligned} g(y|x) &= \frac{\phi(x, y)}{m(x)} = \frac{\int \phi(x, y, \theta) d\theta}{m(x)} \\ \frac{\phi(x, y)}{m(x)} &= \frac{\int \phi(x, y, \theta) d\theta}{m(x)} = \frac{\int g(y|x, \theta) \phi(x, \theta) d\theta}{m(x)} \\ &= \frac{\int g(y|x, \theta) f(x|\theta) \pi(\theta) d\theta}{m(x)} \end{aligned}$$

Predictive Distribution (2)

i.e.,

$$\begin{aligned} g(y|x) &= \frac{\int g(y|x, \theta) f(x|\theta) \pi(\theta) d\theta}{m(x)} \\ &= \frac{\int g(y|x, \theta) f(x|\theta) \pi(\theta) d\theta}{\int f(x|\theta) \cdot \pi(\theta) d\theta} \\ &= \int g(y|x, \theta) \frac{f(x|\theta) \pi(\theta)}{\int f(x|\theta) \cdot \pi(\theta) d\theta} d\theta \end{aligned}$$



Predictive Distribution (3)

The predictive distribution $g(y|x)$ is thus

$$g(y|x) = \int g(y|x, \theta) \pi(\theta | x) d\theta.$$



Predictive Distribution (4)

If y and x are conditionally independent given θ , then

$$g(y|x, \theta) = g(y|\theta)$$

and

$$g(y|x) = \int g(y|\theta) \pi(\theta | x) d\theta.$$



Predictive Distribution (5)

Introduce order (e.g., in time). $x = x^{(n)} = (x_1, \dots, x_n)$ $y = x_{n+1}$.

$$g(x_{n+1}|x^{(n)}) = \int g(x_{n+1}|x^{(n)}, \theta) \pi(\theta | x^{(n)}) d\theta.$$

We assume conditional independence

$$\begin{aligned} \phi(x, y|\theta) &= \phi(x^{(n)}, x_{n+1}|\theta) = (f(x_1|\theta) \cdots f(x_n|\theta)) \cdot f(x_{n+1}|\theta) \\ &= \phi(x^{(n)}|\theta) \cdot f(x_{n+1}|\theta) \end{aligned}$$



Predictive Distribution (6)

$$\phi \left(x^{(n)}, x_{n+1} \mid \theta \right) = \phi \left(x^{(n)} \mid \theta \right) \cdot f \left(x_{n+1} \mid \theta \right)$$

Then

$$g \left(x_{n+1} \mid x^{(n)}, \theta \right) = \frac{\phi \left(x^{(n)}, x_{n+1} \mid \theta \right)}{\phi \left(x^{(n)} \mid \theta \right)} = f \left(x_{n+1} \mid \theta \right)$$

and

$$g \left(x_{n+1} \mid x^{(n)} \right) = \int f \left(x_{n+1} \mid \theta \right) \pi \left(\theta \mid x^{(n)} \right) d\theta$$

Predictive Distribution (7)

$$g\left(x_{n+1}|x^{(n)}\right)=\int f\left(x_{n+1}|\theta\right) \pi\left(\theta \mid x^{(n)}\right) d \theta$$

We often need to update $g\left(x_{n+1}|x^{(n)}\right)$ for new observations. We need clearly only to update $\pi\left(\theta \mid x^{(n)}\right)$.



Up-date of Posterior Distribution (1)

$$\begin{aligned}\pi(\theta \mid x^{(n+1)}) &= \frac{f(x^{(n+1)} \mid \theta) \pi(\theta)}{\int f(x^{(n+1)} \mid \theta) \cdot \pi(\theta) d\theta} \\&= \frac{f(x_{n+1} \mid \theta) f(x_1 \mid \theta) \cdots f(x_n \mid \theta) \pi(\theta)}{\int f(x_{n+1} \mid \theta) f(x_1 \mid \theta) \cdots f(x_n \mid \theta) \cdot \pi(\theta) d\theta} \\&= \frac{f(x_{n+1} \mid \theta) \frac{f(x_1 \mid \theta) \cdots f(x_n \mid \theta) \pi(\theta)}{m(x^{(n)})}}{\int f(x_{n+1} \mid \theta) \frac{f(x_1 \mid \theta) \cdots f(x_n \mid \theta) \pi(\theta)}{m(x^{(n)})} d\theta} \\&= \frac{f(x_{n+1} \mid \theta) \pi(\theta \mid x^{(n)})}{\int f(x_{n+1} \mid \theta) \pi(\theta \mid x^{(n)}) d\theta}\end{aligned}$$

Up-date of Posterior Distribution (2)

$$\pi(\theta | x^{(n+1)}) = \frac{f(x_{n+1}|\theta) \pi(\theta|x^{(n)})}{\int f(x_{n+1}|\theta) \pi(\theta|x^{(n)}) d\theta}$$

Hence, under the assumptions made, we can update posterior distribution in a sequential manner. Or, we can use the posterior of $\pi(\theta|x^{(n)})$ as a new prior for computing $\pi(\theta | x^{(n+1)})$.

Fundamental Task of Statistical Inference

The fundamental task of statistical inference is to pass from one set of observations x to express an opinion about another, as yet unobserved set y .

We have above accomplished this in terms of the predictive distribution

$$g(y|x) = \int f(y|\theta) \pi(\theta | x) d\theta,$$

which shows the role of learning about parameters in accomplishing this task



An Example $f(x | \theta) \leftrightarrow N(\mu, \sigma^2)$

$f(x | \theta) \leftrightarrow N(\mu, \sigma^2)(x)$, the mean μ and variance σ^2 are unknown, i.e., $\theta = (\mu, \sigma^2)$. The (improper) prior density is taken as

$$\pi(\theta) \propto d\mu \frac{1}{\sigma} d\sigma.$$

Let $x = \underline{x}_n = (x^{(1)}, \dots, x^{(n)})$, $x^{(i)} \in N(\mu, \sigma^2)$. We have the estimates $\hat{\mu} = \frac{1}{n} \sum_{i=1}^n x^{(i)}$, and $s^2 = \frac{1}{n-1} \sum_{i=1}^n (x^{(i)} - \hat{\mu})^2$.



An Example of Bayesian prediction $f(x | \theta) \leftrightarrow \mathcal{N}(\mu, \sigma^2)$

$\pi(\theta) \propto d\mu \frac{1}{\sigma} d\sigma$

The predictive density is

$$g\left(x^{(n+1)} | \underline{x}_n\right) = \sqrt{\frac{n}{(n^2 - 1)\pi}} \cdot \frac{\Gamma\left(\frac{n}{2}\right)}{\Gamma\left(\frac{1}{2}(n-1)s\right)} \cdot \left(1 + \frac{n\left(x^{(n+1)} - \hat{\mu}\right)^2}{(n^2 - 1)s^2}\right)^{-n/2}$$

The predictive density

$$\sqrt{\frac{n}{(n^2 - 1)\pi}} \cdot \frac{\Gamma\left(\frac{n}{2}\right)}{\Gamma\left(\frac{1}{2}(n - 1)s\right)} \cdot \left(1 + \frac{n(x^{(n+1)} - \hat{\mu})^2}{(n^2 - 1)s^2}\right)^{-n/2}$$

is known in Bayesian statistics as 't-like distribution'¹.

¹J.M. Dickey (1968): Three Multidimensional-Integral Identities with Bayesian Applications. *The Annals of Mathematical Statistics*, 39, pp. 1615–1628.

Predictive Distribution for the next m rolls of Bayes' Orange Ball

The predictive distribution of y positions of O left of W in m additional rolls is

$$\begin{aligned} g(y|x) &= \int_0^1 f(y|p) \pi(p|x) dp \\ &= \binom{m}{y} \int_0^1 p^y \cdot (1-p)^{m-y} \pi(p|x) dp \end{aligned}$$

$$y = 0, 1, \dots, m$$



Predictive Distribution for the next m rolls of Bayes' Ball

$$\begin{aligned} \int_0^1 p^y \cdot (1-p)^{m-y} \pi(p \mid x) dp &= \\ \int_0^1 p^y \cdot (1-p)^{m-y} \frac{1}{B(x+1, n-x+1)} \cdot p^x (1-p)^{n-x} dp &= \\ = \frac{1}{B(x+1, n-x+1)} \int_0^1 p^{y+x} \cdot (1-p)^{m+n-y-x} dp &= \\ = \frac{B(y+x+1, m+n-x-y+1)}{B(x+1, n-x+1)}. \end{aligned}$$

by the Beta integral.

Predictive Distribution for the next m rolls of Bayes' Ball

$$\begin{aligned} g(y|x; m) &= \binom{m}{y} \int_0^1 p^y \cdot (1-p)^{m-y} \pi(p|x) dp \\ &= \binom{m}{y} \frac{B(y+x+1, m+n-x-y+1)}{B(x+1, n-x+1)}. \\ &= \frac{m!}{(m-y)!y!} \frac{\Gamma(n+2)\Gamma(y+x+1)\Gamma(m+n-x-y+1)}{\Gamma(x+1)\Gamma(n-x+1)\Gamma(m+n+2)}. \end{aligned}$$

Predictive Distribution for $y = 1$ in the next $m = 1$ roll of Bayes' Ball

$$\begin{aligned} g(1|x;1) &= \frac{\Gamma(x+2)\Gamma(n-x+1)\Gamma(n+2)}{\Gamma(x+1)\Gamma(n-x+1)\Gamma(n+3)} \\ &= \frac{(x+1)!(n-x)!(n+1)!}{x!(n-x)!(n+2)!} = \frac{x+1}{n+2} \end{aligned}$$

A famous predictive probability, known as *Laplace's rule of succession*,

$$\frac{x+1}{n+2} = \frac{x+1}{(n+1)+1}$$

The prior knowledge can be translated into equivalent sample information.



The Maximum Likelihood Estimate of p in Bayes' Billiard

$$\begin{aligned}\hat{p}_{\text{ML}} &= \operatorname{argmin}_{0 \leq p \leq 1} \ell(p \mid x) \\ &= \operatorname{argmin}_{0 \leq p \leq 1} \left[-\log \binom{n}{x} - x \log p - (n - x) \log (1 - p) \right] . \\ &= \operatorname{argmin}_{0 \leq p \leq 1} (-x \log p - (n - x) \log (1 - p)) . \\ &\Rightarrow \\ \hat{p}_{\text{ML}} &= \frac{x}{n}\end{aligned}$$

If you observed $x = 0$, would you believe in the estimate $\hat{p} = 0$ for all future purposes ?



MLE and Predictive Probability in Bayes' Billiard

The predictive probability found above

$$\frac{x+1}{n+2} = \frac{x+1}{(n+1)+1}$$

is a maximum likelihood estimate of p when $n+1$ rolls of the ball O and the first roll of the ball W are included in the data.



Q: How do we choose $\pi(\theta)$?

- Assessment (by Questionnaires)
- Conjugate prior
- Non-informative or reference prior
 - Laplace's prior
 - Jeffreys' prior
- Maximum entropy prior



(One form of) Bayesian statistics relies upon a **personalistic theory of probability** for quantification of prior knowledge. In such a theory

- probability measures the confidence that a particular individual (assessor) has in the truth of a particular proposition
- no attempt is made to specify which assessments are correct
- personal probabilities should satisfy certain postulates of coherence.

Choice of prior distributions: questionnaires

R.L.Winkler (work published in

- Robert L. Winkler: The Assessment of Prior Distributions in Bayesian Analysis

Journal of the American Statistical Association, Vol. 62, No. 319.
(Sep., 1967), pp. 776-800.)

devises questionnaires (or interviews) to elicit information to write down a prior distribution. Students of Univ. of Chigago were asked to, e.g., assess the uncertainty about the probability of a randomly chosen student of Univ. of Chigago being Roman Catholic using a probability distribution. The assessment was done by four different methods, like giving fractiles, making bets, assessing impact of additional data, drawing graphs. One interesting finding is that the assessments by the same person using different methods may be conflicting.



Diffuse/Non-diffuse prior distributions by assessment

The priors in Winkler's study are not diffuse: the students of Univ. of Chicago have, since they have been around, an idea about the number of Roman Catholics at the campus of Univ. of Chicago.



The interviews by Winkler were mathematically speaking all concerned with assessing the prior of θ in a Bernoulli $\text{Be}(\theta)$ – I.I.D. process. Winkler claims a sensitivity analysis (loc.cit p. 791) showing that the prior distributions assessed by the interviews yielded posterior distributions that were ‘only little’ different (by a test of goodness-of-fit) from those obtained from Beta densities on θ .

Choice of prior distributions by elicitation

- A. O'Hagan: Eliciting Expert Beliefs in Substantial Practical Applications. *The Statistician* , 47, pp. 21–35, 1998.

Not only priors are elicited in

- R.L. Keeney & D. von Winterfeldt: Eliciting Probabilities in Complex Technical Problems. *IEEE Transactions on Engineering Management*, 38, pp.191–201, 1991.

Choice of prior distributions by assessment

But this line of study can evolve rapidly to a topic of research in psychology or (economic) behaviour, c.f.,

- C-A. S. Stael von Holstein: *Assessment and Evaluation of Subjective Probability Distributions*. 1970, Stockholm School of Economics.
- A. G. Wilson: *Cognitive factors affecting subjective probability assessments*.

ISDS Discussion Paper # 94-02,
<http://www.isds.duke.edu/>

so it feels safe to leave the matter at rest here.



Parametric Statistical Model, n I.I.D. $|\theta$ rv's

$$x_i|\theta \in f(x|\theta), \text{ I.I.D. },$$

or **independent, identically, distributed conditional on θ**

$$x^{(n)} = (x_1, x_2, \dots, x_n) \in \mathcal{X}^n$$

$f(x|\theta)$ is a probability density on R^p . $f(x|\theta)$ is a known function of x and θ . θ is an unknown parameter $\in \Theta =$ a vector space of finite dimension.



Asymptotic Shape of the Posterior

Let us assume that $f(x|\theta)$ is a density with a scalar parameter (for simplicity of notation), and that $f(x|\theta)$ is some $k \geq 2$ times differentiable in θ . We let $\hat{\theta}_{ML}$ be the maximum likelihood estimate of θ . We expand the log likelihood function around $\hat{\theta}_{ML}$

$$\begin{aligned}\log f\left(x^{(n)}|\theta\right) = & \\ \log f\left(x^{(n)}|\hat{\theta}_{ML}\right) + \left(\theta - \hat{\theta}_{ML}\right) \frac{d}{d\theta} \log f\left(x^{(n)}|\hat{\theta}_{ML}\right) & \\ + \frac{1}{2} \left(\theta - \hat{\theta}_{ML}\right)^2 \frac{d^2}{d\theta^2} \log f\left(x^{(n)}|\hat{\theta}_{ML}\right) + R_n(\theta) &\end{aligned}$$



Asymptotic Shape of the Posterior

But here $\hat{\theta}_{ML}$ is a solution of the equation

$$\frac{d}{d\theta} \log f \left(x^{(n)} | \hat{\theta}_{ML} \right) = 0$$

Hence

$$\begin{aligned} \log f \left(x^{(n)} | \theta \right) = \\ \log f \left(x^{(n)} | \hat{\theta}_{ML} \right) + \frac{1}{2} \left(\theta - \hat{\theta}_{ML} \right)^2 \frac{d^2}{d\theta^2} \log f \left(x^{(n)} | \hat{\theta}_{ML} \right) + R_n(\theta) \end{aligned}$$



Asymptotic Shape of the Posterior: Law of Large Numbers

We have by assumption of I.I.D. data

$$\frac{d^2}{d\theta^2} \log f \left(x^{(n)} | \theta \right) = \sum_{l=1}^n \frac{d^2}{d\theta^2} \log f (x_l | \theta)$$

We set $Y_l = \frac{d^2}{d\theta^2} \log f (x_l | \theta)$. Then the Law of Large Numbers says that

$$\frac{1}{n} \sum_{l=1}^n Y_l \rightarrow E[Y], n \rightarrow \infty$$

where

$$E[Y] = \int_{\mathcal{X}} \frac{d^2}{d\theta^2} \log f (x | \theta) f (x | \theta) dx$$



Asymptotic Shape of the Posterior: Fisher Information

The integral

$$I(\theta) = - \int_{\mathcal{X}} \frac{d^2}{d\theta^2} \log f(x|\theta) f(x|\theta) dx$$

is called *Fisher information*.



Asymptotic Shape of the Posterior: Fisher Information

Then we may feel inclined to believe that

$$\frac{d^2}{d\theta^2} \log f \left(x^{(n)} | \hat{\theta}_{ML} \right) = \sum_{l=1}^n \frac{d^2}{d\theta^2} \log f \left(x_l | \hat{\theta}_{ML} \right) \approx -n \cdot I \left(\hat{\theta}_{ML} \right)$$

Note that even $\hat{\theta}_{ML}$ depends on n .



Asymptotic Shape of the Posterior

This gives

$$\log f\left(x^{(n)}|\theta\right) \approx \log f\left(x^{(n)}|\hat{\theta}_{ML}\right) - \frac{1}{2}\left(\theta - \hat{\theta}_{ML}\right)^2 n \cdot I\left(\hat{\theta}_{ML}\right)$$

The first term does not involve θ .

Asymptotic Shape of the Posterior

Then

$$f\left(x^{(n)}|\theta\right) \approx e^{-\frac{n}{2}(\theta-\hat{\theta}_{ML})^2 \cdot I(\hat{\theta}_{ML})}$$

The interpretation of the relation is that the likelihood function can be for large n be approximated by a normal density for which the mean is $\hat{\theta}_{ML}$ and the variance is $\frac{1}{nI(\hat{\theta}_{ML})}$.



Finally: Jeffreys' prior

Let $I(\theta)$ be the Fisher information of a parametric model. Take the prior density as

$$\pi(\theta) \stackrel{\text{def}}{=} \frac{\sqrt{I(\theta)}}{\int \sqrt{I(\theta')} d\theta'},$$

assuming the integral exists. This choice of prior is known as Jeffreys' prior. This prior is invariant to monotonous transformations of θ .



Finally: Jeffreys' prior

It turns out that Jeffreys' prior for a binomial likelihood is obtained by $\alpha = 1/2$ and $\beta = 1/2$ in

$$\pi(p) = \begin{cases} \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha)\Gamma(\beta)} p^{\alpha-1} (1-p)^{\beta-1} & 0 < p < 1 \\ 0 & \text{elsewhere.} \end{cases}$$

and with $\Gamma(1/2 + 1/2) = 1$, $\Gamma(1/2) = \sqrt{\pi}$,

$$\pi_{\text{Jeffreys}}(p) = \begin{cases} \frac{1}{\pi} p^{-1/2} (1-p)^{-1/2} & 0 < p < 1 \\ 0 & \text{elsewhere.} \end{cases}$$

