

today: statistics frequentist vs. Bayesian

Statistical concepts



example: model N indep coin-flips

prob. theory $\rightarrow$ prob. k heads in a row

statistics: I have observed k heads, $N-k$ tails, what is Θ ?

frequentist vs. Bayesian

semantic of prob.: meaning of a prob.?

a) (traditional) frequentist semantic

$P\{X=x\}$ represents the limiting frequency of observing $X=x$
 $\text{if I could repeat } N\# \text{ of } \underline{\text{i.i.d}}$ experiments

b) Bayesian (subjective) semantic

$P\{X=x\}$ encodes an agent "belief" that $X=x$

laws of prob. characterizes a "rational" way to combine "beliefs"
 and "evidence" [observations]

[has motivation in terms of gambling, utility/decision theory, etc...]

operationally:

Bayesian approach: $\otimes$ very simple philosophically

treat all uncertain quantities as R.V.

i.e. - encode all knowledge about the system ("beliefs")
 as a "prior" on probabilistic models

and then use laws of prob. (and Bayes rule) to
 get updated beliefs and answer?

justification for frequentist semantic

• for discrete R.V. X , suppose $P\{X=x\} = \theta$

$$\Rightarrow P\{X \neq x\} = 1 - \theta$$

$$B \triangleq \mathbb{1}_{\{X=x\}} \Rightarrow B \sim \text{Bern}(\theta) \text{ R.V.}$$

Indicator fct. $\mathbb{1}_A(u) = \begin{cases} 1 & \text{if } u \in A \\ 0 & \text{o.w.} \end{cases}$

repeat i.i.d. i.e. $B_i \stackrel{\text{i.i.d.}}{\sim} \text{Bern}(\theta)$

by L.L.N. law of large numbers $\frac{1}{n} \sum_{i=1}^n B_i \xrightarrow{\text{a.s.}} E[B_1] = \theta$

Limiting frequency

by C.L.T. $\sqrt{n} \left(\frac{1}{n} \sum_{i=1}^n B_i - \theta \right) \xrightarrow{d} N(0, \theta(1-\theta))$
 $\sim \text{Bin}(n, \theta)$

Coin flips - Bayesian approach

baised coin flips

unknown $\Rightarrow$ model it as R.V.

we believe $X \sim \text{Bin}(n, \theta)$

$\Rightarrow$ need $a_p(\theta)$ "prior distribution"

$$\Omega_{\theta} = [0, 1]$$

suppose we observe $X=x$ (result of n coin flips)

then we can "update" our belief about θ by using Bayes rule

$$p(\theta = \theta | X=x) = \frac{p(X=x | \theta = \theta) p(\theta = \theta)}{p(X=x)}$$

posterior belief $\quad$ observation model / likelihood $\quad$ prior belief $\quad$ normalization "marginal likelihood"

[note: $p(x|\theta) \rightarrow$ pmf $\quad p(x, \theta)$ is a "mixed distribution"]
 $p(\theta) \rightarrow$ pdf

example:

suppose $p(\theta)$ is uniform on $[0, 1]$ "no specific preference"

$$p(\theta|x) \propto p(x|\theta) p(\theta)$$

"proportional to"

$$p(\theta|x) \propto \theta^x (1-\theta)^{n-x} \cdot \mathbb{1}_{[0,1]}(\theta)$$

up to a constant in θ

scaling: $\int_0^1 \theta^x (1-\theta)^{n-x} d\theta = B(x+1, n-x+1)$

scaling: $\int_0^1 \theta^x (1-\theta)^{n-x} d\theta = B(x+1, n-x+1)$

normalization constant $\int_0^1 p(\theta|x) d\theta = 1$

$B(a,b) \triangleq \frac{\Gamma(a)\Gamma(b)}{\Gamma(a+b)}$

beta fct. $\Gamma(a) = \int_0^\infty u^{a-1} e^{-u} du$ gamma fct.

here $p(\theta|x)$ is called a "beta distribution"

$$B(\theta|\alpha, \beta) \triangleq \frac{\theta^{\alpha-1} (1-\theta)^{\beta-1} \mathbb{1}_{[0,1]}(\theta)}{B(\alpha, \beta)}$$

parameters

• uniform distribution: $B(\theta|1,1)$

posterior density: $B(\theta|x+1, n-x+1)$ as prior counts

exercise to the reader: if we use $B(\alpha_0, \beta_0)$ as prior

posterior will be $B(x+\alpha_0, n-x+\beta_0)$

15h32

(*) posterior $p(\theta|X=x)$ contains all the info from data x that we need to answer questions about θ

e.g. question: what is prob. of head ($F=1$) on the next flip?

as a frequentist $P(F=1|\text{data}) = \hat{\theta}$ relation to mean "estimate"

as a Bayesian $P(F=1|X=x) = \int P(F=1, \theta=\theta|X=x) d\theta$

product rule $= \int_{\mathcal{E}} \underbrace{P(F=1|\theta=\theta, X=x)}_{=\theta \text{ (by our model)}} \underbrace{P(\theta|X=x)}_{\text{posterior}} d\theta$

$= \int_{\mathcal{E}} \theta p(\theta|X=x) d\theta = \mathbb{E}[\theta|X=x]$

* a meaningful "Bayesian" estimator of θ

$\hat{\theta}_{\text{Bayes}}(x) \triangleq \mathbb{E}[\theta|X=x]$ (posterior mean)

relation: $\hat{\theta}$: observation $\rightarrow \mathbb{R}_{[0,1]}$

our coin example: $p(\theta|x) = \text{Beta}(\theta|\alpha=x+1, \beta=n-x+1)$

mean of a beta R.V. $\frac{\alpha}{\alpha+\beta}$

thus $\hat{\theta}_{\text{Bayes}}(x) = \mathbb{E}[\theta|x] = \frac{x+1}{n+2}$

$$\text{thus } \hat{\theta}_{\text{Bayes}}(x) = \mathbb{E}[\theta | x] = \frac{x+1}{n+2}$$

here, biased estimator $\mathbb{E}_x[\hat{\theta}(x)] \neq \theta$

$$= \mathbb{E}\left[\frac{X+1}{n+2}\right] = \frac{\mathbb{E}X+1}{n+2} = \frac{n\theta+1}{n+2} \neq \theta$$

but asymptotically unbiased $\xrightarrow{n \rightarrow \infty} \theta$

compare & contrast with $\hat{\theta}_{\text{MLE}}(x) = \frac{x}{n}$ [unbiased $\mathbb{E}\left[\frac{X}{n}\right] = \frac{\mathbb{E}X}{n} = \theta$]

to summarize:

- as a Bayesian: get a posterior + use law of probabilities
- in "frequentist statistics"

consider multiple estimators

MLE
moment matching
Bayesian posterior mean
MAP
regularized MLE

and then analyze the statistical properties of estimator:

- biased?
- variance?
- consistent?
- frequentist risk?

Maximum likelihood principle

setup: given a parametric family $p(x; \theta)$ for $\theta \in \Theta$

we want to estimate/learn θ from x

$$\hat{\theta}_{\text{MLE}}(x) \triangleq \underset{\theta \in \Theta}{\operatorname{argmax}} p(x; \theta)$$

$L \triangleq L(\theta)$
"likelihood function" of θ

$\hat{\theta}_{\text{MLE}}(x)$ maximizes $p(x, \cdot)$

MLE example I: binomial:

n coin flips $\Omega_X = 0:n$

$$X \sim \text{Bin}(n, \theta) \quad p(x; \theta) = \binom{n}{x} \theta^x (1-\theta)^{n-x}$$

trick: to maximize log $L(\theta)$ instead of $L(\theta)$

trick: to maximize $\log L(\theta)$ instead of $L(\theta)$
 $\stackrel{!}{=} \ell(\theta)$ log likelihood

justification: $\log(\cdot)$ is strictly increasing

$$\text{i.e. } a < b \Leftrightarrow \log a < \log b \quad (a, b > 0)$$

$$\Rightarrow \arg\max_{\theta \in \Theta} \log p(x; \theta) = \arg\max_{\theta \in \Theta} p(x; \theta)$$

$$\log p(x; \theta) = \underbrace{\log \binom{n}{x}}_{\text{constant w.r. to } \theta} + x \log \theta + (n-x) \log (1-\theta) = \ell(\theta)$$



$$\text{look for } \theta \text{ st. } \frac{\partial \ell}{\partial \theta} = 0$$

$$\text{want: } \frac{x}{\theta} - \frac{n-x}{1-\theta} = 0$$

$$\begin{aligned} x(1-\theta) &= \theta(n-x) \\ x - x\theta &= n\theta - \theta x \end{aligned}$$

$$\text{hence } \boxed{\hat{\theta}_{MLE}(x) = \frac{x}{n}}$$

Used often
as solution in
optimization

$$\boxed{\theta^* = \frac{x}{n}}$$