

INTRODUCTION AND BACKGROUND RESULTS

$y \in \mathcal{Y}$ (Sample space)
 $\Theta \in \Omega$ (parameter space)
 $p(y|\Theta)$ (data density)
 $S_n = S(y_1, \dots, y_n)$ (statistic)

The probability distribution of S_n is called its "sampling distribution".

Our task is to use the data, via S_n , to draw inferences about Θ .

Estimators and tests are decision rules.

The characteristics of the sampling distribution of S_n will essentially determine its quality as a basis for inference.

e.g.: $y_i \sim (\mu, \sigma^2)$ (i.i.d.)

$$S_n = \bar{Y} = \frac{1}{n} \sum_{i=1}^n y_i \sim (\mu, \sigma^2/n)$$

"Moments" of a Distribution:

Let Y be a random variable with density $p(y|\Theta)$. Then the "moments" of Y are defined as:

$$(i) \quad E(Y^k) = \mu_k' = \int_{-\infty}^{\infty} y^k p(y|\Theta) dy; \quad k=1, 2, \dots$$

("Raw moments" or "moments about origin")

$$(ii) \quad E(Y - \mu_1')^k = \mu_k = \int_{-\infty}^{\infty} (y - \mu_1')^k p(y|\Theta) dy$$

($k=1, 2, \dots$)

("moments about the mean")

$$\text{So: } \mu_1' = E(Y) = \text{mean}$$

$$\mu_2 = E(Y - \mu_1')^2 = \text{variance}$$

$$= E(Y^2 + (\mu_1')^2 - 2Y\mu_1') = E(Y^2) - 2\mu_1' E(Y)$$

$$= E(Y^2) - (\mu_1')^2$$

$$= \mu_2' - (\mu_1')^2$$

We can always write the moments about mean in terms of moments about zero, & vice versa.

Why are moments important?

If all of the moments exist, then knowledge of the moments provides knowledge of the underlying distribution.

If all of moments exist, two distributions will be identical if their moments match.

Uniqueness requires existence of moments.

There is a convenient way of generating the moments of a distribution, without any need to integrate. We take the Fourier transform of the density, which is called the Characteristic Function of the random variable.

Definition: The characteristic Function of Y

$$\text{is } \phi_Y(t) = E(e^{itY}) ; \quad i^2 = -1 \\ = \int_{-\infty}^{\infty} e^{itY} p(Y) dY$$

Example: $Y \sim N(0, 1)$

$$\text{So, } \phi_Y(t) = E(e^{itY}) = \int_{-\infty}^{\infty} e^{itY} \frac{1}{\sqrt{2\pi}} e^{-y^2/2} dy \\ = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-(y+it)^2/2} \cdot e^{-t^2/2} dy \\ = e^{-t^2/2} \quad (= e^{(it)^2/2})$$

Show that:

(i) If $Y \sim N(\mu, \sigma^2)$, then

$$\phi_Y(t) = \exp[i\mu t - \frac{\sigma^2 t^2}{2}] .$$

(ii) If Y_1 & Y_2 are independent, then

the c.f. of their sum is the product of their individual c.f.'s.
[$\phi_{Y_1+Y_2}(t) = \phi_{Y_1}(t) \cdot \phi_{Y_2}(t)$]

Example: Let $Z \sim N(0, 1)$ & $Y = Z^2$.

$$\phi_Y(t) = \int_{-\infty}^{\infty} e^{itZ^2} \frac{1}{\sqrt{2\pi}} e^{-z^2/2} dz \\ = \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-z^2/2(1-2it)} dz$$

$$= (1-2it)^{\frac{1}{2}} \int_{-\infty}^{\infty} \frac{1}{(1-2it)^{-1/2} \sqrt{2\pi}} e^{-\frac{z^2}{2(1-2it)}} dz$$

$$= (1-2it)^{-1/2}$$

So, the c.f. for $\chi_{(1)}^2$ is $(1-2it)^{-1/2}$.
Similarly, the c.f. for $\chi_{(p)}^2$ is $(1-2it)^{-p/2}$.
There are various ways of getting the moments of the distribution from the c.f.:

(a) Differentiation -

$$\mu_k' = \left[\frac{d^k \phi(t)}{dt^k} \right]_{t=0} / i^k$$

e.g. : $Y \sim \chi_{(p)}^2$; $\phi_Y(t) = (1-2it)^{-p/2}$

$$\phi_Y'(t) = (-p/2)(-2i)(1-2it)^{-p/2-1}$$

$$= ip(1-2it)^{-(p/2+1)}$$

$$\mu_1' = [\phi_Y'(t)|_{t=0}] / i = p.$$

$$\phi_Y''(t) = -(ip)(p/2+1)(-2i)(1-2it)^{-(p/2+2)}$$

$$\mu_2' = [\phi_Y''(t)|_{t=0}] / i^2 = 2p(p/2+1)$$

$$= 2p(p+2)/2 = p(p+2)$$

$$\text{So, var}(Y) = \mu_2 - (\mu_1')^2 = 2p.$$

(b) Taylor-Series Expansion -

$\mu_k' =$ coefficient of $(it)^k/k!$ in series expansion.

e.g. $Y \sim N(0,1)$; $\phi_Y(t) = e^{-t^2/2}$

$$\text{Now, } e^x = 1 + x + \frac{x^2}{2!} + \dots = e^{(it)^2/2}$$

$$\text{So, } \phi_Y(t) = 1 + \frac{(it)^2}{2} + \left[\frac{(it)^2}{2} \right]^2 / 2! + \left[\frac{(it)^2}{2} \right]^3 / 3! + \dots$$

$$= 1 + \frac{(it)^2}{2!} + \frac{(it)^4}{4!} + \frac{(it)^6}{6!} + \dots$$

So, $\mu_k' = 0$; for all odd k .

$$\mu_2' = 1$$

$$\mu_4' = 3$$

$$\left(\frac{1}{8} = 3/4! \right)$$

$$\mu_2 = E(Y) = 0 ; \text{var}(Y) = 1 - 0^2 = 1.$$

Definition: Skewness = $\frac{\mu_3}{(\mu_2)^{3/2}}$

$$\mu_3 = E(Y-\mu)^3 = \mu_3' + 2(\mu_1')^3 - 3\mu_1'\mu_2'$$

e.g. $Y \sim N(0, 1)$

$\mu_1' = \mu_3' = 0$; $\mu_2' = 1$; $\mu_4' = 3$
So, Skewness = 0

Definition: Kurtosis = $[\mu_4 / \mu_2^2] - 3$

$$Y \mu_4 = E(Y - \mu)^4 = \mu_4' - 4\mu_3'\mu_1' + 6\mu_2'(\mu_1')^2 - 3(\mu_1')^4$$

e.g. $Y \sim N(0, 1)$

$$\mu_4 = 3 - 0 + 0 - 0 = 3$$

(Also true for $N(\mu, \sigma)$.)

Normal Dist'n. is "Mesokurtic"

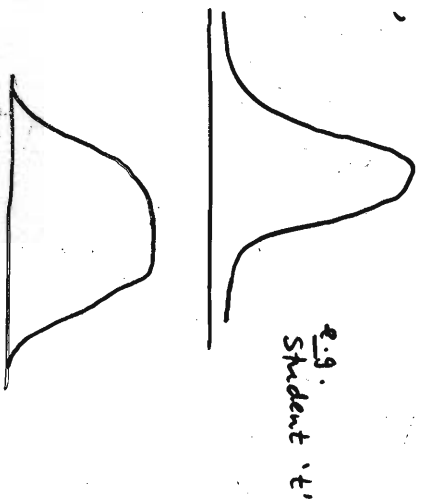
If kurtosis > 3 ,

"Leptokurtic"

If kurtosis < 3 ,

"Platykurtic"

e.g. Uniform



So, we often report the "excess kurtosis":
 $[\mu_4 / \mu_2^2] - 3$

Summary:

- * Moments of a distribution fully describe that dist'n. (if they exist).
- * Characteristic fctn. provides a painless way of obtaining the moments.
- * The moments are used to determine properties of a statistic when used as estimator or test statistic.

Estimator Properties:

Focus on "exact" (finite-sample) properties. Consider "asymptotic properties" after introducing Maximum Likelihood estimation.

(i) Bias:

$$\text{Bias}(\hat{\theta}_n) = E(\hat{\theta}_n) - \theta$$

("Unbiased" if Bias = 0)

Compares first moment of $\hat{\theta}_n$ with location

(iii) (Relative) Efficiency:

Let $\hat{\theta}_n$ & $\tilde{\theta}_n$ be 2 unbiased estimators of θ . Then $\hat{\theta}_n$ is efficient relative to $\tilde{\theta}_n$ if $[V(\tilde{\theta}_n) - V(\hat{\theta}_n)]$ is p.s.d. (scalar: $\text{var}(\tilde{\theta}_n) \geq \text{var}(\hat{\theta}_n)$)

If one or both of the estimators is biased, $\hat{\theta}_n$ is relatively more efficient than $\tilde{\theta}_n$ if $[M(\tilde{\theta}_n) - M(\hat{\theta}_n)]$ is p.s.d., where $M(\theta_n^*) = V(\theta_n^*) + \text{Bias}(\theta_n^*)\text{Bias}(\theta_n^*)'$ is the matrix MSE. (1st & 2nd moments)

(iii) Sufficiency:

If we can find a statistic, S_n , that tells us as much about the population parameter(s) as does the full sample, this statistic will be "sufficient" for estimation purposes. More formally -

Let $\{y_i\}_{i=1}^n$ be a random sample from $p(y|\theta)$. Then S_n is a "sufficient statistic" iff $p(y|S_n)$ does not depend on θ .

(If you know S_n , the sample values themselves add no information about θ .) This idea extends to "jointly sufficient statistics". And, any 1-1 transformation of a set of jointly sufficient statistics is also jointly sufficient.

Example: If $\sum y_i$ & $\sum y_i^2$ are jointly sufficient, then so are $\bar{y}$ and $\sum (y_i - \bar{y})^2$.

Now, how do we find a sufficient statistic in practice?

Factorization Theorem:

$\tilde{S}_n$ is sufficient iff

$$p(y_1, \dots, y_n | \theta) = g(\tilde{S}_n | \theta) h(y_1, \dots, y_n)$$

where $h(\cdot) \geq 0$ & does not involve θ ,
and $g(\cdot) \geq 0$ & depends on the y_i 's only
through $\tilde{S}_n$.

For jointly sufficient statistics, this
becomes:

$$p(y_1, \dots, y_n | \theta) = g(\tilde{S}_n, \dots, \tilde{S}_{rn} | \theta) h(y_1, \dots, y_n)$$

Example: Let $y_1, \dots, y_n$ be drawn randomly

from a Bernoulli density:

$$\begin{aligned} p(y | \theta) &= \prod_{i=1}^n p(y_i | \theta) \\ &= \prod_{i=1}^n \theta^{y_i} (1-\theta)^{1-y_i} ; 0 \leq \theta \leq 1 \\ &= \theta^{\sum y_i} (1-\theta)^{n - \sum y_i} \cdot 1 \end{aligned}$$

So, $g(\cdot) = \theta^{\sum y_i} (1-\theta)^{n - \sum y_i}$; $h(\cdot) = 1$

& $\sum y_i$ is sufficient for θ (so is $\bar{y}$)

Example:

$$f(y | \mu) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}} \exp\left[-\frac{1}{2}(y_i - \mu)^2\right]$$

$$\begin{aligned} &= \left(\frac{1}{2\pi}\right)^{n/2} \exp\left[-\frac{1}{2} \sum_i (y_i - \mu)^2\right] \\ &= \left(\frac{1}{2\pi}\right)^{n/2} \exp\left[-\frac{1}{2} \left(\sum_i y_i^2 + 2\mu \sum_i y_i + n\mu^2\right)\right] \\ &= \left(\frac{1}{2\pi}\right)^{n/2} \exp\left[\mu \sum_i y_i - \frac{n}{2}\mu^2\right] \exp\left[-\frac{1}{2} \sum_i y_i^2\right] \end{aligned}$$

$$\text{Let } g(\cdot) = \exp\left[\mu \sum_i y_i - \frac{n}{2}\mu^2\right]$$

$$h(\cdot) = \left(\frac{1}{2\pi}\right)^{n/2} \exp\left[-\frac{1}{2} \sum_i y_i^2\right]$$

Then, $\sum_i y_i$ is a sufficient statistic,

& so is $\bar{y}$.

Intuitively, it looks as if sufficient stats.
(or simple fctns. of them) may lead to
familiar & sensible estimators - this
will in deed turn out to be the case.

Now let's pursue the idea of estimation as a decision-making procedure & introduce some ideas from statistical decision theory.

(iv) Risk:

Suppose we have a loss function,

$$L(\hat{\theta}_n; \theta) \geq 0 \quad ; \quad \text{all } \theta, \hat{\theta}_n$$

$$= 0 \quad ; \quad \text{if } \hat{\theta}_n = \theta.$$

(N.B. : Loss is random)

To get a single value, average:

$$\text{Risk}(\hat{\theta}_n; \theta) = R(\hat{\theta}_n; \theta)$$

$$= E[L(\hat{\theta}_n; \theta)]$$

$$= \int_{\mathcal{Y}} L(\hat{\theta}_n; \theta) p(y|\theta) dy$$

e.g. Quadratic Loss -

(i) Scalar : $L = c(\hat{\theta}_n - \theta)^2$; $c > 0$

$R = MSE$ if $c=1$.

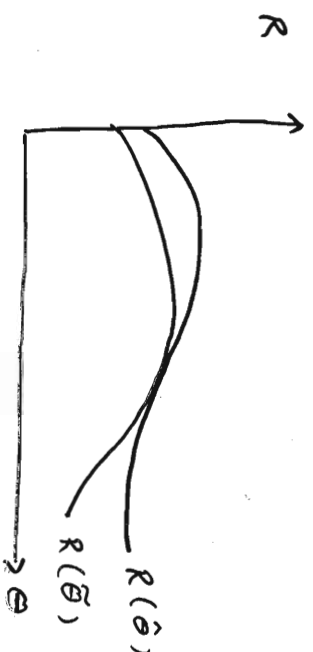
(ii) Vector : $L = (\hat{\theta}_n - \theta)' W (\hat{\theta}_n - \theta)$; W p.d.

$R = \text{tr}(MSE)$, if $W=I$.

(v) Admissibility -

A risk seems desirable. $\hat{\theta}_n$, an estimator (decision rule), is inadmissible if $\exists$ any $\tilde{\theta}_n$ such that $R(\tilde{\theta}_n) \leq R(\hat{\theta}_n)$ for all θ , and $R(\tilde{\theta}_n) < R(\hat{\theta}_n)$ for some θ .

An Admissible estimator is one which is not inadmissible.

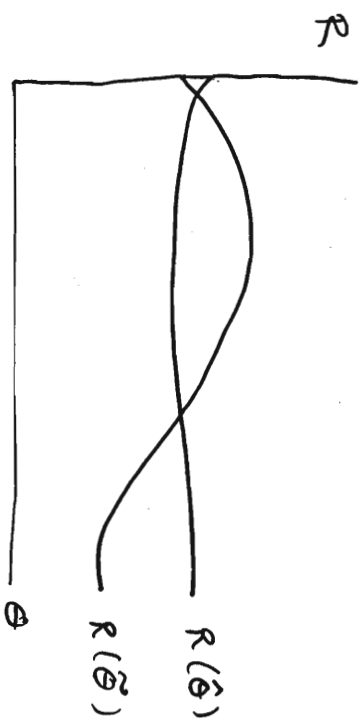


Many common estimators are inadmissible!

e.g. usual OLS estimators of β & σ^2 in linear regression, under quadratic loss, if $K \geq 3$! (Stein, 1956)

(vi) Mini-Max -

Typically, risk fctns. of estimators of interest cross -



Which estimator is preferred? We might average the risk (over θ), but what weights?

(Return to this as Bayesians!)

One (conservative) option -

$\hat{\theta}$ is Mini-max within family of estimators, for specific loss fctn. if

$$\max. R(\hat{\theta}) = \min. \{ \max. R(\theta^*) \},$$

$\forall \theta^*$ in family of interest.

15

Some Questions:

- * Are mini-max. estimators admissible?
- * Are admissible estimators mini-max.?
- * Are there any systematic, general, principles for constructing estimators that ensure that the estimators will have "good" properties of the types discussed so far?

We'll be considering 3 general principles -

- (a) Maximum Likelihood Estimation
- (b) Method of Moments Estimation
- (c) Bayesian Estimation

(& the associated testing principles).

As a second-best solution, if it is hard to find principles guaranteed to yield "good" estimators in the above senses, maybe we can find principles that do well for large n

16