

ERASMUS B4 Course

Winter Semester 2020/2021, 21 JANUARY

Subject: topics in statistical inference for given general population with the feature X .

Introduction

Let the general population P with the feature X be given.

We know that the main goal of M.S. is to establish the probability distribution of X .

To do this, we need the result of observation of P . It means, that we have to choose from P its representation and examine it.

In M.S. we assume that the result of the above procedure has a form

$$\mathbb{R}^n \ni (X_1, X_2, \dots, X_n) = (X_1, X_2, \dots, X_n)(U_0),$$

where $d(X_j) = d(X)$, $j = 1 \dots n$,

and $X_1, \dots, X_n$ are stoch. indep.

then we say that we have the SIMPLE SAMPLE.

The last lecture showed us that the presentation of S.S. is very important in statistical inference.

It is clear that S.S. describes $P_0 \subset \mathcal{P}$. We are going to explain below that the information for P_0 can be extrapolated to the whole $\mathcal{P}$.

The concept of Empirical Cumulative Prob. Dist. (ECPD)

Fix the S.S.

$$(X_1, X_2, \dots, X_n) = (X_1, X_2, \dots, X_n)(\omega).$$

The function

$$\mathbb{R} \ni t \longmapsto F_n(t, \omega) \stackrel{\text{def}}{=} \frac{1}{n} \# \{i: X_i \leq t\}$$

is called ECPD corresponding to $\mathcal{P}$ with X for given S.S.

We know, that ECPD we can obtain from histogram (or frequency diagram) and vice versa. Moreover, the graph of ECPD is a "step function".

Why ECPD is so important in M.S.?

It is explained by the following theorem

Theorem (on ECPD) (GNIEDENKO)

For given G.P. $\{X_i\}$ with the feature X let us consider the ~~sequence of~~ functions

$$\mathbb{R} \times \Omega \ni (t, \omega) \longrightarrow F_n(t, \omega) \stackrel{\text{def}}{=} \frac{1}{n} \# \{i: X_i(\omega) < t\},$$

when for $n \geq 2$

$X_1, X_2, \dots, X_n$ are st. independent with the same distribution X .

If F_X is (unknown) a cumulative prob. dist. f. corresponding to X , then for the events

$$A_t = \left\{ \omega \in \Omega: F_n(t, \omega) \longrightarrow F_X(t) \right\}, \quad t \in \mathbb{R}$$

we have

$$P(A_t) = 1, \quad t \in \mathbb{R}.$$

Remarks

1st. Each member $F_n(t, u)$ of the above sequence is a generalisation of ECDF, namely

$$F_n(t, w_0) = F_n(t, u) \Big|_{u=w_0}.$$

So if we take $(F_n(t, u))_{n \geq 2}$ it means we consider (from theoretical point of view!) all S.S. (with respect of n and $w \in \mathbb{R}$).

2nd. Since, for given t , $P(A_t) = 1$, we can assume that the definition of S.S. belongs to A_t .

There we have an approximation rule

$$\boxed{F_n(t, w_0) \approx F_X(t)}$$

which is the basic argument in extrapolation procedure.

The concept of estimation

The stage of the statistical inference using the concept of ECDF generally does not solve the main problem of M.S.

On the output of this stage the typical solution is as follows :

$F_X(t, \theta)$ - we know the "shape" of F_X , but the value of the parameter θ is unknown.

So, we need the next stage of S.I, to ~~estimate~~ establish the value θ .

It is called the ESTIMATION PROCEDURE.

The main idea of E.P. says :

For given $\alpha \in (0, 1)$ (significant level),
for instance $\alpha = 0.05$ or 0.1

We try to construct of two r.v. Z_1, Z_2 such that:

(i) Z_1, Z_2 is constructed by using the random vector (X_1, X_2) which defines the S.S.

(ii) $Z_1(u) < Z_2(u)$ with probability 1

→ (iii) $P(\text{Real: } Z_1(u) < \theta < Z_2(u)) = 1 - \alpha$.

Then we say we have a random interval (Z_1, Z_2) .

Now, if we use $\omega \in \Omega$ from the S.S.

We obtain the real interval $(Z_1(\omega), Z_2(\omega)) \subset \mathbb{R}$,

and finally ^{(by (iii))} we have the following conclusion

$\theta \in (Z_1(\omega), Z_2(\omega))$ with prob. $1 - \alpha$

Example

Suppose that we know (for instance from analysis of a histogram) that the S.I.

$(X_{n-1}, X_n) = (X_{n-1}, X_n)(U_0)$, let $X \in N(m, \sigma^2)$, and the value σ^2 is fixed.

If we fix $\alpha \in (0, 1)$ and we take

$$Z_1 = \bar{X}_n - n_\alpha \frac{\sigma}{\sqrt{n}}, \quad Z_2 = \bar{X}_n + n_\alpha \frac{\sigma}{\sqrt{n}},$$

where $\bar{X}_n = \frac{1}{n}(X_{n-1} + X_n)$, we can find n_α if we resolve the equation

$$(\#) \quad P(\text{I.V.} : Z_1(U) < m < Z_2(U)) = 1 - \alpha.$$

Namely,

$$(\#) \Leftrightarrow P(\text{I.V.} : \bar{X}_n(U) - n_\alpha \frac{\sigma}{\sqrt{n}} < m < \bar{X}_n(U) + n_\alpha \frac{\sigma}{\sqrt{n}}) = 1 - \alpha$$

$(=)$

$$\text{Proven: } -n_\alpha < \frac{\bar{X}_n - m}{\frac{\sigma}{\sqrt{n}}} < n_\alpha \Big) = 1 - \alpha.$$

But, according to our assumption of X ,

$$\frac{\bar{X}_n - m}{\frac{\sigma}{\sqrt{n}}} \in N(0,1), \text{ so}$$

$$(\#) \Rightarrow \Phi(n_\alpha) - \Phi(-n_\alpha) = 1 - \alpha \Rightarrow$$

$$2\Phi(n_\alpha) - 1 = 1 - \alpha \Rightarrow$$

$$\boxed{\Phi(n_\alpha) = 1 - \frac{\alpha}{2}}$$

$$\text{For instance, for } \alpha = 0.05, \quad \underline{n_\alpha = 1.96}.$$

The above ~~pro~~ shows, that by using z_1, z_2 given as above, we can construct the random interval and ~~finally~~ finally to achieve E.P.

Let's realize the following numerical simulation.

$$\rightarrow (x_1, x_2, x_3, x_4) = (2.13, 2.12, 1.97, 2.07)$$

for $X \in N(m, (0.25)^2)$ and $\alpha = 0.05$.

Since for $\alpha = 0.05$, $n\alpha = 1.96$, we have

$$Z_1(u) = \bar{X}_n(u) - 1.96 \frac{0.25}{2}$$

$$Z_2(u) = \bar{X}_n(u) + 1.96 \frac{0.25}{2}, \quad u \in U.$$

So, for $u = u_0$

$$\bar{X}_n(u_0) = 2.000 \text{ and}$$

$$Z_1(u_0) = 2.000 - 1.96 \frac{0.25}{2} = 1.938$$

$$Z_2(u_0) = 2.000 + 1.96 \frac{0.25}{2} = 2.028.$$

Therefore, in this case

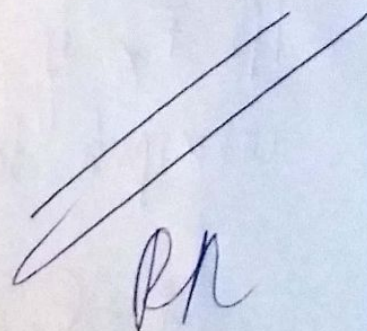
$$\text{we get } (1.938, 2.028) \text{ with prob. } 0.95$$

Conclusion

The ideas presented above belong to many others that make up the stochastic inference procedure (S.I.P.).

The main goal of S.I.P. is to establish the probability distribution of the feature X for given general population P .

But to present these methods, you need another course that I invite you to take in the future.


RN