

Tema 1

30/01/2025

0 = Introducción y motivación ✓

1 = Estadística descriptiva → Erick

2 = Modelos estadísticos

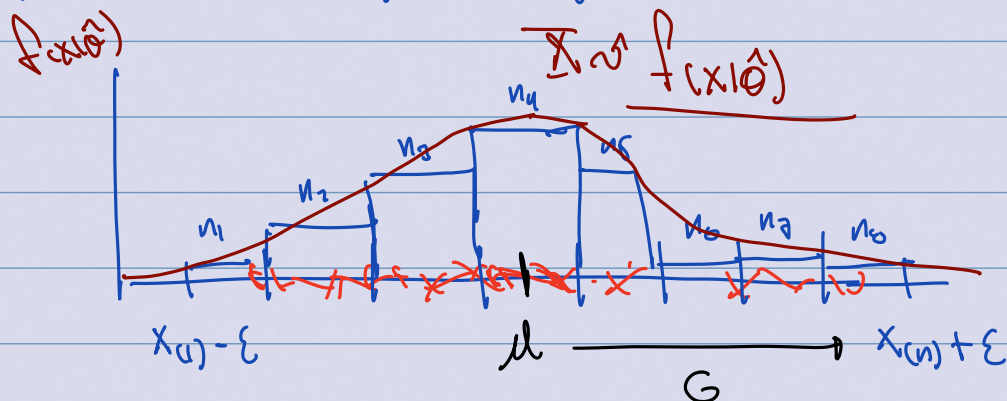
Lo que pasará viene "el proceso de la inferencia estadística"

① F.A. que nos interesa conocer / describir / predecir
y poder recolectar información generada por F.A., i.e.
datos $X_1, X_2, X_3, \dots, X_n$

② Asumir que el F.A. puede describirse mediante un v.o. X
¿Qué significa esto?

el comportamiento agregado de $X_1, X_2, \dots, X_n$
puede describirse mediante un v.o. X

pasar en un histograma o gráfico de barras



pero aquí necesitamos decidir si $X \sim f(x|\theta)$ ó $X \sim f(x)$

③ Para poder echar andar la maquinaria de la inferencia estadística, necesitamos asumir un modelo de recolección de datos, i.e.

los datos $x_1, x_2, \dots, x_n$ fueran el resultado de algún mecanismo aleatorio $\rightarrow$ probabilidades

④ Con estos supuestos podremos estimar

$$\mu = \mathbb{E}(X)$$

$$\sigma = \sqrt{\text{Var}(X)}$$

$$P(X > c)$$

$$P(a < X < b) \quad \leftarrow \text{resolver fijando } \underline{a \times b}$$

$$\text{o } P(a < X < b) = 0.95$$

fijar

y encontrar a y b

$$\text{Quizá nos interesen } \underline{f(x|\theta)} \rightarrow f(x|\hat{\theta}_n)$$

$$f(x) \rightarrow \hat{f}_n(x)$$

Vamos a describir con mayor detalle los pasos ② y ③

En el paso ②, necesitamos elegir/asumir un modelo estadístico que describa razonablemente bien a FA, de manera general
know 2 cosas

- Modelo estadístico paramétrico
- Modelo estadístico no paramétrico

Por ejemplo

$$M = \left\{ f(x|\mu, \sigma^2) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2\sigma^2}(x-\mu)^2}, \mu \in \mathbb{R} \text{ y } \sigma^2 > 0 \right\}$$

en este caso M es un modelo paramétrico pues lo podemos describir mediante un número finito de parámetros $\rightarrow (\mu, \sigma^2)$

En general a los modelos paramétricos los describiremos como

$$M = \{ f(x|\theta) \mid \theta \in \Omega \}$$

en donde $\theta = (\theta_1, \theta_2, \dots, \theta_K)$ es el parámetro que gobierna el comportamiento de $f(x|\theta)$ y $K < \infty$

Un modelo estadístico NO paramétrico, es un conjunto M que no se puede describir mediante un número finito de parámetros

$$M = \{ \text{todas las funciones de distribución acumulada } F(x) \}$$

$$M = \{ \text{todas las funciones " " } F(x) \text{ continuas} \}$$

$$M = \{ \dots, \text{simétricas} \}$$

En los viejos días se decía que un modelo paramétrico tenía parámetro y un no paramétrico, no tenía parámetro.

Entonces elegir las estadísticas bayesianas no paramétricas
y de alguna manera definir

$$f(x) = \sum_{j=1}^{\infty} \omega_j N(x | \mu_j, \sigma_j^2)$$

utilizando el proceso Dirichlet

Ferguson T. S (1973) A Bayesian analysis of
some non-parametric problems

Ahora vamos a describir con mayor detalle el punto ③, en este
caso necesitamos definir un modelo de recolección de información
que considere que los datos se generan $x_1, x_2, \dots, x_n$

- 1) Están sujetos a variaciones independientes
- 2) Si obtenemos más información de FA, $x_1^*, \dots, x_n^*$
no será la misma información que $x_1, \dots, x_n$
- 3) La v.o. X describe el comportamiento
agregado de los datos generados por FA

Un modelo de recolección que cumple con lo anterior y que
es sencillo de trabajar es el de

muestra aleatoria = m.a. $\equiv$ v.a.i.i.d

Entonces asumir que

$$1) \underbrace{X_1, X_2, \dots, X_n}_{\text{v.o.}} \text{ es un m.c. de } \underline{X} \rightsquigarrow \begin{cases} f(x|\theta) \\ f(x) \end{cases}$$

$$2) \underbrace{x_1, x_2, \dots, x_n}_{\substack{\text{valores observados} \\ \text{al realizar la medición}}} \text{ es una realización de } X_1, X_2, \dots, X_n$$

Del curso de probabilidad

$$X_1, X_2, \dots, X_n \text{ es un m.c. de } X \rightsquigarrow f(x|\theta) \text{ si}$$

$$1) X_i \rightsquigarrow f(x|\theta), i=1, \dots, n \leftarrow \text{identicamente distribuidos}$$

$$2) X_i \perp X_j, \forall i \neq j$$

Bajo nuestro modelo de recolección de información, la densidad conjunta de los datos $x_1, x_2, \dots, x_n$ es dada por

$$\begin{aligned} f_{X_1, X_2, \dots, X_n}(x_1, x_2, \dots, x_n | \theta) &= \prod_{i=1}^n f_{X_i}(x_i | \theta) \quad \leftarrow \text{independencia} \\ &= \prod_{i=1}^n f_X(x_i | \theta) \quad \leftarrow \text{identicamente distribuidos} \\ &= \prod_{i=1}^n f(x_i | \theta) \quad \leftarrow \text{por Abuses} \end{aligned}$$

Ahora vamos a describir con mayor detalle el punto (4)

De manera muy general, podemos decir que la inferencia estadística nos da herramientas para resolver 3 problemas

- i) Estimación puntual
- ii) Estimación por intervalo
- iii) Pruebas de hipótesis

i) Estimación puntual. El objetivo es obtener una única mejor estimación del parámetro o cantidad de interés.

Ej. $X \sim P(x|\theta) \rightarrow$ la mejor estimación $\hat{\theta}_n = h(x_1, \dots, x_n)$ de θ .

$X \sim f(x) \rightarrow$ la mejor aproximación a $f(x)$ por x en $\mathbb{R}$,

$$\Rightarrow \underline{\underline{\hat{f}_n(x)}}$$

$X \sim f(x) \rightarrow$ la mejor aproximación a $f(x)$, por x en $\mathbb{R}$
 $\hat{f}_n(x)$

Vamos a concentrarnos en el caso en el que $X_1, X_2, \dots, X_n$
es una m.c. de $X \sim f(x|\theta)$ y queremos encontrar el "mejor"
estimador de θ

De alguna muestra vamos a obtener un estimador

$\hat{\theta}_n = h(X_1, X_2, \dots, X_n) \leftarrow$ es una función de la m.c.
que usamos para estimar θ , es una V.O.

Como el estimador es una V.O. podemos analizar algunas
propiedades usando teoría de v.e.s. Pero antes definamos
la estimación

$\hat{\theta}_n = h(X_1, X_2, \dots, X_n)$ es el valor que usamos
en la práctica para
aproximar θ y
se obtiene incertidumbre
los valores observados
 $x_1, x_2, \dots, x_n$ en n

$\Rightarrow$ Estimador $\equiv$ V.O. $\rightarrow$ probabilidad

Estimación $\equiv$ valor que usamos para aproximar θ .