

• Rebecca Jörnsten, jornsten@chalmers

(1)

email best mode of communication

• Class MW 13.15-15 - perhaps change later if needed

• Course plan - preliminary only

Masters, PhDs, PhDs other fields

can adapt to some requests, spec. re
modeling section

General idea - { theoretical foundations, emph. on approx
assumptions.
w1-w4

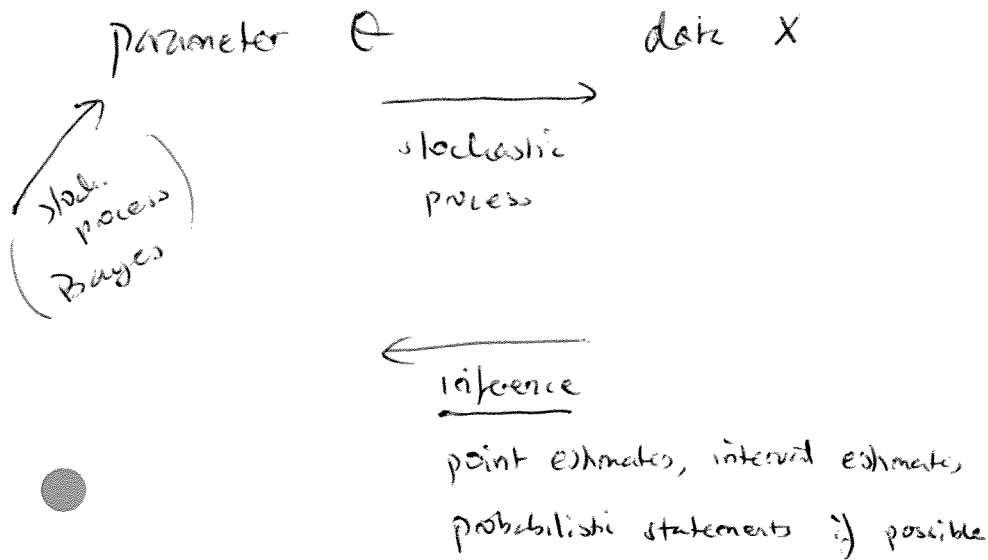
- { w5-w7 - model selection/distances
- likelihood in more complex
data structures

• Today - likelihood, setup & notation, definitions

- Basic likelihood inference

Intro

②



Frequentists

- think of observed data as an instance of something that can be repeated under identical conditions.
- Statistics computed from observed X
⇒ long-run properties of

Example $CI = [\hat{\theta} \pm 1.96 \text{se}(\hat{\theta})]$

covers true θ in 95% of repeated experiments but we have no idea whether our specific interval covers or not.

Bayesian

- inherent uncertainty/variability in real world too
⇒ prior belief $\pi(\theta)$
- observed data ⇒ update of belief into posterior ~~$\pi(\theta|x)$~~

Both prior and posterior - probabilistic statements about θ (3)

$\Rightarrow$ credible intervals $[\hat{\theta} \pm w(\hat{\theta})]$

posterior probability e.g. 95%. That θ is covered by this interval.

Fisher, likelihood based inference

- The likelihood based on observed data x

- $\rightarrow$ provides some idea about uncertainty of θ

- $\Rightarrow$ if possible convert to an exact probabilistic statement
(for some data outcomes we may be able to do so)

- O/w use large-sample approximations, or approximations to particular likelihood function (quadratic = normal more later)

Notation setup

9

- Stoch. process for data generation $p_\theta(x)$ or $f_\theta(x)$ or $p(x; \theta)$
- $X_1^n = \underline{x}$ observed — joint density $p_\theta(x_1, \dots, x_n)$

(could have $\theta_i = \theta(z_i)$ for x_i)
→ GLM

$$p_\theta(x_1, \dots, x_n) = p_\theta(x_1) p_\theta(x_2 | x_1) p_\theta(x_3 | x_1, x_2) \dots$$

/ /
marginal conditional

$$= \prod_{i=1}^n p_\theta(x_i | x_{0:i-1})$$

Markov
1st order

$$\prod_{i=1}^n p_\theta(x_i | x_{i-1})$$



note $p_\theta(x_i)$ different from
other $p_\theta(x_i | x_{i-1})$

⇒ Time series

if X_i 's are independent

Likelihood (joint density)

$$L(\theta | \underline{x}) = \prod_{i=1}^n p_\theta(x_i)$$

Bayes

$$\text{posterior } p(\theta|x) \propto \pi(\theta) L(\theta|x)$$

⑤

direct math
expression for
uncertainty of θ

prior belief

likelihood

As $n \rightarrow \infty$, $L(\theta|x)$ dominates and two schools coincide

For smaller n , $\pi(\theta)$ can dominate so prior belief contributes largely to posterior

$\Rightarrow$ for this reason, Uninformative priors

popular use

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

$\pi(\theta) \sim \text{Uniform on } [0, 1]$

Other popular examples - regularization priors

$$\begin{cases} X \sim N(z'\beta, \sigma^2) & \text{regression model} \\ \pi(\beta) \sim N(0, c) \end{cases} \quad X_i = z_i'\beta + \varepsilon_i, \varepsilon_i \sim N(0, \sigma^2)$$

Small $c \rightarrow \beta$ shrunk toward 0

large $c \rightarrow$ more like standard regression
 $c \rightarrow \infty$ uninformative

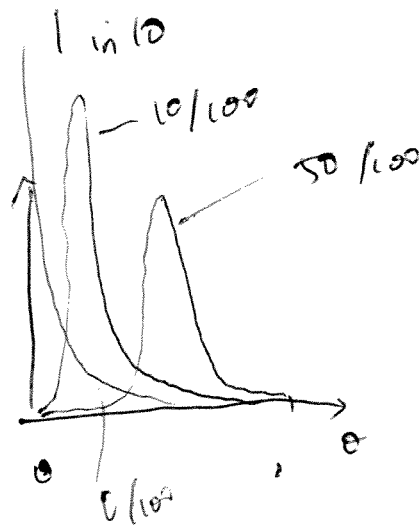
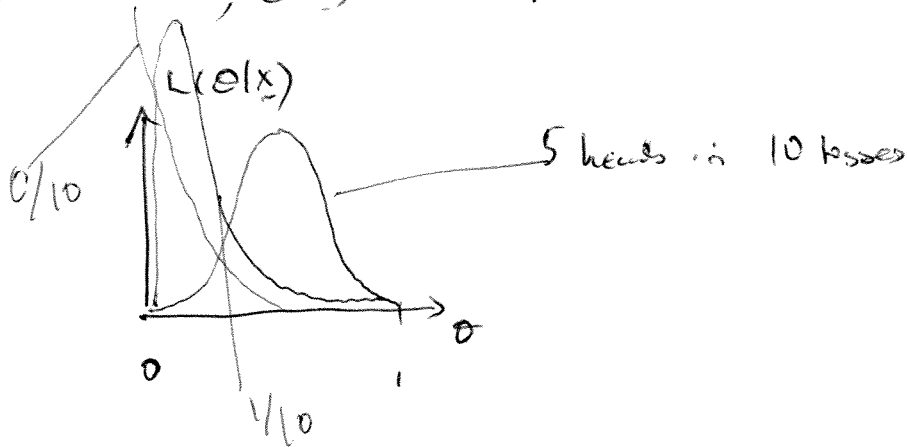
The likelihood function $L(\theta|x)$

⑥

unknown
parameter.

fixed data

In scalar case, 'easy' to plot $L(\theta|x)$ as a function of θ



The likelihood is more 'flat' if data carries limited info about θ .

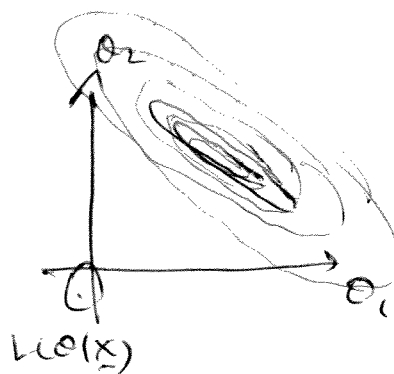
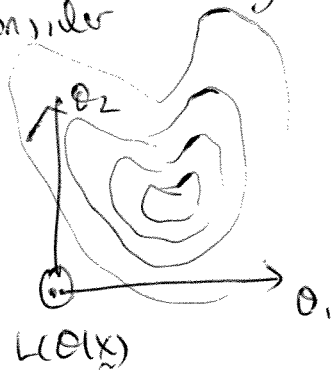
Thinking of $L(\theta|x)$ as measuring likely, data generates θ

- logical choice for estimate is the maximum

(but) Shape of $L(\theta|x)$ is what tells us about the uncertainty of this estimate

In non-scalar cases, not so easy to plot $L(\theta|x)$
and shape can be even more important to

consider = by contours.



redundancy?

Working with the likelihood

Usually on a log-scale

$$\log L(\theta|x) = \ell(\theta|x) = \sum_{i=1}^n \log f_{\theta}(x_i)$$

(can usually be converted
to summary statistics
like $\bar{x}$ and $\bar{x}^2$ etc)

$\Rightarrow$ sufficiency later

So - to draw inference on θ we want
to study $L(\theta|x)$ or $\ell(\theta|x)$

(but) we want to make sure our inference
is not affected by the scale of the data we choose
(perhaps arbitrarily) (e.g. inches vs cm)

2) we transform $X \Rightarrow y$, $x = x(y)$

the change-of-variables $\Rightarrow p_\theta(y) = p_\theta(x(y)) \left| \frac{\partial x}{\partial y} \right|$

$$\Rightarrow L(\theta|y) = L(\theta|x) \left| \frac{\partial x}{\partial y} \right|$$

To remove effect of variable transformation then $\Rightarrow$ look

- at likelihood ratios

- $$\frac{L(\theta|y)}{L(\theta'|y)} = \frac{L(\theta|x)}{L(\theta'|x)}$$

Another very important question is the choice of parameterization $\eta = g(\theta)$

- $\Rightarrow$ data should carry the same amount of info about η as it does about θ .

Invariance

(commonly used parameterization in $x \sim \text{Bin}(M, \theta)$ is the logit $\eta = \log \frac{\theta}{1-\theta}$)

This is the for the likelihood ratio comparison.

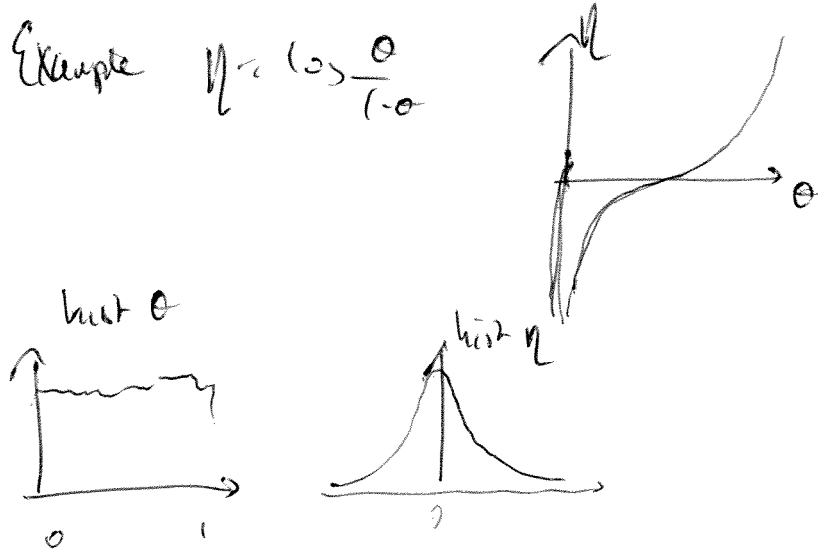
Interesting to note though is if we're Bayesian
— the parametrization does matter

(9)

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

/
if $\pi(\theta)$ uniform $\Rightarrow$ this may not
be the case for $\eta = g(\theta)$!

Example $\eta = \cos \frac{\theta}{1-\theta}$



Summarizing the likelihood function

(10)

→ logical to use the location of the maximum
for our estimate $\hat{\theta}_{MLE}$ = point estimate of θ

→ but provides no measure of uncertainty

Idea - look at how sharp the likelihood function is
around the maximum, i.e. the curvature of $L(\theta|x)$
near $\hat{\theta}$.

Indeed, if $L(\theta|x)$ is quadratic in θ this
would be sufficient to describe the entire function.

We call the (log) likelihood function $l(\theta|x)$

• regular if it is well approximated by

• a quadratic function (in which case $\hat{\theta} +$
the curvature provides
the complete information!)

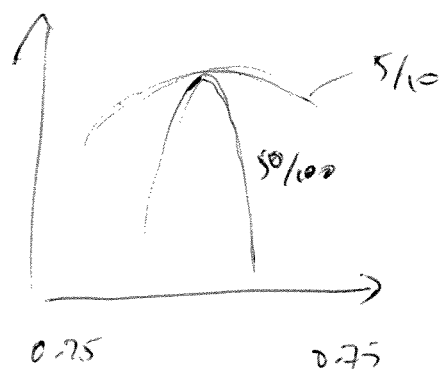
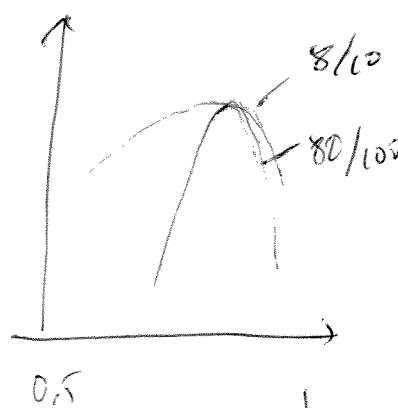
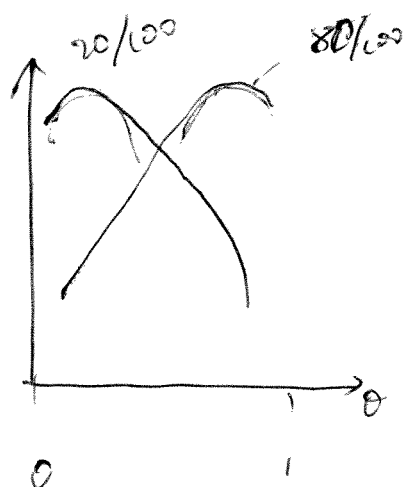
(checking the quadratic approximation

(11)

Plot the normalized loglikelihood

$$\log \frac{L(\theta|x)}{L(\hat{\theta}|x)}$$

(makes things comparable
for different x as well
(max @ $\hat{\theta} = 0$)



As you can see, how well approximated the likelihood can be by a quadratic depends on sample size and outcome.

In general, for large n , approx works well.

Some concepts we will come back to later as well

(12)

Likelihood $L(\theta|x)$ or $L(\theta)$

loglikelihood $\ell(\theta|x)$ or $\ell(\theta)$

The score function $S(\theta) = \frac{\partial}{\partial \theta} \log L(\theta)$

• The MLE $\{\theta: S(\theta) = 0\}$, denote by $\hat{\theta}$

• Curvature $I(\theta) = -\frac{\partial^2}{\partial \theta^2} \log L(\theta)$ (positive quantity)

Curvature @ MLE $I(\hat{\theta})$ - 'measures' the sharpness of the peak
- also called the Observed Fisher Information

(as opposed to the expected

$$E \left[-\frac{\partial^2}{\partial \theta^2} \log L(\theta|x) \right])$$

... more later

So is the quadratic approx OK?

$$\log L(\theta) \approx \log L(\hat{\theta}) + S(\hat{\theta})(\theta - \hat{\theta}) - \frac{1}{2} I(\hat{\theta})(\theta - \hat{\theta})^2$$

(Taylor expansion)

Since $S(\hat{\theta}) = 0$ by def

$$\Rightarrow \log \frac{L(\theta)}{L(\hat{\theta})} \approx -\frac{1}{2} I(\hat{\theta}) (\theta - \hat{\theta})^2$$

(Special case - normal mean model

$$p_{\mu}(x) = \frac{1}{\sqrt{2\pi}} \exp\left\{-\frac{1}{2}(x-\mu)^2\right\}$$

$$\log L(\mu) \propto -\frac{1}{2}(x-\mu)^2 \quad \text{exactly quadratic}$$

Normal approx of $\hat{\theta} \approx$ Quad approx of $L(\theta)$ near $\hat{\theta}$.

$$\frac{d}{d\theta} \log \frac{L(\theta)}{L(\hat{\theta})} = \boxed{S(\theta) \approx I(\hat{\theta})(\hat{\theta} - \theta)}$$

Now, for nice distribution families (more later)

$$\text{One can show } E_{\theta}[S(\theta)] = E_{\theta}\left[\frac{d}{d\theta} \log L(\theta|X)\right] = 0$$

$$\text{and } V_{\theta}[S(\theta)] = E[I(\theta)]$$

$$\Rightarrow \underbrace{I^{-1/2}(\hat{\theta})}_{\text{dimensionless}} S(\theta) \approx I^{-1/2}(\hat{\theta}) (\hat{\theta} - \theta)$$

$$\Rightarrow \text{Var}(\hat{\theta}) \approx I^{-1}(\hat{\theta})$$

[Later, go through this in more rigour, CR lower bound etc]

Summary

14

γ $l(\theta)$ approx quadratic near $\hat{\theta}$

$\Rightarrow$ can summarize $l(\theta)$ by the
MLE $\hat{\theta}$ and the curvature $I(\hat{\theta})$

And Have
$$\begin{cases} E(\hat{\theta}) = \theta \\ V(\hat{\theta}) \approx I(\hat{\theta})^{-1} \end{cases}$$

Alternative - working with the likelihood function directly.

Values of θ that ^{are} supported by the data - large values
for $L(\theta)$, or large values of LR $\frac{L(\theta)}{L(\hat{\theta})}$ (normalized L)

Suggests interval estimate $\left\{ \theta : \frac{L(\theta)}{L(\hat{\theta})} > c \right\}$

But how do we choose c ?

Where likelihood at least 5% of the maximum? or 10%?

This is tricky - we have no obvious c s

here since $L(\theta)$ is data dependent not absolute.

$\Rightarrow$ try to convert this to a long-run freq
or probabilistic statement.