

Gov 2001 Section 5: Theory of Maximum Likelihood Estimation

Konstantin Kashin

February 28, 2013

OUTLINE

INTRODUCTION

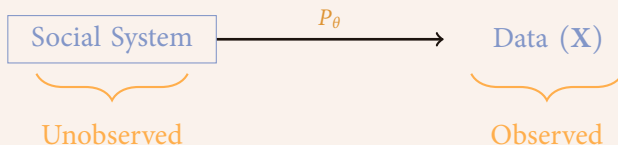
LIKELIHOOD

EXAMPLES OF MLE

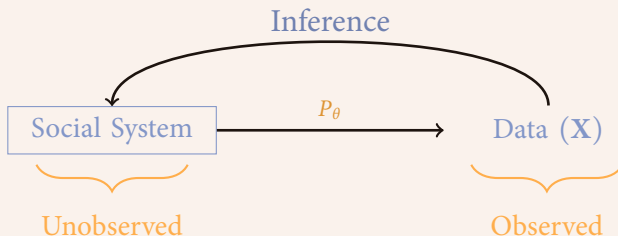
VARIANCE OF MLE

ASYMPTOTIC PROPERTIES

WHAT IS STATISTICAL INFERENCE?



WHAT IS STATISTICAL INFERENCE?



MODEL-BASED (PARAMETRIC) INFERENCE

MODEL-BASED (PARAMETRIC) INFERENCE

- ▶ We assume that the data we observe comes from a model / family of distributions:

$$X \sim f(x|\theta)$$

MODEL-BASED (PARAMETRIC) INFERENCE

- ▶ We assume that the data we observe comes from a model / family of distributions:

$$X \sim f(x|\theta)$$

- ▶ Not right or wrong, but useful, representation of data generating process (DGP)

MODEL-BASED (PARAMETRIC) INFERENCE

- ▶ We assume that the data we observe comes from a model / family of distributions:

$$X \sim f(x|\theta)$$

- ▶ Not right or wrong, but useful, representation of data generating process (DGP)
- ▶ Goal of inference is to use the sample we observe $\mathbf{x} = (x_1, x_2, \dots, x_n)$ to say something about θ (the parameter that completely specifies the DGP) under our model assumptions

MODEL-BASED (PARAMETRIC) INFERENCE

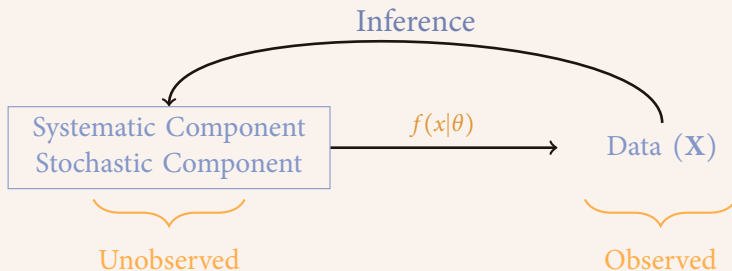
- ▶ We assume that the data we observe comes from a model / family of distributions:

$$X \sim f(x|\theta)$$

- ▶ Not right or wrong, but useful, representation of data generating process (DGP)
- ▶ Goal of inference is to use the sample we observe $\mathbf{x} = (x_1, x_2, \dots, x_n)$ to say something about θ (the parameter that completely specifies the DGP) under our model assumptions

There are two main theories of doing this:
Frequentist (likelihood) and Bayesian

WHAT IS STATISTICAL INFERENCE?



BAYES' RULE

Intuitively, we would like to know the probability density over the unknown parameter θ conditional on the data we observe:

$$\xi(\theta|\mathbf{x})$$

BAYES' RULE

Intuitively, we would like to know the probability density over the unknown parameter θ conditional on the data we observe:

$$\xi(\theta|\mathbf{x})$$

By Bayes' Rule, we can write this probability as:

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})} = \frac{f_n(\mathbf{x}|\theta)\xi(\theta)}{\int_{\Omega} f_n(\mathbf{x}|\theta)\xi(\theta)}, \text{ for } \theta \in \Omega$$

FREQUENTIST / LIKELIHOOD INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

FREQUENTIST / LIKELIHOOD INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We absorb everything that is a constant of the data in $k(x)$:

$$L(\theta|\mathbf{x}) = k(x)f_n(\mathbf{x}|\theta)$$

FREQUENTIST / LIKELIHOOD INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We absorb everything that is a constant of the data in $k(x)$:

$$L(\theta|\mathbf{x}) = k(x)f_n(\mathbf{x}|\theta)$$

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta)$$

FREQUENTIST / LIKELIHOOD INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We absorb everything that is a constant of the data in $k(x)$:

$$L(\theta|\mathbf{x}) = k(x)f_n(\mathbf{x}|\theta)$$

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta)$$

- Parameter θ is an unknown constant

FREQUENTIST / LIKELIHOOD INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We absorb everything that is a constant of the data in $k(x)$:

$$L(\theta|\mathbf{x}) = k(x)f_n(\mathbf{x}|\theta)$$

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta)$$

- ▶ Parameter θ is an unknown constant
- ▶ All fundamental variability (uncertainty) comes from sampling

FREQUENTIST / LIKELIHOOD INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We absorb everything that is a constant of the data in $k(x)$:

$$L(\theta|\mathbf{x}) = k(x)f_n(\mathbf{x}|\theta)$$

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta)$$

- ▶ Parameter θ is an unknown constant
- ▶ All fundamental variability (uncertainty) comes from sampling
- ▶ Everything we know about the parameter based on the data is summarized in the likelihood function

FREQUENTIST / LIKELIHOOD INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We absorb everything that is a constant of the data in $k(x)$:

$$L(\theta|\mathbf{x}) = k(x)f_n(\mathbf{x}|\theta)$$

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta)$$

- ▶ Parameter θ is an unknown constant
- ▶ All fundamental variability (uncertainty) comes from sampling
- ▶ Everything we know about the parameter based on the data is summarized in the likelihood function
- ▶ Focus of inference: characterize likelihood $L(\theta|\mathbf{x})$

FREQUENTIST / LIKELIHOOD INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We absorb everything that is a constant of the data in $k(x)$:

$$L(\theta|\mathbf{x}) = k(x)f_n(\mathbf{x}|\theta)$$

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta)$$

- ▶ Parameter θ is an unknown constant
- ▶ All fundamental variability (uncertainty) comes from sampling
- ▶ Everything we know about the parameter based on the data is summarized in the likelihood function
- ▶ Focus of inference: characterize likelihood $L(\theta|\mathbf{x})$
- ▶ Point summary: maximum likelihood estimate

BAYESIAN INFERENCE

BAYESIAN INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

BAYESIAN INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We can drop the proportionality constant $g_n(\mathbf{x})$ since it's not a function of θ :

$$\underbrace{\xi(\theta|\mathbf{x})}_{\text{posterior}} \propto \underbrace{f_n(\mathbf{x}|\theta)}_{\text{likelihood}} \underbrace{\xi(\theta)}_{\text{prior}}$$

BAYESIAN INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We can drop the proportionality constant $g_n(\mathbf{x})$ since it's not a function of θ :

$$\underbrace{\xi(\theta|\mathbf{x})}_{\text{posterior}} \propto \underbrace{f_n(\mathbf{x}|\theta)}_{\text{likelihood}} \underbrace{\xi(\theta)}_{\text{prior}}$$

- Parameter θ is a latent (unobserved) random variable

BAYESIAN INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We can drop the proportionality constant $g_n(\mathbf{x})$ since it's not a function of θ :

$$\underbrace{\xi(\theta|\mathbf{x})}_{\text{posterior}} \propto \underbrace{f_n(\mathbf{x}|\theta)}_{\text{likelihood}} \underbrace{\xi(\theta)}_{\text{prior}}$$

- ▶ Parameter θ is a latent (unobserved) random variable
- ▶ All fundamental variability (uncertainty) comes from sampling & parameter (through the prior)

BAYESIAN INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We can drop the proportionality constant $g_n(\mathbf{x})$ since it's not a function of θ :

$$\underbrace{\xi(\theta|\mathbf{x})}_{\text{posterior}} \propto \underbrace{f_n(\mathbf{x}|\theta)}_{\text{likelihood}} \underbrace{\xi(\theta)}_{\text{prior}}$$

- ▶ Parameter θ is a latent (unobserved) random variable
- ▶ All fundamental variability (uncertainty) comes from sampling & parameter (through the prior)
- ▶ Probabilities still relative because we don't truly know $\xi(\theta)$

BAYESIAN INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We can drop the proportionality constant $g_n(\mathbf{x})$ since it's not a function of θ :

$$\underbrace{\xi(\theta|\mathbf{x})}_{\text{posterior}} \propto \underbrace{f_n(\mathbf{x}|\theta)}_{\text{likelihood}} \underbrace{\xi(\theta)}_{\text{prior}}$$

- ▶ Parameter θ is a latent (unobserved) random variable
- ▶ All fundamental variability (uncertainty) comes from sampling & parameter (through the prior)
- ▶ Probabilities still relative because we don't truly know $\xi(\theta)$
- ▶ Focus of inference: estimate posterior $\xi(\theta|\mathbf{x})$

BAYESIAN INFERENCE

$$\xi(\theta|\mathbf{x}) = \frac{f(\mathbf{x}|\theta)\xi(\theta)}{g_n(\mathbf{x})}$$

We can drop the proportionality constant $g_n(\mathbf{x})$ since it's not a function of θ :

$$\underbrace{\xi(\theta|\mathbf{x})}_{\text{posterior}} \propto \underbrace{f_n(\mathbf{x}|\theta)}_{\text{likelihood}} \underbrace{\xi(\theta)}_{\text{prior}}$$

- ▶ Parameter θ is a latent (unobserved) random variable
- ▶ All fundamental variability (uncertainty) comes from sampling & parameter (through the prior)
- ▶ Probabilities still relative because we don't truly know $\xi(\theta)$
- ▶ Focus of inference: estimate posterior $\xi(\theta|\mathbf{x})$
- ▶ Point summary: maximum a posteriori (MAP) or posterior mean (PM)

OUTLINE

INTRODUCTION

LIKELIHOOD

EXAMPLES OF MLE

VARIANCE OF MLE

ASYMPTOTIC PROPERTIES

LIKELIHOOD

For an i.i.d. sample of $\mathbf{X} = X_1, \dots, X_n$, we define the likelihood of parameter θ as:

LIKELIHOOD

For an i.i.d. sample of $\mathbf{X} = X_1, \dots, X_n$, we define the likelihood of parameter θ as:

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta) = \prod_{i=1}^n f(x_i|\theta)$$

LIKELIHOOD

For an i.i.d. sample of $\mathbf{X} = X_1, \dots, X_n$, we define the likelihood of parameter θ as:

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta) = \prod_{i=1}^n f(x_i|\theta)$$

Conceptually, $L(\theta|\mathbf{x})$, is a function that assigns a value to each point in parameter space Ω that indicates how likely each value of the parameter is to have generated the data.

LIKELIHOOD

For an i.i.d. sample of $\mathbf{X} = X_1, \dots, X_n$, we define the likelihood of parameter θ as:

$$L(\theta|\mathbf{x}) \propto f_n(\mathbf{x}|\theta) = \prod_{i=1}^n f(x_i|\theta)$$

Conceptually, $L(\theta|\mathbf{x})$, is a function that assigns a value to each point in parameter space Ω that indicates how likely each value of the parameter is to have generated the data.

For a variety of reasons, we work with the log-likelihood:

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

MLE

The MLE is defined as:

$$\hat{\theta}_{MLE} = \max_{\theta \in \Omega} L(\theta | \mathbf{x}) = \max_{\theta \in \Omega} \prod_{i=1}^n f(x_i | \theta)$$

MLE

The MLE is defined as:

$$\hat{\theta}_{MLE} = \max_{\theta \in \Omega} L(\theta|\mathbf{x}) = \max_{\theta \in \Omega} \prod_{i=1}^n f(x_i|\theta)$$

Alternatively, we define the MLE in terms of maximizing the log-likelihood:

$$\hat{\theta}_{MLE} = \max_{\theta \in \Omega} \log L(\theta|\mathbf{x}) = \max_{\theta \in \Omega} \ell(\theta|\mathbf{x}) = \max_{\theta \in \Omega} \sum_{i=1}^n \log(f(x_i|\theta))$$

FINDING THE MLE

FINDING THE MLE

- ▶ **Analytic:** solve first order condition for critical points, then check that second derivative at critical point is negative

FINDING THE MLE

- ▶ **Analytic:** solve first order condition for critical points, then check that second derivative at critical point is negative
 - ▶ Define the score as:

$$s(\theta) = \frac{\partial \ell(\theta|\mathbf{x})}{\partial \theta}$$

FINDING THE MLE

- ▶ **Analytic:** solve first order condition for critical points, then check that second derivative at critical point is negative

- ▶ Define the score as:

$$s(\theta) = \frac{\partial \ell(\theta|\mathbf{x})}{\partial \theta}$$

- ▶ Find critical values by setting score to 0 and solving for θ

FINDING THE MLE

- ▶ **Analytic:** solve first order condition for critical points, then check that second derivative at critical point is negative

- ▶ Define the score as:

$$s(\theta) = \frac{\partial \ell(\theta | \mathbf{x})}{\partial \theta}$$

- ▶ Find critical values by setting score to 0 and solving for θ- ▶ **Numeric:** `optim()` in R

OTHER QUANTITIES WE CAN CALCULATE FROM THE SAMPLE...

OTHER QUANTITIES WE CAN CALCULATE FROM THE SAMPLE...

- Score evaluated at a given θ :

$$S(\theta) = \ell'(\theta|\mathbf{x})$$

OTHER QUANTITIES WE CAN CALCULATE FROM THE SAMPLE...

- Score evaluated at a given θ :

$$S(\theta) = \ell'(\theta|\mathbf{x})$$

- Second derivative of log-likelihood (hessian if multiple parameters) at any θ :

$$\ell''(\theta|\mathbf{x})$$

OTHER QUANTITIES WE CAN CALCULATE FROM THE SAMPLE...

- Score evaluated at a given θ :

$$S(\theta) = \ell'(\theta|\mathbf{x})$$

- Second derivative of log-likelihood (hessian if multiple parameters) at any θ :

$$\ell''(\theta|\mathbf{x})$$

- Observed Fisher information: negation of second derivative at any θ

$$J_n(\theta) = -\ell''(\theta|\mathbf{x}) = -\frac{\partial^2}{\partial \theta^2} \log f(\mathbf{x}|\theta)$$

OTHER QUANTITIES WE CAN CALCULATE FROM THE SAMPLE...

- Score evaluated at a given θ :

$$S(\theta) = \ell'(\theta|\mathbf{x})$$

- Second derivative of log-likelihood (hessian if multiple parameters) at any θ :

$$\ell''(\theta|\mathbf{x})$$

- Observed Fisher information: negation of second derivative at any θ

$$J_n(\theta) = -\ell''(\theta|\mathbf{x}) = -\frac{\partial^2}{\partial \theta^2} \log f(\mathbf{x}|\theta)$$

- Inverse of observed Fisher information: inverse of negation of second derivative at any θ

IN MULTIPLE DIMENSIONS...

If we have multiple parameters (a vector θ of length k):

IN MULTIPLE DIMENSIONS...

If we have multiple parameters (a vector θ of length k):

- Log-likelihood:

$$\ell(\theta|\mathbf{x}) = \log L(\theta|\mathbf{x})$$

IN MULTIPLE DIMENSIONS...

If we have multiple parameters (a vector $\boldsymbol{\theta}$ of length k):

- ▶ Log-likelihood:

$$\ell(\boldsymbol{\theta}|\mathbf{x}) = \log L(\boldsymbol{\theta}|\mathbf{x})$$

- ▶ Score:

$$S(\boldsymbol{\theta}) = \nabla \ell(\boldsymbol{\theta}) = \begin{pmatrix} \frac{\partial \ell(\boldsymbol{\theta})}{\partial \theta_1} \\ \frac{\partial \ell(\boldsymbol{\theta})}{\partial \theta_2} \\ \vdots \\ \frac{\partial \ell(\boldsymbol{\theta})}{\partial \theta_k} \end{pmatrix}$$

IN MULTIPLE DIMENSIONS...

If we have multiple parameters (a vector $\boldsymbol{\theta}$ of length k):

- Log-likelihood:

$$\ell(\boldsymbol{\theta}|\mathbf{x}) = \log L(\boldsymbol{\theta}|\mathbf{x})$$

- Score:

$$S(\boldsymbol{\theta}) = \nabla \ell(\boldsymbol{\theta}) = \begin{pmatrix} \frac{\partial \ell(\boldsymbol{\theta})}{\partial \theta_1} \\ \frac{\partial \ell(\boldsymbol{\theta})}{\partial \theta_2} \\ \vdots \\ \frac{\partial \ell(\boldsymbol{\theta})}{\partial \theta_k} \end{pmatrix}$$

- Observed fisher information:

$$J(\boldsymbol{\theta}) = -\nabla \nabla^T \ell(\boldsymbol{\theta}|\mathbf{x}) = - \begin{pmatrix} \frac{\partial^2}{\partial \theta_1^2} & \frac{\partial^2}{\partial \theta_1 \partial \theta_2} & \cdots & \frac{\partial^2}{\partial \theta_1 \partial \theta_k} \\ \frac{\partial^2}{\partial \theta_2 \partial \theta_1} & \frac{\partial^2}{\partial \theta_2^2} & \cdots & \frac{\partial^2}{\partial \theta_2 \partial \theta_k} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2}{\partial \theta_k \partial \theta_1} & \frac{\partial^2}{\partial \theta_k \partial \theta_2} & \cdots & \frac{\partial^2}{\partial \theta_k^2} \end{pmatrix} \ell(\boldsymbol{\theta}|\mathbf{x})$$

OUTLINE

INTRODUCTION

LIKELIHOOD

EXAMPLES OF MLE

VARIANCE OF MLE

ASYMPTOTIC PROPERTIES

A GEOMETRIC MLE

$X_1 \dots X_n$ form a random sample from a geometric distribution with unknown parameter $0 \leq p \leq 1$. A geometric distribution describes the number of failures before we observe a success in a series of independent Bernoulli trials, each with probability p of a success. We need to find the MLE estimator for p .

A GEOMETRIC MLE

$X_1 \dots X_n$ form a random sample from a geometric distribution with unknown parameter $0 \leq p \leq 1$. A geometric distribution describes the number of failures before we observe a success in a series of independent Bernoulli trials, each with probability p of a success. We need to find the MLE estimator for p .

The pdf of a geometric distribution is:

$$f(x|p) = (1 - p)^x p$$

A GEOMETRIC MLE

$X_1 \dots X_n$ form a random sample from a geometric distribution with unknown parameter $0 \leq p \leq 1$. A geometric distribution describes the number of failures before we observe a success in a series of independent Bernoulli trials, each with probability p of a success. We need to find the MLE estimator for p .

The pdf of a geometric distribution is:

$$f(x|p) = (1-p)^x p$$

Thus, the likelihood for n i.i.d. draws is:

$$L(p|\mathbf{x}) \propto \prod_{i=1}^n (1-p)^{x_i} p = p^n \prod_{i=1}^n (1-p)^{x_i}$$

A GEOMETRIC MLE

$X_1 \dots X_n$ form a random sample from a geometric distribution with unknown parameter $0 \leq p \leq 1$. A geometric distribution describes the number of failures before we observe a success in a series of independent Bernoulli trials, each with probability p of a success. We need to find the MLE estimator for p .

The pdf of a geometric distribution is:

$$f(x|p) = (1-p)^x p$$

Thus, the likelihood for n i.i.d. draws is:

$$L(p|\mathbf{x}) \propto \prod_{i=1}^n (1-p)^{x_i} p = p^n \prod_{i=1}^n (1-p)^{x_i}$$

Taking the log:

$$\ell(p|\mathbf{x}) = n \log(p) + \sum_{i=1}^n x_i \log(1-p) = n \log(p) + (1-p) \sum_{i=1}^n x_i$$

A GEOMETRIC MLE

The log likelihood function is:

$$\ell(p|\mathbf{x}) = n \log(p) + \sum_{i=1}^n x_i \log(1-p) = n \log(p) + (1-p) \sum_{i=1}^n x_i$$

A GEOMETRIC MLE

The log likelihood function is:

$$\ell(p|\mathbf{x}) = n \log(p) + \sum_{i=1}^n x_i \log(1-p) = n \log(p) + (1-p) \sum_{i=1}^n x_i$$

The score, which is the first derivative of the log-likelihood, is:

$$S(p) = \ell'(p|\mathbf{x}) = \frac{n}{p} - \frac{1}{1-p} \sum_{i=1}^n x_i$$

A GEOMETRIC MLE

To find the MLE, we set the score equal to 0 and solve for p :

$$S(p) = \frac{n}{p} - \frac{1}{1-p} \sum_{i=1}^n x_i = 0$$

A GEOMETRIC MLE

To find the MLE, we set the score equal to 0 and solve for p :

$$S(p) = \frac{n}{p} - \frac{1}{1-p} \sum_{i=1}^n x_i = 0$$

$$\frac{n}{p} = \frac{1}{1-p} \sum_{i=1}^n x_i$$

A GEOMETRIC MLE

To find the MLE, we set the score equal to 0 and solve for p :

$$S(p) = \frac{n}{p} - \frac{1}{1-p} \sum_{i=1}^n x_i = 0$$

$$\frac{n}{p} = \frac{1}{1-p} \sum_{i=1}^n x_i$$

A GEOMETRIC MLE

To find the MLE, we set the score equal to 0 and solve for p :

$$S(p) = \frac{n}{p} - \frac{1}{1-p} \sum_{i=1}^n x_i = 0$$

$$\frac{n}{p} = \frac{1}{1-p} \sum_{i=1}^n x_i$$

$$p \sum_{i=1}^n x_i = n - np$$

A GEOMETRIC MLE

To find the MLE, we set the score equal to 0 and solve for p :

$$S(p) = \frac{n}{p} - \frac{1}{1-p} \sum_{i=1}^n x_i = 0$$

$$\frac{n}{p} = \frac{1}{1-p} \sum_{i=1}^n x_i$$

$$p \sum_{i=1}^n x_i = n - np$$

$$p = \frac{n}{n + \sum_{i=1}^n x_i}$$

A GEOMETRIC MLE

To find the MLE, we set the score equal to 0 and solve for p :

$$S(p) = \frac{n}{p} - \frac{1}{1-p} \sum_{i=1}^n x_i = 0$$

$$\frac{n}{p} = \frac{1}{1-p} \sum_{i=1}^n x_i$$

$$p \sum_{i=1}^n x_i = n - np$$

$$p = \frac{n}{n + \sum_{i=1}^n x_i}$$

$$\hat{p}_{MLE} = \frac{1}{1 + \bar{x}}$$

OBSERVED FISHER INFORMATION

First, let's find the second derivative of the log-likelihood function:

$$\ell''(p|\mathbf{x}) = -\frac{n}{p^2} - \frac{1}{(1-p)^2} \sum_{i=1}^n x_i$$

Note that the second derivative captures the steepness of the curvature around the point p . A more negative second derivative implies that the function is more steeply concave down around the point p .

The observed Fisher information is just the negation of the second derivative:

$$J_n(p) = -\ell''(p|\mathbf{x}) = \frac{n}{p^2} + \frac{1}{(1-p)^2} \sum_{i=1}^n x_i$$

OUTLINE

INTRODUCTION

LIKELIHOOD

EXAMPLES OF MLE

VARIANCE OF MLE

ASYMPTOTIC PROPERTIES

VARIANCE OF THE MLE

We are interested in calculating a measure of uncertainty of our MLE.
That is, we are after the following:

VARIANCE OF THE MLE

We are interested in calculating a measure of uncertainty of our MLE. That is, we are after the following:

$$\text{Var}(\hat{\theta}_{MLE})$$

VARIANCE OF THE MLE

We are interested in calculating a measure of uncertainty of our MLE. That is, we are after the following:

$$\text{Var}(\hat{\theta}_{MLE})$$

Conceptually, what is this quantity? We want to understand it conceptually before we can talk about calculating it!

HOW DO WE UNDERSTAND VARIANCE IN FREQUENTIST FRAMEWORK?

HOW DO WE UNDERSTAND VARIANCE IN FREQUENTIST FRAMEWORK?

In terms of drawing infinite samples with sample size n from the distribution of interest!

HOW DO WE UNDERSTAND VARIANCE IN FREQUENTIST FRAMEWORK?

In terms of drawing infinite samples with sample size n from the distribution of interest!

Specifically, we can think of our sample as one of many possible samples we can draw from our population:

HOW DO WE UNDERSTAND VARIANCE IN FREQUENTIST FRAMEWORK?

In terms of drawing infinite samples with sample size n from the distribution of interest!

Specifically, we can think of our sample as one of many possible samples we can draw from our population:

Samples from Pop.					
	1	2	3	...	∞
X_1	$X_{1,1}$	$X_{1,2}$	$X_{1,3}$	$\cdots$	$x_{1,\infty}$
X_2	$X_{2,1}$	$X_{2,2}$	$X_{2,3}$	$\cdots$	$x_{2,\infty}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
X_n	$X_{n,1}$	$X_{n,2}$	$X_{n,3}$	$\cdots$	$x_{n,\infty}$
$\hat{\theta}_{MLE}$	$\hat{\theta}_{MLE,1}$	$\hat{\theta}_{MLE,2}$	$\hat{\theta}_{MLE,3}$	$\cdots$	$\hat{\theta}_{MLE,\infty}$
$J_n(\hat{\theta}_{MLE})$	$J_{n,1}(\hat{\theta}_{MLE})$	$J_{n,2}(\hat{\theta}_{MLE})$	$J_{n,3}(\hat{\theta}_{MLE})$	$\cdots$	$J_{n,\infty}(\hat{\theta}_{MLE})$

HOW DO WE UNDERSTAND VARIANCE IN FREQUENTIST FRAMEWORK?

	Samples from Pop.				
	1	2	3	...	∞
X_1	$X_{1,1}$	$X_{1,2}$	$X_{1,3}$	$\cdots$	$x_{1,\infty}$
X_2	$X_{2,1}$	$X_{2,2}$	$X_{2,3}$	$\cdots$	$x_{2,\infty}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
X_n	$X_{n,1}$	$X_{n,2}$	$X_{n,3}$	$\cdots$	$x_{n,\infty}$
$\hat{\theta}_{MLE}$	$\hat{\theta}_{MLE,1}$	$\hat{\theta}_{MLE,2}$	$\hat{\theta}_{MLE,3}$	$\cdots$	$\hat{\theta}_{MLE,\infty}$
$J_n(\hat{\theta}_{MLE})$	$J_{n,1}(\hat{\theta}_{MLE})$	$J_{n,2}(\hat{\theta}_{MLE})$	$J_{n,3}(\hat{\theta}_{MLE})$	$\cdots$	$J_{n,\infty}(\hat{\theta}_{MLE})$

We see that $\hat{\theta}_{MLE}$ and $J_n(\hat{\theta}_{MLE})$ are random variables! Thus, they have some theoretical distributions.

$Var(\hat{\theta}_{MLE})$ is just the variance of this theoretical distribution!

EXPECTATIONS OF RANDOM VARIABLES

In fact, the following are all random variables (vary across samples):

EXPECTATIONS OF RANDOM VARIABLES

In fact, the following are all random variables (vary across samples):

- ▶ MLE: $\hat{\theta}_{MLE}$

EXPECTATIONS OF RANDOM VARIABLES

In fact, the following are all random variables (vary across samples):

- ▶ MLE: $\hat{\theta}_{MLE}$
- ▶ Log-likelihood evaluated at MLE (or any θ): $\ell(\hat{\theta}_{MLE}|\mathbf{x})$

EXPECTATIONS OF RANDOM VARIABLES

In fact, the following are all random variables (vary across samples):

- ▶ MLE: $\hat{\theta}_{MLE}$
- ▶ Log-likelihood evaluated at MLE (or any θ): $\ell(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Score evaluated at the MLE (or any θ): $S(\hat{\theta}_{MLE}) = \ell'(\hat{\theta}_{MLE}|\mathbf{x})$

EXPECTATIONS OF RANDOM VARIABLES

In fact, the following are all random variables (vary across samples):

- ▶ MLE: $\hat{\theta}_{MLE}$
- ▶ Log-likelihood evaluated at MLE (or any θ): $\ell(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Score evaluated at the MLE (or any θ): $S(\hat{\theta}_{MLE}) = \ell'(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Observed Fisher information evaluated at MLE (or any θ):
 $J_n(\hat{\theta}_{MLE}) = -\ell''(\hat{\theta}_{MLE}|\mathbf{x})$

EXPECTATIONS OF RANDOM VARIABLES

In fact, the following are all random variables (vary across samples):

- ▶ MLE: $\hat{\theta}_{MLE}$
- ▶ Log-likelihood evaluated at MLE (or any θ): $\ell(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Score evaluated at the MLE (or any θ): $S(\hat{\theta}_{MLE}) = \ell'(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Observed Fisher information evaluated at MLE (or any θ):
 $J_n(\hat{\theta}_{MLE}) = -\ell''(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Inverse observed Fisher information: $J_n(\hat{\theta}_{MLE})^{-1}$

EXPECTATIONS OF RANDOM VARIABLES

In fact, the following are all random variables (vary across samples):

- ▶ MLE: $\hat{\theta}_{MLE}$
- ▶ Log-likelihood evaluated at MLE (or any θ): $\ell(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Score evaluated at the MLE (or any θ): $S(\hat{\theta}_{MLE}) = \ell'(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Observed Fisher information evaluated at MLE (or any θ):
 $J_n(\hat{\theta}_{MLE}) = -\ell''(\hat{\theta}_{MLE}|\mathbf{x})$
- ▶ Inverse observed Fisher information: $J_n(\hat{\theta}_{MLE})^{-1}$

We thus often talk about the expectation of these random quantities across infinite samples. We can denote this, for the example of the MLE, as:

$$E[\hat{\theta}_{MLE}]$$

or

$$E_{\theta_0}[\hat{\theta}_{MLE}]$$

VARIANCE OF THE MLE

Recall that the variance of the MLE is:

VARIANCE OF THE MLE

Recall that the variance of the MLE is:

$$\text{Var}(\hat{\theta}_{MLE})$$

VARIANCE OF THE MLE

Recall that the variance of the MLE is:

$$\text{Var}(\hat{\theta}_{MLE})$$

Now that we understand it conceptually, how do we estimate it?

ASYMPTOTIC DISTRIBUTION OF THE MLEs

It can be shown that under certain regularity conditions, the MLE is distributed normally with a mean equal to the true parameter (θ_o) and the variance equal to the inverse of the expected sample Fisher information at the true parameter (denoted as $I_n(\theta_o)$):

$$\hat{\theta}_{MLE} \sim N\left(\theta_o, \underbrace{\left(-E\left[\frac{\partial^2 \ell(\theta|\mathbf{x})}{\partial \theta^2} \middle| \theta = \theta_o\right]\right)^{-1}}_{I_n(\theta_o)}\right)$$

Let's focus on understanding the variance for now...

EXPECTED FISHER INFORMATION

EXPECTED FISHER INFORMATION

- Asymptotically: $\text{Var}(\hat{\theta}_{MLE}) = I_n(\theta_o)$

EXPECTED FISHER INFORMATION

- ▶ Asymptotically: $\text{Var}(\hat{\theta}_{MLE}) = I_n(\theta_o)$
- ▶ What is $I_n(\theta_o)$?

$$E[J_n(\theta_o)] = -E\left[\ell(\theta|\mathbf{x})\right]$$

EXPECTED FISHER INFORMATION

- ▶ Asymptotically: $\text{Var}(\hat{\theta}_{MLE}) = I_n(\theta_o)$
- ▶ What is $I_n(\theta_o)$?

$$E[J_n(\theta_o)] = -E \left[\ell(\theta|\mathbf{x}) \right]$$

- ▶ That is, it's the expectation of the observed Fisher information evaluated at the true parameter θ_o

EXPECTED FISHER INFORMATION

- ▶ Asymptotically: $\text{Var}(\hat{\theta}_{MLE}) = I_n(\theta_o)$
- ▶ What is $I_n(\theta_o)$?

$$E[J_n(\theta_o)] = -E \left[\ell(\theta|\mathbf{x}) \right]$$

- ▶ That is, it's the expectation of the observed Fisher information evaluated at the true parameter θ_o
- ▶ Conceptually, this is the *expected* curvature of the log-likelihood curve (or surface) across repeated samples at the point θ_o (the true parameter)

EXPECTED FISHER INFORMATION

- ▶ Asymptotically: $\text{Var}(\hat{\theta}_{MLE}) = I_n(\theta_o)$
- ▶ What is $I_n(\theta_o)$?

$$E[J_n(\theta_o)] = -E \left[\ell(\theta|\mathbf{x}) \right]$$

- ▶ That is, it's the expectation of the observed Fisher information evaluated at the true parameter θ_o
- ▶ Conceptually, this is the *expected* curvature of the log-likelihood curve (or surface) across repeated samples at the point θ_o (the true parameter)
- ▶ As $n \rightarrow \infty$, the observed Fisher information converges to the expected Fisher information and the $\hat{\theta}_{MLE}$ converges to θ_o

EXPECTED FISHER INFORMATION

- ▶ Asymptotically: $\text{Var}(\hat{\theta}_{MLE}) = I_n(\theta_o)$
- ▶ What is $I_n(\theta_o)$?

$$E[J_n(\theta_o)] = -E \left[\ell(\theta|\mathbf{x}) \right]$$

- ▶ That is, it's the expectation of the observed Fisher information evaluated at the true parameter θ_o
- ▶ Conceptually, this is the *expected* curvature of the log-likelihood curve (or surface) across repeated samples at the point θ_o (the true parameter)
- ▶ As $n \rightarrow \infty$, the observed Fisher information converges to the expected Fisher information and the $\hat{\theta}_{MLE}$ converges to θ_o

In practice, we use the inverse of the observed Fisher information evaluated at the MLE to approximate the true variance of the MLE!

LET'S DO AN EXAMPLE...

LET'S DO AN EXAMPLE...

- Suppose $X \sim \text{Geom}(0.5)$

LET'S DO AN EXAMPLE...

- Suppose $X \sim \text{Geom}(0.5)$
- The true parameter is $p_0 = 0.5$

LET'S DO AN EXAMPLE...

- ▶ Suppose $X \sim \text{Geom}(0.5)$
- ▶ The true parameter is $p_0 = 0.5$
- ▶ We want to estimate $\hat{p}_{MLE}$ and find the uncertainty around it

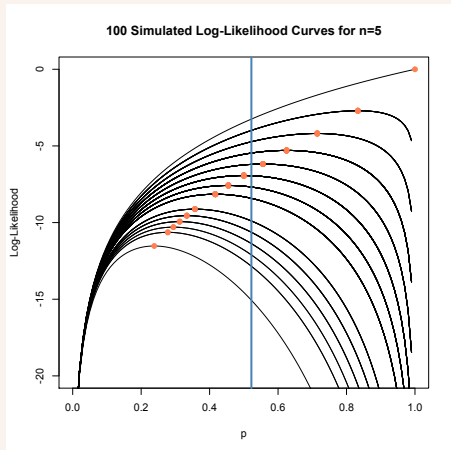
LET'S DO AN EXAMPLE...

- ▶ Suppose $X \sim \text{Geom}(0.5)$
- ▶ The true parameter is $p_o = 0.5$
- ▶ We want to estimate $\hat{p}_{MLE}$ and find the uncertainty around it

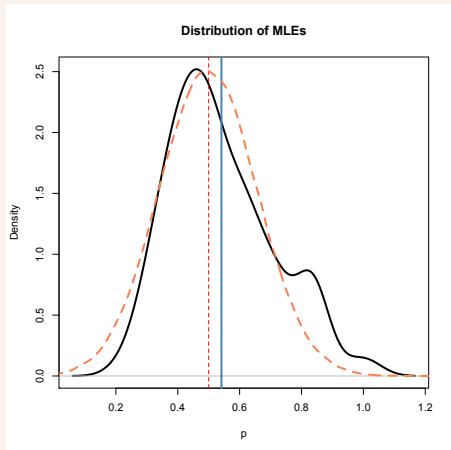
We shall see how as sample size (n) gets larger, $\hat{p}_{MLE} \approx p_o$ and

$$[-\ell''(\hat{p}_{MLE}|\mathbf{x})]^{-1} \approx \text{Var}(\hat{p}_{MLE})$$

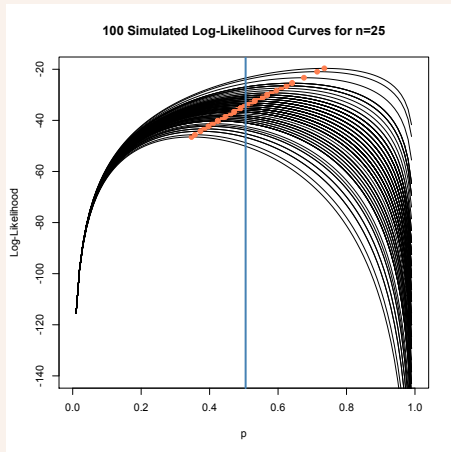
LOG-LIKELIHOODS FOR $n = 5$



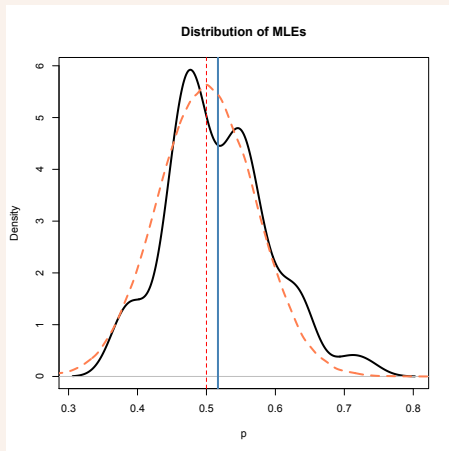
DISTRIBUTION OF MLEs FOR $n = 5$



LOG-LIKELIHOODS FOR $n = 25$



DISTRIBUTION OF MLEs FOR $n = 25$

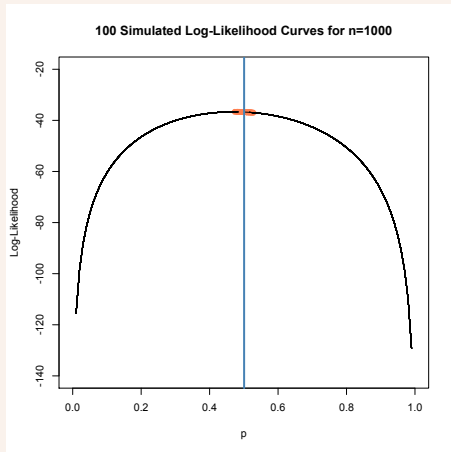


$$\text{Var}_{1000}(\hat{\theta}_{MLE}) = 0.005532$$

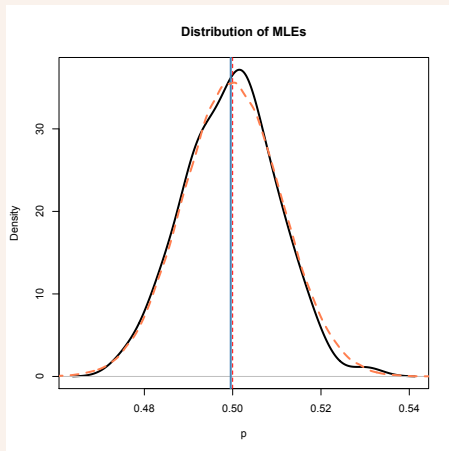
$$J(\hat{\theta}_{MLE})^{-1} = 0.00503$$

$$I(\theta_o)^{-1} = 0.005$$

LOG-LIKELIHOODS FOR $n = 1000$



DISTRIBUTION OF MLEs FOR $n = 1000$



$$\text{Var}_{1000}(\hat{\theta}_{MLE}) = 0.000106$$

$$J(\hat{\theta}_{MLE})^{-1} = 0.000124$$

$$I(\theta_0)^{-1} = 0.000125$$

OUTLINE

INTRODUCTION

LIKELIHOOD

EXAMPLES OF MLE

VARIANCE OF MLE

ASYMPTOTIC PROPERTIES

SUMMARY OF ASYMPTOTIC PROPERTIES

- **Consistency:**

$$\hat{\theta}_{MLE} \xrightarrow{p} \theta_o$$

- **Normality:**

$$\hat{\theta}_{MLE} \sim N\left(\theta_o, \underbrace{\left(-E\left[\frac{\partial^2 \ell(\theta|\mathbf{x})}{\partial \theta^2} \middle| \theta=\theta_o\right]\right)^{-1}}_{I_n(\theta_o)}\right)$$

- **Efficiency:** lowest mean square error amongst asymptotically unbiased estimators

CONSISTENCY

As sample size (n) increases, the MLE ($\hat{\theta}_{MLE}$) converges to the true parameter, θ_o :

$$\hat{\theta}_{MLE} \xrightarrow{p} \theta_o$$

Proof relies upon the uniform law of large numbers.

REGULARITY CONDITIONS FOR CONSISTENCY

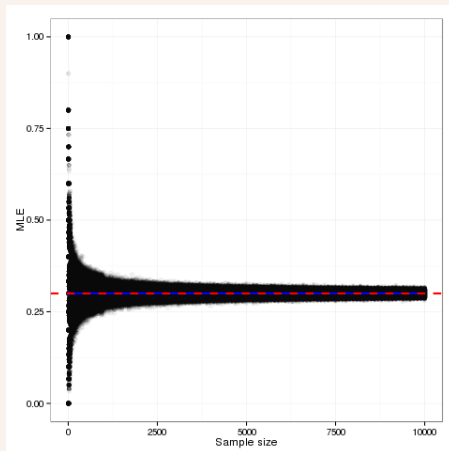
- ▶ Model identification: the true parameter θ_0 is the unique global maximizer of $E_{\theta_0}[\ell(\theta|\mathbf{x})]$
- ▶ Compactness: the parameter space Ω must be a bounded and closed set. That is, Ω must be a *compact subset*.
- ▶ The log-likelihood function $\ell(\theta|\mathbf{x})$ is continuous in θ
- ▶ Note that compactness condition is sufficient but not necessary: it can be replaced with other conditions for non-compact and infinite parameter spaces

VISUALIZING CONSISTENCY OF THE MLE

Simple simulation study to look at properties of MLE:

- ▶ $X \sim \text{Bern}(0.3)$
- ▶ Vary sample size n
- ▶ For each n , simulate 10,000 datasets of size n and calculate MLE and observed Fisher information
- ▶ Take mean across MLEs for each n and compare to true value of the parameter

VISUALIZING CONSISTENCY OF THE MLE



For each n , 10,000 MLEs are plotted with black dots. The dotted red line denotes the true parameter ($p = 0.3$), while the dotted blue line represents the mean MLE across the 10,000 samples at each value of n .

NORMALITY

The standardized MLE is distributed normally with a mean of 0 and a variance equal to the expected unit Fisher information:

$$\sqrt{n}(\hat{\theta}_{MLE} - \theta_0) \sim N\left(0, \underbrace{\left(-E\left[\frac{\partial^2 \ell(\theta|x)}{\partial \theta^2} \middle| \theta = \theta_0\right]\right)^{-1}}_{I_1(\theta_0)}\right)$$

Proof relies upon the Central Limit Theorem.

REGULARITY CONDITIONS FOR NORMALITY

We need all conditions needed for consistency as well as:

- ▶ The true value of the parameter, θ_o , must be an interior point of the parameter set: $\theta_o \in \text{int}(\Omega)$. Phrased differently, θ_o cannot be on the boundary of the set.
 - ▶ This is equivalent to saying that Ω doesn't depend on θ
 - ▶ A violation of this is the uniform distribution $\text{Unif}[0, \theta]$ which has a biased MLE that is not asymptotically normally distributed (see Lehmann and Casella 1998)
- ▶ The likelihood function (or probability distribution $f(x|\theta)$) is continuously twice-differentiable in the neighborhood of θ_o
- ▶ Fisher information matrix exists, is non-singular, and finitely bounded
- ▶

$$\left| \frac{d^j \ln f(x|\theta)}{d\theta^j} \right| \leq h_j(x), \text{ for } j = 1, 2, 3$$

$$E[h_j(y)] < \infty, j = 1, 2$$

$$E[h_3(x)] \text{ does not depend on } \theta$$

NORMALITY

Another way to phrase asymptotic normality is:

As sample size (n) increases, the MLE is normally distributed with a mean equal to the true parameter (θ_o) and the variance equal to the inverse of the expected sample Fisher information at the true parameter (denoted as $I_n(\theta_o)$):

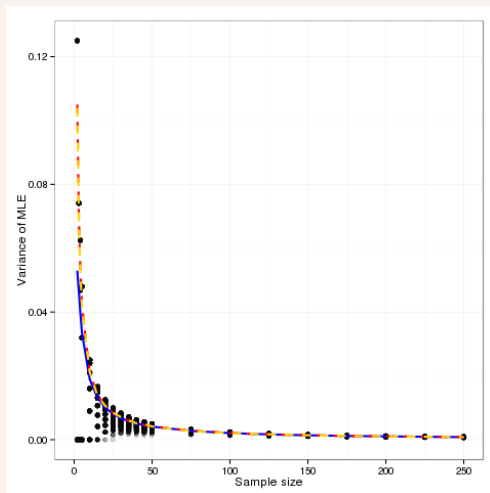
$$\hat{\theta}_{MLE} \sim N\left(\theta_o, \underbrace{\left(-E\left[\frac{\partial^2 \ell(\theta|\mathbf{x})}{\partial \theta^2} \middle| \theta = \theta_o\right]\right)^{-1}}_{I_n(\theta_o)}\right)$$

NORMALITY

However, using the consistency property of the MLE and observed sample Fisher information, we can use the inverse of the observed sample Fisher information evaluated at the MLE, denoted as $J_n(\hat{\theta}_{MLE})$ to approximate the variance:

$$\hat{\theta}_{MLE} \sim N\left(\theta_o, \underbrace{\left(-\left[\frac{\partial^2 \ell(\theta|\mathbf{x})}{\partial \theta^2}\right]_{\theta=\hat{\theta}_{MLE}}\right)^{-1}}_{J_n(\hat{\theta}_{MLE})}\right)$$

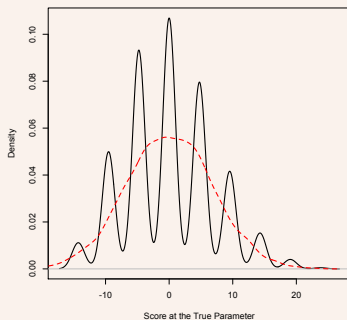
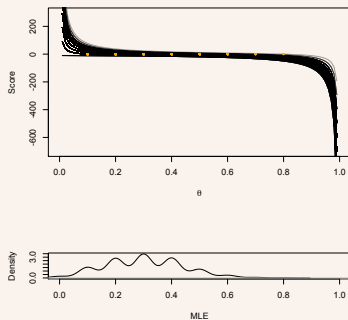
CONSISTENCY OF THE OBSERVED FISHER INFORMATION



Average observed Fisher information in blue, expected Fisher information in red, and simulated variance across the 10,000 MLEs at each n in gold.

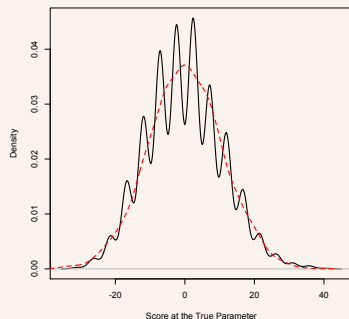
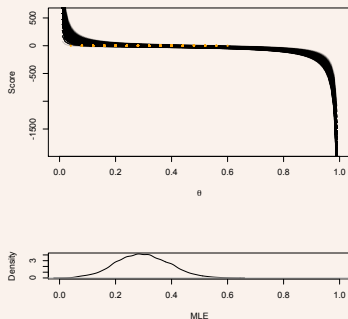
VISUALIZING SCORE FUNCTIONS AND NORMALITY OF MLE

For a sample size of $n = 10$:



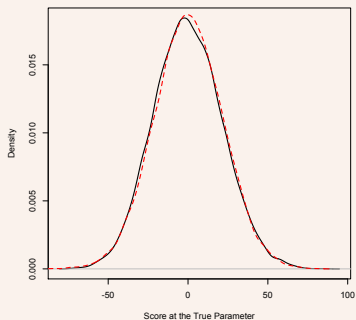
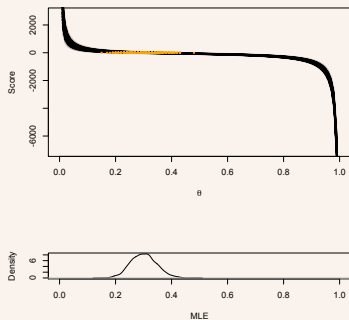
VISUALIZING SCORE FUNCTIONS AND NORMALITY OF MLE

For a sample size of $n = 25$:



VISUALIZING SCORE FUNCTIONS AND NORMALITY OF MLE

For a sample size of $n = 100$:



EFFICIENCY

- ▶ As sample size (n) increases, MLE is the estimation procedure that generally provides the lowest variance (in the class of other consistent and asymptotically normal estimators)
- ▶ ML estimator has a variance equal to Cramer-Rao lower bound