

Decision Theory

- Classical Inference
- Bayesian Inference
- Classical Inference :-

Inferences can be drawn on the basis of sample information by taking parameters as a fixed constants is known as Classical inference.

on the hand, sample values are treated as random variables parameters as constants.

MLE is the wonderful method to establish the same.

Bayesian Inference :-

Inferences in which sample information is treated as fixed constant and parameters as a random variable

(having the certain distribution).

→ Prior Distribution
→ Posterior Distribution

Prior Distribution :-

In Bayesian context, the random variable which are governed with certain distributions are termed as Prior Distribution.

Types of Prior distribution :-

- (i) Informative Priors
- (ii) Non-informative priors

Informative Prior :-

Priors having more or more information about population characteristics is known as informative prior.

For example :- Gamma prior, conjugate prior

Non-Informative Prior :-

A prior containing very little or having

no information 'about population' characteristics is known as non-informative prior.

For example:-

uniform prior, Jeffreys prior

Uniform Prior:-

uniform prior of parameter θ is denoted by $\pi_U(\theta)$ and defined as

$$\pi_U(\theta) \propto 1$$

Jeffreys Prior:-

Jeffreys prior for parameter θ is denoted by $\pi_J(\theta)$ and defined as

$$\pi_J(\theta) \propto [I(\theta)]^{1/2}$$

where $I(\theta)$ is the Fisher's information on θ

$$I(\theta) = -E\left(\frac{\partial^2}{\partial \theta^2} \log L\right)$$

Reminder :-

Bayes

$$P(E_i|A) = \frac{P(E_i) \cdot P(A|E_i)}{\sum_{i=1}^n P(E_i) \cdot P(A|E_i)} \quad i=1, 2, \dots, n$$

$P(E_i)$ — Prior probabilities

$P(A|E_i)$ — likelihood prob.

$P(E_i|A)$ — Posterior prob.

Conjugate Prior :-

If prior and posterior distributions belongs to same family then the associated prior is known as

Conjugate prior

For example:- Gamma prior.

$$\pi_c(\theta) = \frac{a^b e^{-a\theta} \theta^{b-1}}{\Gamma(b)}$$

Posterior Distribution:-

Suppose that the pdf is $f(x, \theta)$ with prior distⁿ $f(\theta)$ then posterior pdf will be -

$$f(\theta | \underline{x}) = \frac{f(\theta) \cdot L(\underline{x} | \theta)}{\int_{\theta} f(\theta) \cdot L(\underline{x} | \theta) d\theta}$$

$$\pi(\theta | \underline{x}) = \frac{\pi(\theta) \cdot L(\underline{x} | \theta)}{\int_{\theta} \pi(\theta) \cdot L(\underline{x} | \theta) d\theta}$$

$$\pi(\theta | \underline{x}) \propto \pi(\theta) \cdot L(\underline{x} | \theta)$$

Loss Function :-

A non-negative function which maps from the Cartesian product of parameter space (Θ) and action space ($\mathcal{A}$) to real line ($\mathbb{R}^+$) and is denoted by $L(\theta, a)$ and defined as

$$L(\theta, a): \Theta \times \mathcal{A} \rightarrow \mathbb{R}^+$$

Types of Loss Function :-

- (1) Squared Error Loss Function (SELF)
OR
Quadratic Loss Function
- (2) General Entropy Loss Function (GELF)
- (3) Absolute Error Loss Function
- (4) LINEX Loss Function.