

Likelihood: Efficiency

Patrick Breheny

October 27, 2025

Introduction

- Today we will prove the information inequality, which establishes a lower bound on the variability of an estimator
- This leads to the idea of an “efficient” estimator, as any estimator that achieves this bound can be considered optimal
- We will then see that the MLE is asymptotically efficient, as are Bayesian estimators, and discuss Bayesian asymptotics a bit

Information inequality: 1D

- First, let's take a look at the information inequality in the case of a scalar estimator
- **Theorem (Information inequality):** Let $\hat{\gamma}$ be a statistic with finite expectation $g(\theta) = \mathbb{E}\hat{\gamma}$. Suppose $X_1, X_2, \dots, X_n \sim p(\cdot|\theta^*)$ and $d/d\theta$ can be passed under the integral sign with respect to both $\int dP$ and $\int \hat{\gamma}dP$. Finally, suppose $\mathcal{I}_n(\theta^*) > 0$. Then

$$\mathbb{V}\hat{\gamma} \geq \frac{\dot{g}(\theta^*)^2}{\mathcal{I}_n(\theta^*)}$$

Remarks

- Keep in mind here that p refers to the joint distribution of $X_1, X_2, \dots, X_n$; we are not assuming iid here but we are assuming that the derivative can be passed inside the integral with respect to this joint distribution
- Furthermore, note that this is *not* an asymptotic theorem – it is an inequality that is true for all values of n (including $n = 1$)

Attainment

- Is it possible for estimators to achieve this bound? (i.e., to have the minimum possible variance?)
- An interesting theorem due to Wijsman (1973) is that equality is only possible in the information inequality if $\hat{\gamma}$ is linearly related to the score
- In other words, the only situation in which the lower bound is attainable (for all θ , for all n) is when $\hat{\gamma}$ is the sufficient statistic of an exponential family

Cramér-Rao lower bound

- The information inequality is often restated in terms of the bias of an estimator $\hat{\theta}$ of θ
- Letting $b(\theta) = g(\theta) - \theta$ denote the bias of $\hat{\theta}$, and assuming we have an iid sample, then the information inequality becomes

$$\mathbb{V}\hat{\theta} \geq \frac{(1 + \dot{b}(\theta^*))^2}{n\mathcal{I}(\theta^*)}$$

or, in the case of an unbiased estimator,

$$\mathbb{V}\hat{\theta} \geq \frac{1}{n\mathcal{I}(\theta^*)}$$

- In this form, the inequality is known as the *Cramér-Rao lower bound*

Remarks

- Recall that the mean squared error of an estimator is

$$\begin{aligned}\text{MSE} &= \mathbb{E}\{(\hat{\theta} - \theta^*)^2\} \\ &= \text{Bias}^2 + \text{Var}\end{aligned}$$

- Thus, among unbiased estimators, the CRLB represents the minimum possible MSE
- However, this requirement is rather artificial: it is often the case that biased estimators can be constructed with a lower MSE than the best unbiased estimator

Example #1

- The CRLB is not always attainable
- For example, if $X_i \stackrel{\text{iid}}{\sim} N(\mu, \sigma^2)$, the CRLB for σ^2 is $2\sigma^4/n$
- It turns out that this bound is unobtainable if μ is unknown; all unbiased estimators have a higher variance than this
- For example, letting s^2 represent the usual unbiased estimator of the variance,

$$\mathbb{V}s^2 = \frac{2\sigma^4}{n-1} > \frac{2\sigma^4}{n}$$

Example #2

- Keep in mind also that the CRLB only applies when we can pass the derivative under the integral
- One common model for which this cannot be done is $X_i \stackrel{\text{iid}}{\sim} \text{Unif}(0, \theta)$
- In this case, one might think that the CRLB is θ^2/n
- However, $\hat{\theta} = (n+1)X_{(n)}/n$ is an unbiased estimate of θ with

$$\mathbb{V}\hat{\theta} = \frac{\theta^2}{n(n+2)} < \frac{\theta^2}{n}$$

- The “real” CRLB here is not well defined

Information inequality: Multiparameter

- Now, let's prove the information inequality for the case of a vector of parameters
- **Theorem (Information inequality):** Suppose $X_1, X_2, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} p(x|\theta^*)$, with $\mathcal{J}(\theta^*)$ positive definite. Let $\hat{\gamma}$ be an estimator with finite expected value $g(\theta)$. If $\nabla_{\theta} p(x|\theta^*)$ exists and can be passed under the integral sign with respect to $\int dP$ and $\int \hat{\gamma} dP$, then

$$\mathbb{V}\hat{\gamma} \succeq \nabla g(\theta^*)^{\top} \mathcal{J}_n(\theta^*)^{-1} \nabla g(\theta^*)$$

- Recall that $\mathbf{A} \succeq \mathbf{B}$ means $\mathbf{A} - \mathbf{B}$ is positive semidefinite

Special case: $g(\theta) = \theta$

- In the special case where we have iid data and an unbiased estimator $\hat{\theta}$ of θ , we have the simple result that:

$$\mathbb{V}\hat{\theta} \succeq \frac{1}{n} \mathcal{J}(\theta^*)^{-1},$$

the Cramér-Rao lower bound in d dimensions

- A related case: suppose we are estimating only a subset of θ , say, θ_1 , with remaining parameters so-called “nuisance parameters”
- What is the impact on the CRLB?

Nuisance parameters

- A common notation convention when dealing with partitions of the information matrix is to let $\mathcal{J}_{11}$ denote the $(1, 1)$ block of the information matrix, and $\mathcal{J}^{11}$ denote the $(1, 1)$ block of $\mathcal{J}^{-1}$ (and so on for other partitions, and for the observed information)
- Using this notation, the CRLB for estimating θ_1 is $\mathcal{J}^{11}/n$, as opposed to the CRLB for estimating θ_1 in the case where θ_2 is known: $\mathcal{J}_{11}^{-1}/n$
- Personally, I don't like this notation and prefer $\mathcal{V}$ to denote $\mathcal{J}^{-1}$ and $\mathcal{V}$ to denote $\mathcal{I}^{-1}$, mainly because $\mathcal{J}^{11}$ tends to cause some confusion as looking like an information, when it is very much *not* an information of any kind

Information loss due to nuisance parameters

- Recall that the relationship between these two quantities is given by the Schur complement, which we restate here in terms of our new information matrix notation (for the sake of compactness, I'm suppressing the dependence on θ here):

$$\mathcal{V}_{11}^{-1} = \mathcal{I}_{11} - \mathcal{I}_{12}\mathcal{I}_{22}^{-1}\mathcal{I}_{21},$$

or, if you prefer the superscript notation,

$$(\mathcal{I}^{11})^{-1} = \mathcal{I}_{11} - \mathcal{I}_{12}\mathcal{I}_{22}^{-1}\mathcal{I}_{21};$$

recall that $\mathcal{I}_{22}^{-1}$ is positive definite, so the term being subtracted cannot be negative ($\mathcal{I}_{11} \succeq \mathcal{V}_{11}^{-1}$)

- In other words, $\mathcal{I}_{12}\mathcal{I}_{22}^{-1}\mathcal{I}_{21}$ is the cost of not knowing θ_2 when estimating θ_1 (i.e., the information we've lost)

Orthogonality

- Only if $\mathcal{J}_{12} = \mathbf{0}$ do we suffer no information loss
- This can indeed happen; when it does, the parameters θ_1 and θ_2 are said to be *orthogonal*
- For example, consider the case where $X_i \stackrel{\text{iid}}{\sim} N(\mu, \sigma^2)$
- Here, $\bar{x}$ is unbiased for μ and achieves the CRLB regardless of whether we know σ^2 or not
- Such situations, however, are more the exception than the rule

Efficiency

- The information inequality and CRLB are of somewhat limited use in finite samples, since they are only achieved in special cases
- Reaching the CRLB *asymptotically*, on the other hand, is a different matter, and a much more attainable goal for a hardworking little estimator
- **Definition:** Let $X_i \stackrel{\text{iid}}{\sim} p(x|\theta^*)$. Suppose a sequence of estimates $\hat{\theta}_n$ for θ satisfies $\sqrt{n}(\hat{\theta}_n - \theta) \xrightarrow{d} N(0, \Sigma(\theta))$. The sequence is said to be *asymptotically efficient* if $\Sigma(\theta) = \mathcal{J}^{-1}(\theta)$ for all θ .
- While “asymptotically efficient” is a more accurate term, it is common to refer to such estimators as “efficient”

Efficiency and maximum likelihood

- As we have already shown, the MLE is asymptotically efficient (under certain regularity conditions)
- Thus, the MLE is in some sense optimal: at least asymptotically, no sequence of unbiased estimators can improve upon the MLE's accuracy
- For a long time in statistics, it was thought that no biased estimators could do better either; this belief, however, was upended by JL Hodges

Superefficiency

- Suppose $X_i \stackrel{\text{iid}}{\sim} N(\theta, 1)$ so that $\sqrt{n}(\hat{\theta} - \theta) \sim N(0, 1)$
- Consider the biased estimator

$$\tilde{\theta} = \begin{cases} 0 & \text{if } |\hat{\theta}| < n^{-1/4} \\ \hat{\theta} & \text{if } |\hat{\theta}| \geq n^{-1/4} \end{cases}$$

- Now, $\mathbb{P}\{|\hat{\theta}| < n^{-1/4}\} \rightarrow 1$ if $\theta = 0$ and $\rightarrow 0$ otherwise
- Thus, $\sqrt{n}(\tilde{\theta} - \theta) \xrightarrow{d} N(0, v)$, where $v = 1$ if $\theta \neq 0$ and $v = 0$ if $\theta = 0$

Superefficiency (cont'd)

- In other words, v improves upon the CRLB; a so-called “superefficient” estimator
- It's a pretty neat counterexample, although not necessarily a serious challenge to likelihood theory, as it can be shown (Le Cam, 1952) that the set of superefficient points always has Lebesgue measure zero
- This is sort of like saying that the MLE achieves the optimal variance almost everywhere, but this would only be a meaningful statement with a Bayesian prior, as otherwise there is no probability distribution associated with θ

Two Cauchy estimators

- To get a sense of why efficiency is a useful concept in terms of understanding the performance of estimators, let's return to our $X_i \stackrel{\text{iid}}{\sim} \text{Cauchy}(\theta)$ example from the previous lecture
- Consider two potential estimators, the sample median $\tilde{\theta}$ and the “one-step” estimator where we solve the likelihood equations using a Taylor series approximation about $\tilde{\theta}$
- Now, it can be shown that

$$\sqrt{n}(\tilde{\theta} - \theta) \xrightarrow{d} N(0, \pi^2/4)$$

$$\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{d} N(0, 2)$$

Asymptotic relative efficiency

- Since $\pi^2/4 = 2.47 > 2$, we can now appreciate the purpose of the one-step estimator: while both estimates are consistent, the one-step estimator is more efficient
- **Definition:** If $\sqrt{n}(\hat{\theta}_1 - \theta) \xrightarrow{d} N(0, \sigma_1^2)$ and $\sqrt{n}(\hat{\theta}_2 - \theta) \xrightarrow{d} N(0, \sigma_2^2)$, the *asymptotic relative efficiency* (ARE) of the two estimators is σ_1^2/σ_2^2
- For the Cauchy estimators, the ARE is $2.47/2 = 1.23$
- In other words, the median estimator requires approximately 23% larger sample size than the one-step estimator: we need $n = 123$ observations with the median estimator to obtain the same amount of information that the one-step estimator has with $n = 100$

Asymptotic relative efficiency: Tests

- This idea can be extended to testing as well
- Since the power of any reasonable test tends to 1 as $n \rightarrow \infty$, one typically considers $H_0 : \theta = \theta_0$ vs $H_a : \theta = \theta_0 + \Delta/\sqrt{n}$
- In this case, if $\beta_1 \rightarrow \Phi(\Delta a_1 - z_{(1-\alpha)})$ and $\beta_2 \rightarrow \Phi(\Delta a_2 - z_{(1-\alpha)})$, where β_i is the power of test i , the asymptotic relative efficiency of the two tests is $(a_1/a_2)^2$
- Generally, if two statistical procedures have the same limit as $n_1 \rightarrow \infty$ and $n_2 \rightarrow \infty$, then the ARE is the limit of the ratio n_1/n_2 ; the estimation and testing definitions we have given are special cases

Asymptotic relative efficiency: Tests (cont'd)

- For example, when $X_i \stackrel{\text{iid}}{\sim} N(\Delta/\sqrt{n}, \sigma^2)$, the one-sample t -test satisfies

$$\beta_1 \rightarrow \Phi(\Delta/\sigma - z_{(1-\alpha)})$$

while the Wilcoxon signed rank test satisfies

$$\beta_2 \rightarrow \Phi\left(\frac{\Delta}{\sigma} \sqrt{\frac{3}{\pi}} - z_{(1-\alpha)}\right)$$

- Thus, the ARE is $\pi/3 = 1.05$; when the data follows the normal distribution assumed by the t -test, the Wilcoxon test requires just 5% more data in order to achieve the same power

Additional remarks

- If the distribution is not normal, then there is no upper bound on the ARE of these two tests – one can always construct a distribution such that the Wilcoxon approach is that many times more efficient than a t -test
- This example illustrates a common use of efficiency: there is often a desire to develop robust nonparametric or semiparametric methods that make less restrictive assumptions than a parametric likelihood model, and efficiency provides something of a gold standard to compare against

Bayesian asymptotics

- We mentioned earlier that maximum likelihood estimation is “optimal” in the sense of being asymptotically efficient, but keep in mind that it is not a unique property – there may be multiple efficient approaches
- For example, Bayesian methods are also asymptotically efficient, as we are now going to see
- First, however, we’ll consider the related (and arguably, more important) question of the asymptotic normality of the *posterior distribution*
- Note that we’re now talking about convergence of a *random measure* (a random element whose outcome is a probability distribution) as opposed to a random variable, which introduces some complications

A rough statement

- Laplace (1820) was the first to observe that the posterior distribution is approximately Gaussian near its mode
- More specifically, the posterior density is approximately normal, centered at the MLE, with variance given by the inverse of the information:

$$\theta | \mathbf{x} \sim N(\hat{\theta}, \frac{1}{n} \mathcal{J}(\theta^*)^{-1})$$

$$\theta | \mathbf{x} \sim N(\hat{\theta}, \mathcal{I}_n(\hat{\theta})^{-1})$$

- However, making this observation rigorous becomes thorny — the target described above (a) depends on n , (b) is random, and (c) is just a point mass at θ^* as $n \rightarrow \infty$

The Bernstein-von Mises theorem

- Instead, we will need to focus on the posterior distribution of $\delta = \sqrt{n}(\theta - \theta^*)$ or equivalently, the posterior distribution of θ evaluated at $\theta^* + \delta/\sqrt{n}$
- The first formal proof of the asymptotic normality of the posterior distribution came from Bernstein (1917) and von Mises (1931), who showed that for any point δ , we have

$$\pi_n(\theta^* + \delta/\sqrt{n}) \xrightarrow{\text{as}} \phi(\delta),$$

where $\phi(\cdot)$ is the $N(\mathbf{0}, \mathcal{J}(\theta^*)^{-1})$ density

- This result is known as the Bernstein-von Mises theorem

Le Cam and the likelihood ratio

- This result, while historically groundbreaking, is slightly unsatisfying — what we really want is

$$\pi_n(\hat{\boldsymbol{\theta}} + \boldsymbol{\delta}/\sqrt{n}) \xrightarrow{\text{as}} \phi(\boldsymbol{\delta}),$$

- This statement is more complicated, because now we're talking about a distribution evaluated at a random point. . . it's not clear what the above statement even means
- To get around this, Le Cam (1953) employed the clever trick of considering the posterior ratio $\pi_n(\hat{\boldsymbol{\theta}} + \boldsymbol{\delta}/\sqrt{n}|\mathbf{x})/\pi_n(\hat{\boldsymbol{\theta}}|\mathbf{x})$, showing that it converged to the kernel of a $N(\mathbf{0}, \mathcal{J}(\boldsymbol{\theta}^*)^{-1})$ distribution

Bernstein-von Mises theorem (pointwise)

Theorem (Bernstein-von Mises): Suppose the prior $\pi(\boldsymbol{\theta})$ is continuous with $\pi(\boldsymbol{\theta}) > 0$ for all $\boldsymbol{\theta} \in \boldsymbol{\Theta}$. Under regularity conditions (A)-(D),

$$\frac{\pi_n(\hat{\boldsymbol{\theta}} + \boldsymbol{\delta}/\sqrt{n}|\mathbf{x})}{\pi_n(\hat{\boldsymbol{\theta}}|\mathbf{x})} \xrightarrow{\text{as}} \exp\left\{-\frac{1}{2}\boldsymbol{\delta}^\top \mathcal{J}(\boldsymbol{\theta}^*)\boldsymbol{\delta}\right\}.$$

Convergence in total variation

- The results we have discussed so far are pointwise — they don't guarantee that if you integrate these densities, you'll get the same probabilities (i.e., that credible sets behave correctly)
- Modern Bayesian asymptotics compares the entire posterior distribution to that of its target, which requires a global notion of distance between distributions
- A standard choice is total variation distance,

$$\int |p(x) - q(x)| dx,$$

which is more natural in Bayesian applications because it controls the largest possible difference in probabilities assigned to any event.

Bernstein-von Mises theorem (total variation)

Theorem (Bernstein-von Mises): Suppose the prior $\pi(\boldsymbol{\theta})$ is continuous with $\pi(\boldsymbol{\theta}) > 0$ for all $\boldsymbol{\theta} \in \boldsymbol{\Theta}$. If regularity conditions (A)-(D) hold, and if

$$\int \frac{\pi_n(\hat{\boldsymbol{\theta}} + \boldsymbol{\delta}/\sqrt{n}|\mathbf{x})}{\pi_n(\hat{\boldsymbol{\theta}}|\mathbf{x})} d\boldsymbol{\delta} \xrightarrow{\text{as}} \int \exp\{-\frac{1}{2}\boldsymbol{\delta}^\top \boldsymbol{\mathcal{J}}(\boldsymbol{\theta}^*)\boldsymbol{\delta}\} d\boldsymbol{\delta},$$

then

$$\int |\pi_n(\boldsymbol{\delta}|\mathbf{x}) - \phi(\boldsymbol{\delta})| d\boldsymbol{\delta} \xrightarrow{\text{as}} 0,$$

where $\phi(\cdot)$ is the $N(\mathbf{0}, \boldsymbol{\mathcal{J}}(\boldsymbol{\theta}^*)^{-1})$ density

Remarks

- The final line of the theorem follows from a result known as *Scheffé's useful convergence theorem*: if $\mathbf{x}_n \xrightarrow{\text{as}} \mathbf{x}$, $\mathbf{x}_n \succeq 0$, and $\mathbb{E}\mathbf{x}_n \rightarrow \mathbb{E}\mathbf{x} < \infty$, then $\mathbf{x}_n \xrightarrow{r} \mathbf{x}$ with $r = 1$
- These results are still usually called the Bernstein-von Mises theorem, even though they are substantially stronger than the original BvM results
- The end result is a nice theoretical justification for a variety of posterior approximation techniques (Laplace approximation, variational inference)

Asymptotic efficiency of the posterior mean

- As a corollary to the Bernstein-von Mises theorem, we can also conclude that Bayes estimators are asymptotically efficient
- To do so, we need the result that if $P_n \xrightarrow{\text{TV}} P$ and $\int \|\delta\| dP \leq \infty$, then

$$\int \delta dP_n \rightarrow \int \delta dP$$

- Applying this result to the Bernstein-von Mises theorem, we have

$$\sqrt{n}(\tilde{\theta} - \hat{\theta}) \xrightarrow{P} \mathbf{0},$$

where $\tilde{\theta}$ denotes the posterior mean

Remarks

- In other words, the posterior mean and the MLE are asymptotically equivalent
- The result is fairly intuitive: unless the prior has ruled θ^* out, eventually we will have enough data that the likelihood dominates the posterior and agrees with maximum likelihood
- Obviously, this does not imply that Bayesian and frequentist methods are equivalent (introducing a prior to improve performance at small sample sizes is a major advantage of Bayesian approaches), but it is reassuring to know that given enough data, both schools of thought will agree on an answer if they are working with the same likelihood model