

Lecture 13

Part I - Estimating P 's

Part II - Model Selection

- Mixture models

- randomness of estimated parameter ✓

- Bayesian estimation ←

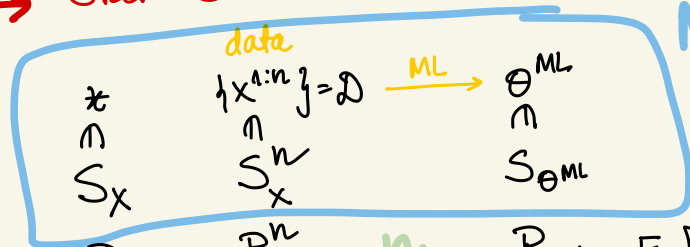
- ML estimation ✓

Poll t.b. posted:
- topics for rest
of 391

✓ Estimation $\begin{cases} \text{discrete} \\ \text{continuous} \end{cases}$

✓ Model Selection

→ Stat. Estimators as R.V



ML Estimation

θ^{ML} guesses !! $\Leftrightarrow \theta^{\text{ML}}$ is a R.V.

assume known P_x $\xrightarrow{\text{iid}}$ P_x^n $\xrightarrow{P_x}$ $P_{\theta^{\text{ML}}}, E[\theta^{\text{ML}}], \text{Var} \theta^{\text{ML}}$

Wanted:

θ^{true} 0
for $n \rightarrow \infty$
consistency

• $E[\theta^{\text{ML}}] = \theta^{\text{true}}$ for finite n
 P_x
Unbiasedness

$$E[\theta^{\text{ML}}] - \theta^{\text{true}} = \text{bias}$$

a special form of bias (refers to a parameter)

Normal (μ, σ^2)

$$\bullet \mu^{ML} = \frac{1}{n} \sum_{i=1}^n x^i$$

$$\left[E[\mu^{ML}] = E\left[\frac{1}{n} \sum_{i=1}^n x^i\right] = \frac{1}{n} \sum_{i=1}^n \underbrace{E[x^i]}_{\mu} \right] \rightarrow \text{unbiased}$$

$$\bullet \left[\text{Var } \mu^{ML} = \text{Var}\left[\frac{1}{n} \sum_{i=1}^n x^i\right] = \frac{1}{n^2} \sum_{i=1}^n \underbrace{\text{Var}(x^i)}_{\sigma^2} = \frac{n\sigma^2}{n^2} = \frac{1}{n} \sigma^2 \right]$$

↑ randomness in sample !!

↑ iid (independent!)

↑ $\text{Var} \rightarrow 0$ for $n \rightarrow \infty$

if σ^2 smaller, n can be smaller too

In general

$$\mu = E[x]$$

P_x ← unknown distribution

$$\hat{\mu} = \frac{1}{n} \sum_{i=1}^n x^i, \quad x^{1:n} \sim P_x \text{ iid} \Rightarrow E[\hat{\mu}] = \mu$$

Estimation of σ^2 for $N(\mu, \sigma^2)$ ← assume $\mu_{\text{true}} = N(\mu, \sigma^2)$

$$(\sigma^2)^{\text{ML}} = \frac{1}{n} \sum_{i=1}^n (x^i - \mu^{\text{ML}})^2 \quad \frac{\sum x^i}{n} \quad \text{or} \quad \begin{cases} \text{Var } \mu_{\text{true}} = \sigma^2 \\ E[x]_{\text{true}} = \mu \end{cases}$$

↑ plug-in estimator

$$\begin{aligned} E[(\sigma^2)^{\text{ML}}] &= E\left[\frac{1}{n} \sum_{i=1}^n (x^i - \mu + \mu - \mu^{\text{ML}})^2\right] \\ &= \frac{1}{n} \sum_{i=1}^n E[(x^i - \mu)^2 + (\mu - \mu^{\text{ML}})^2 + 2(x^i - \mu)(\mu - \mu^{\text{ML}})] \\ &= \frac{1}{n} \sum_{i=1}^n \left\{ E[(x^i - \mu)^2] + E[(\mu^{\text{ML}} - \mu)^2] + 2E[(x^i - \mu)(\mu - \mu^{\text{ML}})] \right\} \\ &= \frac{1}{n} \sum_{i=1}^n \left\{ \underbrace{E[(x^i - \mu)^2]}_{\sigma^2} + \underbrace{E[(\mu^{\text{ML}} - \mu)^2]}_{E[(\mu^{\text{ML}} - E[\mu^{\text{ML}}])^2]} + 2E[(x^i - \mu)(\mu - \mu^{\text{ML}})] \right\} \end{aligned}$$

$$E[(x - \mu)^2] = \text{Var } x = \sigma^2$$

$$\text{Var } \mu^{\text{ML}} = \frac{\sigma^2}{n}$$

$$E[\mu^{\text{ML}}] = \mu$$

$$= \frac{1}{n} \sum_{i=1}^n \left\{ \sigma^2 + \frac{\sigma^2}{n} - 2\frac{\sigma^2}{n} \right\}$$

$$= \sigma^2 \frac{n-1}{n} \neq \sigma^2 \quad \text{BIAS} < 0$$

under estimation

$$\begin{aligned}
 & \star \quad E[(x^i - \mu)(\underbrace{\mu^{ML} - \mu}_{\substack{\underbrace{\frac{1}{n}x^i + \frac{1}{n}\sum_{j \neq i} x^j}_{\mu^{ML}} - \underbrace{\frac{1}{n}\mu - \frac{n-1}{n}\mu}_{\mu}}})] \\
 &= E[(x^i - \mu) \underbrace{\frac{1}{n}(x^i - \mu)}_{\mu^{ML}}] + E[(x^i - \mu) \underbrace{\sum_{j \neq i} (x^j - \mu) \cdot \frac{1}{n}}_{\mu}] \\
 &= E[\underbrace{(x^i - \mu)^2}_{\sigma^2}] \cdot \frac{1}{n} \quad \left(\begin{array}{l} x^i, x^j \text{ independent} \\ \Downarrow \\ \underbrace{E[(x^i - \mu)]}_{0} \underbrace{\sum_{j \neq i} \underbrace{E[(x^j - \mu)]}_{0} \cdot \frac{1}{n}}_{0} \end{array} \right)
 \end{aligned}$$

$$E[z^2] = \mu^2 + \sigma^2$$

$$\Rightarrow E[(\mu^{ML})^2] = \mu^2 + \frac{\sigma^2}{n}$$

$$z : E[z] = \mu$$

$$\text{Var}(z) = \sigma^2$$

$$(\sigma^2)^{ML} \neq E \left[\frac{1}{n} \sum_{i=1}^n (x^i - \mu)^2 \right] = \sigma^2$$

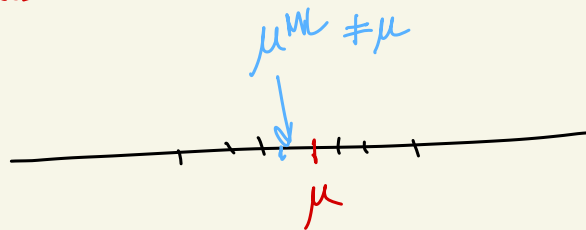
$$E[(\sigma^{ML})^2] = \sigma^2 \left(1 - \frac{1}{n}\right)$$

μ^{ML} "causes" bias

$$\frac{n-1}{n} = 1 - \frac{1}{n}$$

$$\arg \min_a \sum_{i=1}^n (x^i - a)^2 = \frac{1}{n} \sum_{i=1}^n x^i$$

Exercise



What's the point

in practice: $\log \text{ALT}$

$\mu^{ML} \rightarrow \text{Normal}$

$(\sigma^2)^{ML} \rightarrow \text{known distribution } \chi^2$

use $\mu^{ML}, (\sigma^2)^{ML}, n$

$\alpha = \text{confidence level}$

$\alpha \leq 1$

CONFIDENCE INTERVALS

statistical description of uncertainty

Frequentist statistics
Likelihood

$1 - \alpha = \Pr[\text{guess is wrong}]$

Statistics

Bayesian Estimation $\neq$ Philosophy

Bayes' formula

Data $\mathcal{D} = \{(x^i, y^i)\}$

Model family S , $\mathcal{F} = \{P \text{ on } S\}$

Max likelihood $\Rightarrow \theta^{\text{ML}}$, $P_{\theta^{\text{ML}}}$ estimated P

Prior knowledge: $P_{(\theta)}^0 =$ distribution over $\mathcal{F}$
"knowledge before $\mathcal{D}$ is observed"
"prior belief" $\equiv$ distribution over θ parameters

Bayesian estimation ($\equiv$ A PRINCIPLE)

$\Rightarrow P(\theta) =$ posterior distribution
after observing $\mathcal{D}$ and
using $P_{(\theta)}^0$ prior knowledge
"posterior belief"

Example Coin toss

Model Family $S = \{0, 1\}$
 $\mathcal{F} = \{ \theta_1 \in [0, 1], \text{Bernoulli}(\theta_1) \}$

Data $x^{1:n} \in S$

Prior p^0 distribution over $\theta_1 \in [0, 1]$

1) p^0 uniform $\theta_1 \sim \text{unif}[0, 1]$ uninformative prior

wanted $p^{\text{post}} = \Pr[\theta_1 | \mathcal{D}]$ when $\Pr[\theta_1] = p^0 \sim f_{\theta_1}^0$

Bayes' Formula

posterior $f_{\theta_1}(\theta_1 | \mathcal{D}) = \underbrace{f_{\theta_1}^0(\theta_1)}_{\text{Prior}} \underbrace{\prod_{i=1}^n \theta_1^{x_i} (1-\theta_1)^{1-x_i}}_{\text{Likelihood} = \Pr[\mathcal{D} | \theta_1]}$

$\int_0^1 f_{\theta_1}^0(\theta_1) \prod_{i=1}^n \theta_1^{x_i} (1-\theta_1)^{1-x_i} d\theta_1$ $\xrightarrow{\text{normalization}} \int_0^1 \dots d\theta$

(data) observed

posterior $P[A | B] = \frac{P[A] P[B | A]}{\sum_a P[a] P[B | a]}$

variable of interest θ_1

$P[A]$ prior of A
 $P[B | A]$ likelihood $\xrightarrow{\text{evidence}}$ what A predicts about B
 $\sum_a P[a] P[B | a]$ sum over all possible θ_1 's $\equiv$ total probab $\equiv$ independent of θ_1

$$\underline{\underline{f(\theta_1 | \mathcal{D})}} \propto \underbrace{1}_{\text{density of } \theta} \cdot \underbrace{\theta_1^{n_1}}_{\text{unif}[0,1] \text{ prior}} \underbrace{(1-\theta_1)^{n-n_1}}_{\text{likelihood}} = \text{Beta distribution}$$

