

# Frequentist Statistics

## Estimation and Hypothesis testing:

## Sampling distribution of an estimator

$\theta \rightarrow$  parameter       $\hat{\theta}(D) \rightarrow$  estimator  
                                         $\uparrow$   
                                        data

Example:

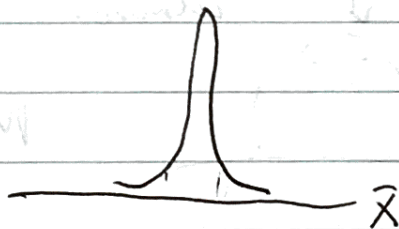


$\mu$  = mean

$$\bar{X} = \frac{X_1 + \dots + X_N}{N}$$

## Parameter

## Estimator



$\bar{X}$  has its own distribution

In general, one can think of datasets  $\mathcal{D}^{(s)}$  from some true

model  $p(\cdot | \theta^*)$ .  $\mathcal{D}^{(s)} = \{x_i^{(s)}\}_{i=1}^N$

Estimator  $\{\hat{\Theta}(D^{(i)})\}_{i=1}^S$  with

$S \rightarrow \infty$  gives us the sampling distribution

## ~~Some~~ <sup>desirable</sup> properties of estimators

~~Q~~ In practice, how do we know what the sampling distribution is?

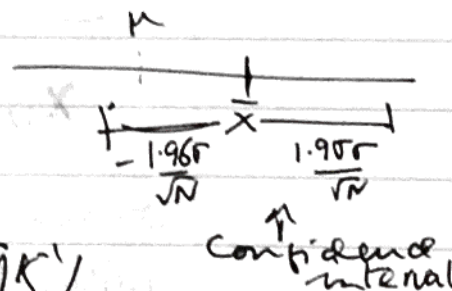
Why is knowing the sampling distribution ~~independent~~ important?

Consider large sample limit of the distribution. If  $X$  is a random variable with finite mean  $\mu$  and variance  $\sigma^2$ ,  $\bar{X}$  is approximately normally distributed with mean  $\mu$  and variance  $\frac{\sigma^2}{N}$ .

$$P\left(\mu - \frac{1.96\sigma}{\sqrt{N}} < \bar{X} < \mu + \frac{1.96\sigma}{\sqrt{N}}\right) = 95\%$$

We can rewrite it as

$$P\left(\bar{X} - \frac{1.96\sigma}{\sqrt{N}} < \mu < \bar{X} + \frac{1.96\sigma}{\sqrt{N}}\right) = 95\%$$



That tells us something about  $\mu$ . If we don't know  $\sigma$  we could replace it by a sample estimate too.

What if we do not have such ~~large~~ large sample theory guarantees?

Bootstrap:



Parametric bootstrap

$$x_i^{(s)} \sim P(\cdot | \theta^*)$$

$$\hat{\theta}^s = f(x_{1:n}^s) \quad \text{estimated parameter}$$

Use  $x_i^{(s)} \sim P(\cdot | \hat{\theta}(\mathcal{D}))$  to generate  $\hat{\theta}^{(s)}$  and its distribution ~~estimator~~.

Non-parametric bootstrap

$$x_1, \dots, x_N$$

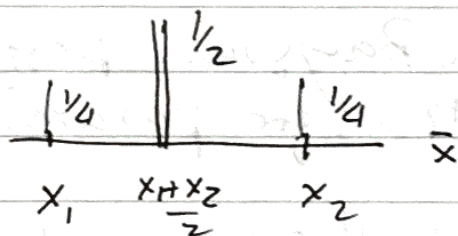
Generate random samples of size  $N$  with ~~replacement~~ replacement.

Each  $x_i$  could now be repeated

How could we ~~we~~ get something almost out of nothing?

Imagine we have only two values  
 $x_1, x_2$

$$\begin{array}{lll} \frac{1}{4} \text{ prob} & x_1, x_1 & \rightarrow \bar{x} = x_1 \\ \frac{1}{2} \text{ prob} & x_1, x_2 & \rightarrow \bar{x} = \frac{x_1 + x_2}{2} \\ \frac{1}{4} \text{ prob} & x_2, x_2 & \rightarrow \bar{x} = x_2 \end{array}$$



Sampling  
 Dist<sup>n</sup> Var  $\sim \frac{1}{4} \left( x_2 - \frac{x_1 + x_2}{2} \right)^2 + \frac{1}{2} 0 + \frac{1}{4} \left( x_1 - \frac{x_1 + x_2}{2} \right)^2$

$$= \frac{1}{2} \left( \frac{x_2 - x_1}{2} \right)^2 = \frac{1}{8} (x_2 - x_1)^2$$

$$\begin{aligned} S_{N=2}^2 &= \frac{\left( x_2 - \frac{x_1 + x_2}{2} \right)^2 + \left( x_1 - \frac{x_1 + x_2}{2} \right)^2}{2} \\ &= \frac{1}{4} (x_1 - x_2)^2 \end{aligned}$$

Essentially we are sampling from the

empirical distribution:



$$\hat{\Theta}^s = \hat{\Theta}(x_{1:N}^s)$$

↑  
Bootstrap samples

Contrast with Bayesian posterior probabilities for parameters.

As we will see local sampling of  $\theta$  by MC is more efficient than having to <sup>find</sup> many samples (Sinh) and extract  $\hat{\theta}$  for each sample.

# Large sample theory for the Maximum Likelihood Estimator (MLE)

Log likelihood function

$$\log P(D|\theta)$$

$$\hat{\theta} = \underset{\theta}{\operatorname{argmax}} \log P(D|\theta)$$

Example  $x_i \sim N(\mu, \sigma^2)$

$$P(D|\theta) = \frac{1}{(2\pi\sigma^2)^{\frac{N}{2}}} e^{-\frac{\sum_{i=1}^N (x_i - \mu)^2}{2\sigma^2}}$$

$$\log P = -\frac{N}{2} \ln(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - \mu)^2$$

Vary  $\mu$   $\sum_{i=1}^N (x_i - \hat{\mu}) = 0$

$$\Rightarrow \hat{\mu} = \frac{\sum_{i=1}^N x_i}{N} = \bar{x}$$

Vary  $\sigma^2$   $-\frac{N}{2} \frac{1}{\sigma^2} + \frac{1}{2(\sigma^2)^2} \sum_{i=1}^N (x_i - \hat{\mu})^2$

$$\hat{\sigma}^2 = \frac{\sum_{i=1}^N (x_i - \hat{\mu})^2}{N} = \frac{\sum_{i=1}^N (x_i - \bar{x})^2}{N} = S_N^2$$

Exponential distribution

$$\frac{1}{b} e^{-\frac{x}{b}} = p(x)$$

$$P(D|b) = \frac{1}{b^N} e^{-\frac{\sum_{i=1}^N x_i}{b}}$$

$$\log P = -N \ln b - \frac{\sum_{i=1}^N x_i}{b}$$

$$\frac{\partial \log P}{\partial b} \bigg|_{b=\hat{b}} = 0$$

$$-\frac{N}{\hat{b}} + \frac{\sum_{i=1}^N x_i}{\hat{b}^2} = 0$$

$$\hat{b} = \frac{\sum_{i=1}^N x_i}{N} = \bar{x}$$

In general

$$s(\theta) = \nabla_{\theta} \log P(D|\theta)$$

Score function

In observed information matrix For MLE  $s(\hat{\theta}, D) = 0$

$$J(\hat{\theta}) = -\nabla s(\theta) \big|_{\theta=\hat{\theta}}$$

$$= -\nabla_{\theta} \nabla_{\theta} \log P(D|\theta)$$

Hessian of log-likelihood

The <sup>matrix</sup>  $J$  tell us how peaked things are in the theta space.

Let us go back to easy example of normal distribution and mean. Assume  $\sigma$  is fixed.

$$\log P(D|\mu) \sim - \frac{\sum (x_i - \mu)^2}{2\sigma^2}$$



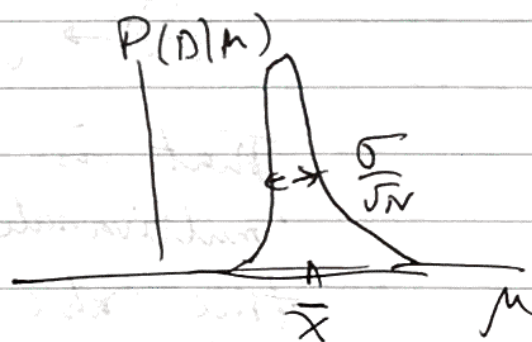
$$S(\mu) = \frac{N}{\sigma^2} (\bar{x} - \mu)$$

$$J(\mu) = - \frac{d}{d\mu} S(\mu) = \frac{N}{\sigma^2}$$

$$P(D|\mu) \sim e^{-\frac{\sum (x_i - \mu)^2}{2\sigma^2}}$$

$$\sim e^{-\frac{\sum (x_i - \bar{x})^2}{2\sigma^2}} e^{-\frac{N(\bar{x} - \mu)^2}{2}}$$

$$\propto e^{-\frac{N(\bar{x} - \mu)^2}{2}}$$



Things are a bit more complicated if we look at the ~~joint variation~~  $\mu, \sigma^2$

~~Plot of~~  
 or Likelihood profile for one sample  $(x_1, \dots, x_N)$   
 $(\hat{\mu}, \hat{\sigma}^2)$  ← True parameters

Average of different samples for true parameters of the observed information matrix  $\rightarrow$  Fisher information

$$I(\hat{\theta} | \theta^*) \triangleq E_{\theta^*} [J(\hat{\theta})]$$

We are ~~still~~ still keeping  $\hat{\theta}$  and  $\theta$  independent, Not using some estimator function of the sample  $x$ 's. Sometimes written as  $I_N(\theta)$ .

Since  $\hat{\theta}$  is built from additive contribution from independent samples after averaging  $I(\hat{\theta}) = NI(\tilde{\theta}) \triangleq NI(\hat{\theta})$   
 $J = -N \frac{\partial}{\partial \theta} \log p(x | \theta) \big|_{\theta = \theta^*}$ ,  $E_{\theta^*}(J) = N \frac{\partial}{\partial \theta} \frac{\partial}{\partial \theta} \log p(x | \theta) \big|_{\theta = \theta^*}$   
 For  $\hat{\theta} = \hat{\theta}_{MLE}(\theta)$

$$\hat{\theta} \rightarrow \mathcal{N}(\theta^*, I_N(\theta^*)^{-1}) = \mathcal{N}\left(\theta^*, \frac{I(\theta^*)^{-1}}{N}\right)$$

That is  $\hat{\theta}$  is a (multivariate) normal with mean at true value and the covariance matrix given by inverse of Fisher information

Note that that the covariance matrix  $\rightarrow 0$  as  $N \rightarrow \infty$

Let us see how that comes out to be. Do it for one parameter first, with an example, the exponential distribution

$$\log P(D|b) = \log \frac{1}{b^N} e^{-\frac{\sum x_i}{b}} = -N \log b - \frac{\sum x_i}{b}$$

$$\partial_{\log} P = S(b, D) = -\frac{N}{b} + \frac{\sum x_i}{b^2} = \sum_{i=1}^N \frac{x_i - b}{b^2}$$

Sum of  $N$  independent random variables

$$J = -\frac{\partial}{\partial b} S = -\frac{N}{b^2} + \frac{2 \sum x_i}{b^3}$$

$$= \sum_i \frac{2x_i - b}{b^3} \quad \leftarrow \text{Sum of } N \text{ independent random variables.}$$

If  $x_i$  are generated from the distribution  $\frac{1}{b} e^{-x_i/b}$ ,  $x_i \geq 0$   $E[x_i] = b$

$$\text{So } I_N(b|b^*) = \sum_i \frac{2E[x_i] - b}{b^3} = \frac{N(2b^* - b)}{b^3}$$

~~$$I(b) = \frac{1}{b^3}$$~~

Now let us look at  $\hat{b}(D)$  distribution when  $\hat{b}(D)$  is the MLE.

$\hat{b}$  is decided by the equation

~~$$\frac{\partial}{\partial b} \log p(x_i | b) = 0$$~~

$$\frac{\partial}{\partial b} \log p(x_i | b) = 0 \text{ or } \sum_i \frac{x_i - b}{b^2} = 0$$

~~If~~ course, we could write  $\hat{b} = \frac{1}{n} \sum x_i$  and analyze it from there

However, here is general procedure

$$\sum_i \frac{\partial}{\partial b} \log p(x_i | b) = 0 \quad \rightarrow \text{Taylor expand:}$$

$$0 \approx \sum_i \left. \frac{\partial}{\partial b} \log p(x_i | b) \right|_{b^*} + (b - b^*) \left[ \sum_i \left. \frac{\partial^2}{\partial b^2} \log p(x_i | b) \right|_{b^*} \right] + o((b - b^*)^2)$$

In particular,

$$0 = \sum_i \frac{x_i - b^*}{b^{*2}} + (b - b^*) \left[ -\sum_{i=1}^N \frac{(2x_i - b^*)}{b^{*3}} \right] + o((b - b^*)^2)$$

Consider the random variable  $S(b, D) = \sum \partial_b \log p(x_i | b)$ . It is a sum of i.i.d variables. We need to know its <sup>mean</sup> ~~expectation~~ and Variance. What is its expectation?

$$\begin{aligned} E_{b^*}[S(b, D)] &= N \partial_b \int \log(p(x|b)) p(x|b^*) dx \\ &= N \int \partial_b p(x|b) \frac{p(x|b^*)}{p(x|b)} dx \end{aligned}$$

Note that ~~the~~  $E_{b^*}[S(b^*, D)]$

$$= N \int \partial_b p(x|b) \Big|_{b^*} dx$$

But  $\int p(x|b) dx = 1$  for every  $b$ ,  
~~so~~ so that  $\int \partial_b p(x|b) dx = 0$

So  $E_{b^*}[S(b^*, D)] = 0$

What is  $\text{Var}_{b^*}[S(b^*, D)]$

~~$$\text{Var}_b[S] = N \text{Var}_b \left[ \frac{1}{N} \sum_{i=1}^N \log p(x_i|b) \right]$$

$$\text{Var}_b[S] = N \text{Var}_b \left[ \frac{1}{N} \sum_{i=1}^N \log p(x_i|b) \right]$$~~

$$\text{Var}_{b^*} [S(b^*, D)] = \left[ \int \left( \partial_b \log p(x|b) \right)^2 p(x|b) dx \right]_{b=b^*}$$

$$= \int \frac{(\partial_b p(x|b))^2}{p(x|b)} dx \Big|_{b=b^*}$$

Remember for two distributions  
 $p_i$  and  $p_i + \delta p_i$  KL divergence

$$= \frac{1}{2} \sum_i \frac{\delta p_i^2}{p_i}$$

Also  $-\int \partial_b^2 [\log p(x|b)] p(x|b) dx$

$$= - \int \partial_b \left( \frac{\partial_b p(x|b)}{p(x|b)} \right) p(x|b) dx$$

$$= - \int \left[ \frac{\partial_b^2 p(x|b)}{p(x|b)} - \frac{\partial_b p(x|b)}{p(x|b)^2} \right] p(x|b) dx$$

$$= -\partial_b^2 \int p(x|b) dx + \int \frac{(\partial_b p(x|b))^2}{p(x|b)} dx$$

zero

First derivative square  
an

$$\text{Var}_{b^*}[S(b^*, D)] = N \int (\partial_b \log p(x|b))^2 p(x|b) dx \Big|_{b^*}$$

$$= -N \int \partial_b^2 [\log p(x|b)] p(x|b) dx \Big|_{b^*}$$

Second deriv ~~version~~ an.

Any way, going back to <sup>the</sup> Taylor expanded equation.

$$0 = S(b^*, D) + (b - b^*) \left[ \sum_i \partial_b \log(x_i|b) \right]_{b^*} + \dots$$

Sum of many variables. Usually non zero average. Replace by

$$N \int dx p(x|b^*) \left[ \partial_b^2 \log p(x|b) \right]_{b=b^*}$$

That is just  $-NI(b^*)$

So  $S(b^*, D) = N(b - b^*) I(b^*) + \dots$

$$\hat{b}(D) - b^* = \frac{1}{N} \underbrace{I(b^*)^{-1}}_{\substack{\uparrow \\ \text{mean zero}}} S(b^*, D)$$

Variance  $\rightarrow \frac{1}{N^2} I(b^*)^{-2} N I(b^*)$   
 $= I(b^*)^{-1} / N$

Thus  $\hat{b}(D) \sim N(b^*, \frac{I(b^*)^{-1}}{N})$

In detail

$$\sum_i \frac{x_i - b^*}{b^{*2}} \approx \sum_i \frac{2x_i - b^*}{b^{*3}} (\hat{b} - b^*)$$

$$\approx \frac{N}{b^{*2}} (\hat{b} - b^*)$$

$$E\left[\frac{X_i - b^*}{b^{*2}}\right] = 0$$

$$\text{Var}\left[\sum_i \frac{X_i - b^*}{b^{*2}}\right] = N \text{Var}\left[\frac{X - b^*}{b^{*2}}\right]$$

$$= N \frac{\text{Var}[X]}{b^{*4}} = N \frac{b^{*2}}{b^{*4}} = \frac{N}{b^{*2}}$$

$$\hat{b} = b^* + \frac{b^{*2}}{N} \sum \frac{X_i - b^*}{b^{*2}}$$

$$= b^* + \frac{1}{N} \sum_i (X_i - b^*)$$

 mean zero

variance  $\frac{b^{*2}}{N}$

$$\hat{b} \sim \mathcal{N}\left(b^*, \frac{b^{*2}}{N}\right)$$

\* In the multiparameter case

$$S(\theta, D) = \sum_i \nabla_{\theta} \log p(x_i | \theta)$$

Sum of independent random vectors.

$$E_{\theta^*}[S(\theta^*, D)] = 0$$

$$\begin{aligned} \text{Cov}_{\theta^*}[S(\theta^*, D)] &= \cancel{N \sum_i \nabla_{\theta} \log p(x_i | \theta^*) \nabla_{\theta} \log p(x_i | \theta^*)} \\ &= N \int dx p(x | \theta^*) [\nabla_{\theta} \log p(x | \theta) \nabla_{\theta} \log p(x | \theta)]_{\theta = \theta^*} \\ &= -N \int dx p(x | \theta^*) [\nabla_{\theta} \nabla_{\theta} \log p(x | \theta)]_{\theta = \theta^*} \\ &= N I_{\theta^*}(\theta^*) \end{aligned}$$

\* Taylor expanded equation  
 $S(\theta, D) = 0$

$$\Rightarrow S(\theta^*, D) + \sum_i \nabla_{\theta} \nabla_{\theta} \log p(x_i | \theta) \big|_{\theta = \theta^*} (\theta - \theta^*) = 0$$

Replace by  $-N I_{\theta^*}(\theta^*)$

$$\text{So } \underbrace{\theta - \theta^*}_{\text{vector}} = \underbrace{\frac{1}{N} \mathbf{I}(\theta^*)^{-1}}_{\text{Mean zero}} s(\theta^*, D)$$

$$\text{Cov}_{\theta^*} = \frac{1}{N} \mathbf{I}(\theta^*)^{-1} \mathbf{I}(\theta^*) \mathbf{I}(\theta^*)^{-1} = \frac{1}{N} \mathbf{I}(\theta^*)^{-1}$$

$$\text{means } \theta - \hat{\theta} \sim \frac{1}{\sqrt{N}}$$

We will also see that the log likelihood at the MLE ~~estimate~~ point goes

as ~~hyper~~ a constant ~~plus~~ minus half a  $\chi^2$  variable.

To see that expand in Taylor Series log likelihood directly around  $\theta = \theta^*$

$$\begin{aligned} \log P(D|\theta) &= \log P(D|\theta^*) + \left[ \nabla_{\theta} \log P(D|\theta) \right]_{\theta=\theta^*}^T (\theta - \theta^*) \\ &\quad + \frac{1}{2} (\theta - \theta^*)^T \left[ \nabla_{\theta} \nabla_{\theta} \log P(D|\theta) \right]_{\theta=\theta^*} (\theta - \theta^*) \\ &\quad + \dots \end{aligned}$$

$$\begin{aligned}
 \log P(D|\theta) & \\
 &\approx \log P(D|\theta^*) \\
 &\quad + S(\theta^*|D)^T (\theta - \theta^*) \\
 &\quad - \frac{1}{2} N (\theta - \theta^*)^T I(\theta^*) (\theta - \theta^*)
 \end{aligned}$$

~~¶~~ If  $\theta = \hat{\theta}$ , the max. likelihood estimator,  $S = N I^{-1}(\hat{\theta}) (\hat{\theta} - \theta^*)$

$$\begin{aligned}
 \log P(D|\theta) &= \log P(D|\theta^*) \\
 &\quad + \frac{N}{2} (\hat{\theta} - \theta^*)^T I(\hat{\theta}) (\hat{\theta} - \theta^*)
 \end{aligned}$$

$$\text{So } 2 \log \frac{P(D|\hat{\theta})}{P(D|\theta^*)} = N (\hat{\theta} - \theta^*)^T I(\hat{\theta}) (\hat{\theta} - \theta^*)$$

This will be important ~~for testing~~ later.

Note that

$$P(\hat{\theta}) \sim e^{-\frac{1}{2} N (\hat{\theta} - \theta^*)^T I(\theta^*) (\hat{\theta} - \theta^*)}$$

If I choose coordinate in the  $\theta$  space, st.  
 $\theta^* \rightarrow 0$        $y \sim O(\theta - \theta^*)$

↑  
orthogonal matrix

so that

$$O I O^T = A = \begin{pmatrix} a_1 & \dots & a_k \end{pmatrix} \quad a_i > 0$$

$$P(\hat{\theta}) \sim e^{-\frac{1}{2} N y^T A y}$$

$$= e^{-\frac{1}{2} \sum_{\alpha} a_{\alpha} y_{\alpha}^2}$$

Define  $z_{\alpha} = \sqrt{N a_{\alpha}} y_{\alpha}$

$$P \sim e^{-\frac{1}{2} \sum_{\alpha} z_{\alpha}^2}$$

$$N(\hat{\theta} - \theta^*)^T I(\theta^*) (\hat{\theta} - \theta^*) = \sum_{\alpha} z_{\alpha}^2$$

the square

This is a sum of  $k$  independent standard normal variables.

Why does max likelihood  
choose the right distribution  
asymptotically?

Consider a discrete distribution.

$$P(D|\theta) = \prod_a p_a(\theta)^{n_a}$$

↑  
events, not samples

$$\sum n_a = N$$

$$\log P(D|\theta) = N \sum_a \frac{n_a}{N} \log p_a(\theta)$$

If events are from  ~~$p_a(\theta)$~~   $p_a(\theta^*)$   
and we have a lot of data

$$\sim N \sum_a p_a(\theta^*) \log p_a(\theta)$$

$$= -N H(p_a(\theta^*), p_a(\theta))$$

↑  
cross entropy

$$KL(P||Q) = \sum p \log \frac{p}{q}$$

$$= H(P, Q) - H(Q)$$

$$- \frac{1}{N} \log P(D|\theta) - H(P(\theta^*))$$

$$\approx KL(P(\theta) || P(\theta^*))$$

$$\delta \theta^T I(\theta) \delta \theta = \int dx \frac{1}{P(x|\theta)} \delta \theta^T \nabla_{\theta} \log P(x) \nabla_{\theta} \log P(x) \delta \theta$$

$$= \int \frac{\delta P(x, \theta)^2}{P(x|\theta)} dx$$

Explicit Example: Normal Distribution

$$\log P(\mathcal{D}|\mu, \sigma^2) = \log \left[ \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{\sum_{i=1}^N (x_i - \mu)^2}{2\sigma^2}} \right]$$

$$= -\frac{N}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - \mu)^2$$

I will treat  $\sigma^2$  as a variable, say  $\theta$

$$\log P = -\frac{N}{2} \log(2\pi\theta) - \frac{1}{2\theta} \sum_{i=1}^N (x_i - \mu)^2$$

$$S(\mu, \theta) = \left( \frac{\partial}{\partial \mu} \log P, \frac{\partial}{\partial \theta} \log P \right)$$

$$= \left( \frac{1}{\theta} \sum_{i=1}^N (x_i - \mu), -\frac{N}{2\theta} + \frac{1}{2\theta^2} \sum_{i=1}^N (x_i - \mu)^2 \right)$$

$$= (S_\mu, S_\theta)$$

$$J(\mu, \theta) = - \begin{pmatrix} \frac{\partial S_\mu}{\partial \mu} & \frac{\partial S_\mu}{\partial \theta} \\ \frac{\partial S_\theta}{\partial \mu} & \frac{\partial S_\theta}{\partial \theta} \end{pmatrix}$$

$$= \begin{pmatrix} \frac{N}{\theta} & \frac{1}{\theta^2} \sum_{i=1}^N (x_i - \mu) \\ \frac{1}{\theta^2} \sum_{i=1}^N (x_i - \mu) & -\frac{N}{2\theta^2} + \frac{1}{\theta^3} \sum_{i=1}^N (x_i - \mu)^2 \end{pmatrix}$$

$$E_{\mu^*, \theta^*} [J(\mu, \theta)] = N \begin{pmatrix} \frac{1}{\theta} & \frac{\mu^* - \mu}{\theta^2} \\ \frac{\mu^* - \mu}{\theta^2} & \frac{2\theta^* - \theta + 2(\mu^* - \mu)^2}{2\theta^3} \end{pmatrix}$$

Since  $\sum_{i=1}^N (x_i - \mu)^2 = \sum_i (x_i - \mu^*)^2 + 2(\mu^* - \mu) \sum_i (x_i - \mu^*) + N(\mu^* - \mu)^2$

At  $\mu = \mu^*, \theta = \theta^*$   $J(\mu^*, \theta^*) = NI(\mu^*, \theta^*)$

$$I(\mu^*, \theta^*) = \begin{pmatrix} \frac{1}{\theta^*} & 0 \\ 0 & \frac{1}{2\theta^{*2}} \end{pmatrix} = \begin{pmatrix} \frac{1}{\sigma^{*2}} & 0 \\ 0 & \frac{1}{2\sigma^{*4}} \end{pmatrix}$$

~~The~~ The max likelihood estimators,

$$S_\mu = 0 \Rightarrow \frac{1}{n} \sum x_i = \hat{\mu}$$

$$S_\theta = 0 \Rightarrow \frac{1}{n} \sum (x_i - \hat{\mu})^2 = \hat{\theta} = \hat{\sigma}^2$$

$\frac{1}{N} \sum x_i = \bar{x}$  ~~The~~ has variance  $\frac{\sigma^2}{N}$

$\frac{1}{N} \sum (x_i - \bar{x})^2$  can be rewritten as  $\frac{1}{n} \sum_{j=1}^{N-1} y_j^2$

where  $y_j \stackrel{i.i.d.}{\sim} N(0, \sigma^2)$ .

$$\begin{aligned}
& \text{Var} \left[ \frac{1}{N} \sum_{i=1}^N (X_i - \bar{X})^2 \right] \\
&= \text{Var} \left[ \frac{1}{N} \sum_{j=1}^{N-1} Y_j^2 \right] \\
&= \sum_{j=1}^{N-1} \frac{1}{N^2} \text{Var}[Y_j^2] \\
&= \sum_{j=1}^{N-1} \frac{1}{N^2} (E[Y_j^4] - E[Y_j^2]^2) \\
&= \frac{N-1}{N^2} (3\sigma^4 - \sigma^4) \\
&= \cancel{\frac{N-1}{N^2}} \frac{2(N-1)\sigma^4}{N^2}
\end{aligned}$$

Asymptotically, this goes as

$$\frac{2\sigma^4}{N}$$

It turns out that  $\bar{X}$  and  $\sum_i (X_i - \bar{X})^2$  are independent  
 So the covariance matrix of  $(\bar{X}, \sigma^2)$  is,

for large  $N$ , given by

$$\begin{pmatrix} \frac{\sigma^2}{N} & 0 \\ 0 & \frac{2\sigma^4}{N} \end{pmatrix}$$

This is just the ~~trans~~ inverse of the Fisher information matrix upto the factor  $N$ .

$$\begin{pmatrix} \frac{N}{\sigma^2} & 0 \\ 0 & \frac{N}{2\sigma^4} \end{pmatrix}$$