

استنباط آماري

پيشرفته (۱)

نادر نعمت الهی

گروه آمار

دانشگاه علامه طباطبائی

فهرست مطالب

عنوان	صفحه
فصل اول: کلیات و مفاهیم اولیه	۱
۱- خلاصه‌ای از نظریه اندازه	۲
۲- خانواده گروهی از توزیع‌ها	۸
۳- خانواده نمایی از توزیع‌ها	۱۲
۴- آماره‌های بسنده	۱۹
۵- کامل بودن	۲۹
۶- توابع زیان محدب	۳۳
فصل دوم: نااریبی	۳۸
۱- برآوردگرهای نااریب	۳۹
۵- نامساوی اطلاع (نامساوی کرامر - راثو)	۵۱
۶- نامساوی اطلاع (کرامر - راثو) در حالت چند پارامتری	۵۸
فصل سوم: پایایی	۶۳
۱- مفهوم پایایی و کلیات	۶۴
۲- برآوردگر MRE پارامتر مکان	۷۸
۳- برآوردگر MRE پارامتر مقیاس	۸۶
۴- برآوردگر MRE پارامترهای مکان-مقیاس	۹۳
فصل چهارم: تصمیم بیزی	۱۰۲
۱- مقدمه	۱۰۳
۲- تصمیم بیزی	۱۰۵
فصل پنجم: مینیماکس بودن و مجاز بودن	۱۲۲
۱- تصمیم مینیماکس	۱۲۳
۲- تصمیم‌های مجاز (پذیرفتنی)	۱۳۸
مراجع	۱۶۴

فصل ۱

کلیات و مفاهیم اولیه

بخش ۱: خلاصه‌ای از نظریه اندازه^۱

در ابتدا خلاصه‌ای از نظریه اندازه که مورد لزوم مباحث این درس می‌باشد را مرور می‌کنیم.

تعریف ۱-۱: سیگما میدان (σ -field)

فرض کنید Ω یک مجموعه باشد و A مجموعه‌ای غیر تهی از زیر مجموعه‌های Ω باشد. A را یک سیگما میدان گویند به شرط آن که شرایط زیر برای آن برقرار باشد

الف) اگر $A \in A$ آنگاه $A^C \in A$

ب) اگر $A_1, A_2, \dots$ مجموعه‌هایی در A باشند آنگاه $\bigcup_{n=1}^{\infty} A_n \in A$ ■

از تعریف فوق نتیجه می‌شود که اگر A یک سیگما میدان باشد آنگاه $\Omega \in A$ ، $\emptyset \in A$ و

$$\bigcap_{n=1}^{\infty} A_n \in A \text{ خواهد بود.}$$

تعریف ۲-۱: یک اندازه μ یک تابع مجموعه‌ای روی سیگما میدان A از زیر مجموعه‌های Ω است

هر گاه در شرایط زیر صدق کند $\mu: A \rightarrow [0, \infty)$

الف) اگر $A \in A$ آنگاه $\mu(A) \geq 0$

ب) اگر $A_1, A_2, \dots \in A$ و A_i ها مجزا باشند، یعنی $\forall i \neq j, A_i \cap A_j = \emptyset$ ، آنگاه

$$\mu\left(\bigcup_{n=1}^{\infty} A_i\right) = \sum_{n=1}^{\infty} \mu(A_n) \quad \blacksquare$$

اگر A یک سیگما میدان از زیر مجموعه‌های Ω باشد آنگاه (Ω, A) را یک فضای اندازه‌پذیر^۲

گویند و مجموعه‌ی $A \in A$ را اندازه‌پذیر گویند. همچنین اگر μ یک اندازه باشد که روی (Ω, A)

تعریف شده باشد آنگاه سه تایی (Ω, A, μ) را یک فضای اندازه^۳ می‌نامند. یک اندازه μ سیگما

^۱ Measure Theory

^۲ Measurable Space

^۳ Measure Space

متناهی (σ -finite) نامیده می‌شود. هر گاه مجموعه‌های $A_i \in \mathcal{A}$ موجود باشند که $\Omega = \bigcup_i A_i$ و $\mu(A_i) < \infty$. تمام اندازه‌هایی که در این درس بررسی می‌شوند، سیگما متناهی هستند و به همین دلیل آن‌ها را تنها اندازه می‌نامیم.

مثال ۱-۱: اندازه‌شمارنده^۱. فرض کنید $\Omega = \{w_1, w_2, w_3, \dots\}$ و تمام زیرمجموعه‌های Ω $\mathcal{A} = \Omega$

اگر A متناهی باشد تعداد اعضای مجموعه $\Rightarrow \mu(A) = \begin{cases} A \\ \infty \end{cases}$ if $A \in \mathcal{A}$ اگر A نامتناهی باشد
این اندازه را اندازه شمارنده گویند. ■

مثال ۲-۱: اندازه لبگ^۲. فرض کنید Ω فضا اقلیدسی n بعدی باشد،

$\Omega = \{(w_1, \dots, w_n) | w_i \in \mathbb{R}, \forall i\}$ و $\mathcal{A}$ کوچکترین سیگما میدانی باشد که شامل تمام مستطیل‌های باز زیر باشد:

$$A = \{(w_1, w_2, \dots, w_n) | a_i < w_i < b_i, -\infty < a_i < b_i < +\infty\}$$

$\mathcal{A}$ شامل تمام زیرمجموعه‌های باز و بسته Ω می‌باشد و به اعضای آن مجموعه‌های برل^۳ گویند. تنها یک اندازه μ می‌توان روی این $\mathcal{A}$ معرفی کرد که به مجموعه A فوق اندازه

$$\mu(A) = (b_1 - a_1)(b_2 - a_2) \dots (b_n - a_n)$$

را نسبت می‌دهد (که حجم مجموعه A است) و به آن اندازه لبگ گویند. ■

توابع اندازه‌پذیر^۴:

یک تابع حقیقی $f: \Omega \rightarrow \mathbb{R}$ یک تابع اندازه‌پذیر است اگر برای هر مجموعه برل B داشته باشیم که

$$f^{-1}(B) \in \mathcal{A} \text{ برای مثال}$$

$$A \in \mathcal{A}, I_A(x) = \begin{cases} 1 & x \in A \\ 0 & x \notin A \end{cases}$$

^۱ Counting Measure

^۲ Lebesgue Measure

^۳ Borel Sets

^۴ Measurable Functions

I_A یک تابع اندازه‌پذیر است. اگر $A_1, A_2, \dots, A_K \in \mathcal{A}$ و $A_i \cap A_j = \emptyset$ آنگاه

$$f_s(x) = a_1 I_{A_1}(x) + a_2 I_{A_2}(x) + \dots + a_k I_{A_k}(x)$$

یک تابع اندازه‌پذیر است که به آن تابع ساده گویند. حال اگر یک دنباله از توابع ساده غیرنزولی

$$f^+(x) = \lim_{n \rightarrow \infty} f_{s_n}(x) \in [0, \infty)$$

آنگاه f^+ یک تابع اندازه‌پذیر است. همچنین یک تابع دلخواه f را می‌توان بصورت $f = f^+ - f^-$ نوشت که در آن

$$f^+(x) = \max(f(x), 0)$$

$$f^-(x) = -\min(f(x), 0)$$

یک تابع $f: \Omega \rightarrow \mathbb{R}$ اندازه‌پذیر خواهد بود اگر قسمت‌های مثبت و منفی آن یعنی f^+ و f^- اندازه‌پذیر باشند و همچنین f^+ و f^- در صورتی اندازه‌پذیر هستند که بصورت حد دنباله‌ای از توابع ساده غیرنزولی باشند.

انتگرال یک تابع نسبت به اندازه μ :

$$I_A(x) = \begin{cases} 1 & x \in A \\ 0 & x \notin A \end{cases} \Rightarrow \int I_A d\mu = \mu(A)$$

$$f_s(x) = a_1 I_{A_1}(x) + \dots + a_k I_{A_k}(x) \Rightarrow \int f_s d\mu = \sum_{i=1}^k a_i \mu(A_i)$$

$$f(x) = \lim_{n \rightarrow \infty} f_{s_n}(x) \Rightarrow \int f d\mu = \lim_{n \rightarrow \infty} \int f_{s_n} d\mu$$

$$f = f^+ - f^- \Rightarrow \int f d\mu = \int f^+ d\mu - \int f^- d\mu$$

مثال ۱-۳: اگر $\Omega = \{x_1, x_2, \dots\}$ و μ اندازه شمارنده باشد آنگاه برای یک تابع $f: \Omega \rightarrow [0, \infty)$ داریم

که

$$f_{s_n}(x) = \begin{cases} f(x_i) & x = x_i \quad i = 1, 2, \dots, n \\ 0 & \text{o.w.} \end{cases}$$

$$= f(x_1)I_{\{x_1\}}(x) + f(x_2)I_{\{x_2\}}(x) + \dots + f(x_n)I_{\{x_n\}}(x)$$

$$f(x) = \lim_{n \rightarrow \infty} f_{s_n}(x) = f(x_1)I_{\{x_1\}}(x) + f(x_2)I_{\{x_2\}}(x) + \dots$$

$$= \begin{cases} f(x_i) & x = x_i \\ 0 & \text{o.w.} \end{cases}$$

$$\int f(x) d\mu = \lim_{n \rightarrow \infty} \int f_{s_n}(x) d\mu = \lim_{n \rightarrow \infty} \sum_{i=1}^n f(x_i) = \sum_{i=1}^{+\infty} f(x_i)$$

پس در صورتی f نسبت به μ انتگرال پذیر است که $\sum_{i=1}^{+\infty} f(x_i)$ متناهی باشد. ■

واژه تقریبا همه جا^۱

فرض کنید یک گزاره P برای تمام نقاط Ω درست باشد بجز نقاط A که در آن $\mu(A) = 0$ (یعنی P برای تمام $\Omega - A$ درست است و $\mu(A) = 0$ می باشد) در این صورت گوئیم این گزاره تقریبا همه جا μ درست است. ($a.e.\mu$)

قضیه همگرایی مغلوب^۲:

اگر f_n دنباله‌ای از توابع اندازه پذیر باشند و $\lim_{n \rightarrow \infty} f_n(x) = f(x)$ ($a.e.\mu$) و یک تابع انتگرال پذیر g چنان موجود باشد که $|f_n(x)| \leq g(x) \quad \forall x$ آنگاه f و f_n انتگرال پذیر هستند و

$$\int f d\mu = \lim_{n \rightarrow \infty} \int f_n d\mu$$

■

لم فاتو^۳

اگر $\{f_n\}$ دنباله‌ای از توابع نامنفی اندازه پذیر باشند، آنگاه

$$\int (\liminf_{n \rightarrow \infty} f_n) d\mu \leq \liminf_{n \rightarrow \infty} \int f_n d\mu$$

■

^۱ almost every where

^۲ Dominated Convergence Theorem

^۳ Fatou's Lemma

قضیه رادون نیکودیم^۱

فرض کنید f یک تابع نامنفی انتگرال پذیر و $A \in \mathcal{A}$ باشد. قرار می دهیم:

$$\nu(A) = \int_A f d\mu \quad (۱)$$

در این صورت $\nu(A)$ یک اندازه روی $(\Omega, \mathcal{A})$ است و واضح است که

$$\mu(A) = 0 \Rightarrow \nu(A) = 0 \quad (۲)$$

اگر رابطه (۲) برقرار باشد گوییم که ν نسبت به μ مطلقاً پیوسته^۲ است. ■

قضیه رادون نیکودیم می گوید که رابطه (۲) شرط لازم و کافی برای آن است که یک تابع f موجود

باشد که در رابطه (۱) صدق کند. تابع f را مشتق رادون نیکودیم ν نسبت به μ گویند $f = \frac{\partial \nu}{\partial \mu}$.

این مشتق به صورت ae یکتا است یعنی اگر g نیز در رابطه (۱) صدق کند آنگاه:

$$f = g \quad ae.\mu$$

قضیه فوبینی^۳

فرض کنید $(\Omega_1, \mathcal{A}, \mu)$ و $(\Omega_2, \mathcal{B}, \nu)$ دو فضای اندازه و f یک تابع غیرمنفی اندازه پذیر روی

$\Omega_1 \times \Omega_2$ باشد. در این صورت:

$$\begin{aligned} \int_{\Omega_1} \left[\int_{\Omega_2} f(x, y) d\nu(y) \right] d\mu(x) &= \int_{\Omega_2} \left[\int_{\Omega_1} f(x, y) d\mu(x) \right] d\nu(y) \\ &= \int_{\Omega_1 \times \Omega_2} f d(\mu \times \nu) \quad \blacksquare \end{aligned}$$

اندازه احتمال^۴

فرض کنید Ω فضای نمونه یک آزمایش تصادفی باشد. یک اندازه P روی فضای اندازه پذیر

$(\Omega, \mathcal{A})$ که در شرط $P(\Omega) = 1$ صدق کند را اندازه ی احتمال گویند و به مقدار $P(A)$ احتمال

A گویند و به $(\Omega, \mathcal{A}, P)$ فضای احتمال گویند.

^۱ Radon Nikodym

^۲ Absolutely Continuous

^۳ Fubini's Theorem

^۴ Probability Measure

اگر P نسبت به اندازه μ مطلقاً پیوسته باشد و p مشتق رادون نیکودیم آن باشد آنگاه

$$P(A) = \int_A p d\mu$$

p را تابع چگالی احتمال P نسبت به μ گویند.

در این درس معمولاً Ω را فضای اقلیدسی در نظر می گیریم. اگر اندازه‌ی احتمال P نسبت به اندازه‌ی لبگ مطلقاً پیوسته با چگالی p باشد آنگاه توزیع مربوطه را یک توزیع پیوسته و اگر P نسبت به اندازه شمارنده مطلقاً پیوسته با چگالی p باشد آنگاه توزیع مربوطه را یک توزیع گسسته گویند. در کاربردهای معمولی این درس داریم که :

$$\int f dP = \int f p d\mu = \begin{cases} \int f(x) p(x) dx & (p \text{ پیوسته}) \text{ اندازه لبگ} \\ \sum f(x) p(x) & (p \text{ گسسته}) \text{ اندازه شمارنده} \end{cases}$$

در استنباط آماری معمولاً با یک خانواده از توزیع‌ها به صورت $P = \{P_\theta : \theta \in \Theta\}$ سر و کار داریم. اگر به ازای هر θ ، P_θ نسبت به اندازه μ مطلقاً پیوسته باشد گوییم که خانواده P تحت تسلط μ قرار دارد.^۱ اگر $P_\theta(N) = 0$ گوئیم $N = P_\theta - \text{null set}$ و اگر $\forall \theta \in \Theta$ ، $P_\theta(N) = 0$ گوئیم $N = P - \text{null set}$.

متغیر تصادفی:

در یک فضای احتمال (Ω, A, P) تابع $X : \Omega \rightarrow R$ را اگر اندازه‌پذیر باشد یک متغیر تصادفی گویند یعنی اگر برای هر مجموعه بُرل B ، $X^{-1}(B) \in A$ باشد آنگاه X را یک متغیر تصادفی گویند

$$A = X^{-1}(B) \quad , \quad P_X(B) = P(X \in B) = P(A)$$

P_X را توزیع متغیر تصادفی X نامیم.

اگر P_X نسبت به اندازه لبگ مطلقاً پیوسته باشد آنگاه X را یک متغیر پیوسته با تابع چگالی p_X

$$P_X(A) = \int_A p_x d\mu \quad \text{گوئیم هر گاه}$$

$$P_X(A) = \int_A p_x d\mu \rightsquigarrow \text{Counting Measure} \quad \text{و } X \text{ را گسسته گوئیم هر گاه}$$

^۱ Dominated by μ

در هر صورت p_X را می‌توان چگالی P_X نامید. گرچه در حالت گسسته آن را تابع احتمال نیز گویند.

اگر $A = (-\infty, x]$ ، $F_X(x) = P_X(A)$ ، آنگاه F_X را تابع توزیع متغیر تصادفی X نامیم. اگر X پیوسته باشد (P_X مطلقاً پیوسته است) آنگاه ثابت می‌شود که تابع $F_X(x)$ نیز (با تعریف توابع مطلقاً پیوسته در آنالیز) مطلقاً پیوسته است.

در آنالیز، تابع g را مطلقاً پیوسته نامیم اگر به ازای هر $\varepsilon > 0$ یک $\delta > 0$ وجود داشته باشد که به

$$\sum_{i=1}^n (b_i - a_i) < \delta \Rightarrow \sum_{i=1}^n |f(b_i) - f(a_i)| < \varepsilon \quad \text{ازای هر } n,$$

بخش ۲: خانواده گروهی از توزیع‌ها^۱

در این درس با دو خانواده از توزیع‌ها سر و کار داریم که یکی خانواده گروهی و دیگری خانواده نمایی از توزیع‌ها می‌باشد.

تعریف ۲-۱: یک خانواده گروهی از توزیع‌ها به وسیله انجام یک خانواده از تبدیلات روی یک متغیر تصادفی با یک توزیع خاص حاصل می‌شود.

در زیر سه خانواده‌ی مهم از خانواده‌ی گروه‌ها را معرفی می‌کنیم.

۱- خانواده‌ی مکانی^۲

فرض کنید U یک متغیر تصادفی با یک توزیع احتمال ثابت $f(u)$ باشد. اگر μ مقداری ثابت باشد

$$f_X(x) = f(x - \mu) \quad \text{از عبارت } X = U + \mu \text{ آنگاه توزیع}$$

توزیع حاصل برای توزیع احتمال ثابت f و $-\infty < \mu < +\infty$ را یک خانواده مکانی از توزیع‌ها گویند.

برای مثال توزیع $N(\mu, 1)$ که از توزیع ثابت $N(0, 1)$ با جمع با مقدار ثابت μ حاصل می‌شود.

$$Z \sim N(0, 1) \Rightarrow X = Z + \mu \sim N(\mu, 1), \quad f_X(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-\mu)^2} = f(x - \mu)$$

^۱ Group Families

^۲ Location Family

۲- خانواده مقیاسی^۱

فرض کنید U یک متغیر تصادفی با یک توزیع احتمال ثابت $f(u)$ باشد. اگر σ مقداری ثابت باشد

$$X = \sigma U \quad \text{عبارت است از} \quad f_X(x) = \frac{1}{\sigma} f\left(\frac{x}{\sigma}\right)$$

توزیع حاصل برای $\sigma > 0$ را یک خانواده مقیاسی از توزیع‌ها گویند.

برای مثال توزیع $E(\theta)$ که از توزیع ثابت $E(1)$ با ضرب در مقدار ثابت θ حاصل می‌گردد.

$$U \sim E(1) \Rightarrow X = \theta U \sim E(\theta), \quad f_X(x) = \frac{1}{\theta} e^{-\frac{x}{\theta}} = \frac{1}{\theta} f\left(\frac{x}{\theta}\right)$$

۳- خانواده مکانی - مقیاسی^۲

فرض کنید U یک متغیر تصادفی با یک توزیع احتمال ثابت $f(u)$ باشد. اگر μ و σ مقادیر ثابتی

$$X = \mu + \sigma U \quad \text{عبارت است از} \quad f_X(x) = \frac{1}{\sigma} f\left(\frac{x - \mu}{\sigma}\right)$$

توزیع حاصل برای $-\infty < \mu < +\infty$ و $\sigma > 0$ را یک خانواده مکانی - مقیاسی از توزیع‌ها گویند.

برای مثال توزیع $N(\mu, \sigma^2)$ از توزیع $N(0, 1)$ به صورت زیر حاصل می‌گردد.

$$U \sim N(0, 1) \Rightarrow X = \mu + \sigma U \sim N(\mu, \sigma^2),$$

$$f_X(x) = \frac{1}{\sigma} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2} = \frac{1}{\sigma} f\left(\frac{x-\mu}{\sigma}\right)$$

در زیر تعدادی از خانواده‌های مکانی مقیاسی آورده شده است.

$$(i) \quad E(\alpha, \beta) \quad f_X(x) = \frac{1}{\beta} e^{-\frac{1}{\beta}(x-\alpha)} I_{(\alpha, \infty)}(x)$$

$$(ii) \quad DE(\alpha, \beta) \quad f_X(x) = \frac{1}{2\beta} e^{-\frac{1}{\beta}|x-\alpha|}$$

$$(iii) \quad C(\alpha, \beta) \quad f_X(x) = \frac{1}{\pi \left\{1 + \left(\frac{x-\alpha}{\beta}\right)^2\right\}}$$

^۱ scale family

^۲ location - scale family

$$(iv) \quad U\left(\theta - \frac{\tau}{2}, \theta + \frac{\tau}{2}\right)$$

در هر یک از حالات فوق ما یک کلاس از تبدیلات G و یک عمل $*$ داریم که دارای دو خاصیت زیر است.

الف- کلاس G تحت عمل ترکیب بسته است یعنی

$$\text{if } g_1, g_2 \in G \quad \text{then} \quad g_1 * g_2 \in G$$

ب- کلاس G تحت عمل وارون بسته است یعنی

$$\text{if } g \in G \quad \text{then} \quad \exists g^{-1} \in G \quad \ni g * g^{-1} = e \in G$$

که e تبدیل همانی نام دارد.

یک کلاس از تبدیلات که تحت عمل ترکیب و وارون بسته باشد، یک گروه تبدیلات نامیده می‌شود. حال فرض کنید P یک توزیع خاص و معلوم مربوط به متغیر تصادفی X باشد. در این صورت مجموعه‌ی $G = \{g \mid Y_g = g(X)\}$ را یک خانواده گروهی از توزیع‌ها گویند.

برای مثال خانواده‌ی مکانی-مقیاسی یک خانواده گروهی از توزیع‌ها می‌باشد.

$$G = \{g \mid g(U) = \mu + \sigma U, \mu \in \mathbb{R}, \sigma > 0\}$$

$$\text{الف) } X_1 = g_1(U) = \mu_1 + \sigma_1 U$$

$$X_2 = g_2(X_1) = g_2(g_1(U)) = \mu_2 + \sigma_2(g_1(U)) = \mu_2 + \sigma_2(\mu_1 + \sigma_1 U)$$

$$= (\mu_2 \sigma_1 + \mu_1 \sigma_2) + \sigma_1 \sigma_2 U \Rightarrow g_2 * g_1 \in G$$

$$X = g(U) = \mu + \sigma U \Rightarrow g^{-1}(U) = \frac{U - \mu}{\sigma}$$

$$Y = g(g^{-1}(U)) = \mu + \sigma\left(\frac{U - \mu}{\sigma}\right) = U \Rightarrow g * g^{-1} = e \in G$$

یک خانواده از تبدیلات که در شرط $g_1 * g_2 = g_2 * g_1 \quad \forall g_1, g_2 \in G$ صدق کند، یک گروه

تبدیلات جابجایی^۱ نامیده می‌شود. برای مثال خانواده‌ی مکانی و خانواده‌ی مقیاسی یک گروه تبدیلات جابجایی هستند، اما خانواده مکانی - مقیاسی چنین نیست.

^۱ Commutative

مثال ۱-۲: فرض کنید $\underline{U} = (U_1, \dots, U_n)$ یک بردار تصادفی n بعدی باشد در این صورت می‌توان گروه تبدیلات زیر را روی این بردار تصادفی انجام داد.

$$\underline{U} + \underline{a} = (U_1 + a, U_2 + a, \dots, U_n + a) \quad a \in R$$

$$b\underline{U} = (bU_1, bU_2, \dots, bU_n)$$

$$a + b\underline{U} = (a + bU_1, a + bU_2, \dots, a + bU_n) \quad a \in R, b > 0 \quad \blacksquare$$

مثال ۲-۲: توزیع چند متغیره نرمال

فرض کنید $\underline{U} = (U_1, \dots, U_p)$ و $U_i \sim N(0, 1), i = 1, \dots, p$ و $U_1, \dots, U_p$ از یکدیگر مستقل باشند و تبدیل زیر را در نظر بگیرید

$$\begin{pmatrix} X_1 \\ \vdots \\ X_p \end{pmatrix} = \begin{pmatrix} a_1 \\ \vdots \\ a_p \end{pmatrix} + B \begin{pmatrix} U_1 \\ U_2 \\ \vdots \\ U_p \end{pmatrix} \Leftrightarrow \underline{X} = \underline{a} + B\underline{U}$$

که در آن B یک ماتریس وارون پذیر $p \times p$ است. توزیع $\underline{X}$ را توزیع نرمال چند متغیره با میانگین و واریانس زیر گویند

$$E(\underline{X}) = \underline{a} + BE(\underline{U}) = \underline{a}$$

$$\Sigma = V(\underline{X}) = BV(\underline{U})B' = BB'$$

برای بدست آوردن توزیع $\underline{X}$ توجه کنید که

$$f(\underline{u}) = (\pi)^{-\frac{p}{2}} e^{-\frac{1}{2} \sum_{i=1}^p u_i^2} = (\pi)^{-\frac{p}{2}} e^{-\frac{1}{2} \underline{u}' \underline{u}}$$

$$\underline{X} = \underline{a} + B\underline{U} \Rightarrow \underline{U} = B^{-1}(\underline{X} - \underline{a}) \Rightarrow |J| = |B^{-1}|$$

$$f_{\underline{X}}(\underline{x}) = |B^{-1}| f_{\underline{U}}(B^{-1}(\underline{x} - \underline{a})) = |B^{-1}| \left\{ e^{-\frac{1}{2}(\underline{x} - \underline{a})'(B^{-1})'B^{-1}(\underline{x} - \underline{a})} \right\} (\pi)^{-\frac{p}{2}}$$

$$= (\pi)^{-\frac{p}{2}} |B|^{-\frac{1}{2}} e^{-\frac{1}{2}(\underline{x} - \underline{a})' \Sigma^{-1}(\underline{x} - \underline{a})} \quad \blacksquare$$

بخش ۳: خانواده نمایی از توزیع‌ها

تعریف ۳-۱: یک خانواده $\{P_\theta: \theta \in \Theta\}$ از توزیع‌ها را یک خانواده‌ی نمایی k پارامتری گویند اگر توزیع P_θ برای هر $\theta \in \Theta$ و هر $x \in S_X^*$ (تکیه‌گاه X) دارای تابع چگالی بفرم زیر نسبت به اندازه‌ی μ باشد

$$p_\theta(x) = \left\{ e^{\sum_{j=1}^k c_j(\theta) T_j(x) - B(\theta)} \right\} h(x) \quad (1)$$

که در آن بایستی

۱- S_X^* بستگی به مقادیر $\theta = (\theta_1, \theta_2, \dots, \theta_k)$ نداشته باشد.

۲- $c_j(\theta)$ ، $j = 1, \dots, k$ توابعی غیر صفر و پیوسته از $\theta = (\theta_1, \dots, \theta_k)$ باشند.

۳- در حالت پیوسته $T_j'(x)$ ، $j = 1, 2, \dots, k$ توابعی پیوسته روی S_X^* و در حالت گسسته

$T_j(x)$ ، $j = 1, 2, \dots, k$ توابعی غیر صفر روی S_X^* باشند.

۴- در حالت پیوسته $h(x)$ تابعی پیوسته روی S_X^* باشد.

خانواده p_θ در (۱) را یک خانواده نمایی پر رتبه^۱ با بُعد k گویند هرگاه شرایط ۱ تا ۴ بالا برقرار باشد و علاوه بر آن

الف- توابع $c_j(\theta)$ ، $j = 1, 2, \dots, k$ به طور تابعی از یکدیگر مستقل باشند.

ب- توابع $T_j(x)$ ، $j = 1, 2, \dots, k$ به طور خطی از یکدیگر مستقل باشند.

ج- فضای پارامتری Θ یک شامل یک مستطیل K بعدی باشد. ■

مثال ۳-۱: خانواده توزیع‌های زیر به یک خانواده نمایی یک پارامتری متعلق می‌باشند.

$$(a) \quad B(1, \theta) \quad \theta \in (0, 1)$$

$$(g) \quad N(\cdot, \sigma^2) \quad \sigma > 0$$

$$(b) \quad B(n, \theta) \quad \theta \in (0, 1) \quad n \text{ is known}$$

$$(h) \quad N(\theta, \theta) \quad \theta > 0$$

$$(c) \quad NB(n, \theta) \quad \theta \in (0, 1) \quad n \text{ is known}$$

$$(i) \quad \Gamma(\alpha, \beta) \quad \beta > 0, \alpha \text{ known}$$

^۱ Full Rank

$$(d) \quad Ge(\theta) \quad \theta \in (0, 1)$$

$$(j) \quad \Gamma(\alpha, \beta) \quad \alpha > 0, \beta \text{ known}$$

$$(e) \quad P(\theta) \quad \theta \in (0, +\infty)$$

$$(k) \quad Be(\alpha, \beta) \quad \alpha > 0, \beta \text{ known}$$

$$(f) \quad N(\mu, 1) \quad \mu \in R$$

$$(l) \quad Be(\alpha, \beta) \quad \beta > 0, \alpha \text{ known} \quad \blacksquare$$

مثال ۳-۲: خانواده توزیع‌های زیر به یک خانواده نمایی دو پارامتری پررتبه متعلق است.

$$(a) \quad N(\mu, \sigma^2) \quad \mu \in R, \sigma > 0$$

$$(b) \quad \Gamma(\alpha, \beta) \quad \alpha > 0, \beta > 0$$

$$(c) \quad Be(\alpha, \beta), \alpha > 0, \beta > 0 \quad \blacksquare$$

مثال ۳-۳: خانواده‌ی $N(\theta, \theta^2), \theta > 0$ به یک خانواده نمایی متعلق است اما پررتبه نیست زیرا

$$p_{\theta}(x) = \frac{1}{\sqrt{2\pi\theta^2}} e^{-\frac{1}{2\theta^2}(x-\theta)^2} = \frac{1}{\sqrt{2\pi\theta^2}} e^{-\frac{1}{2\theta^2}x^2 + \frac{x}{\theta} - \frac{1}{2}} = \left\{ e^{-\frac{1}{2\theta^2}x^2 + \frac{x}{\theta} - \frac{1}{2}\ln(2\pi\theta^2)} \right\} e^{-\frac{1}{2}}$$

$$c_1(\theta) = -\frac{1}{2\theta^2} \quad T_1(x) = x^2 \quad h(x) = e^{-\frac{1}{2}}$$

$$c_2(\theta) = \frac{1}{\theta} \quad T_2(x) = x \quad B(\theta) = +\frac{1}{2}\ln(2\pi\theta^2)$$

پارامتر دو بُعدی $(-\frac{1}{2\theta^2}, \frac{1}{\theta})$ روی یک منحنی در فضای R^2 قرار می‌گیرد و به همین دلیل این

خانواده را خانواده نمایی منحنی^۱ گویند. $\blacksquare$

مثال ۳-۴: خانواده توزیع‌های زیر به یک خانواده نمایی متعلق نیستند.

$$(a) \quad U(\cdot, \theta), \theta > 0$$

$$(b) \quad E(\alpha, \beta), \alpha \in \mathbb{R}, \beta > 0 \quad \blacksquare$$

خانواده‌ی نمایی:

۱- اگر $X_1, X_2, \dots, X_n$ یک نمونه تصادفی (i.i.d) از خانواده نمایی k پارامتری باشند آنگاه

توزیع توأم $\underline{X} = (X_1, \dots, X_n)$ نیز به یک خانواده نمایی k -پارامتری متعلق است، زیرا

^۱ Curved Exponential Family

$$p_{\theta}(x_1, \dots, x_n) = \prod_{i=1}^n p_{\theta}(x_i) = \prod_{i=1}^n e^{\sum_{j=1}^k c_j(\theta) T_j(x_i) - B(\theta)} h(x_i)$$

$$= \dots = e^{\sum_{j=1}^k c_j(\theta) T_j^*(\underline{x}) - nB(\theta)} h^*(\underline{x})$$

$$T_j^*(\underline{x}) = \sum_{i=1}^n T_j(x_i) \quad \text{and} \quad h^*(\underline{x}) = \prod_{i=1}^n h(x_i) \quad \text{که در آن}$$

۲- در یک خانواده نمایی یک پارامتری با تابع چگالی

$$p_{\theta}(x) = e^{c(\theta)T(x) - B(\theta)} h(x)$$

$$E_{\theta}(T(x)) = \frac{B'(\theta)}{C'(\theta)} \quad \text{داریم که}$$

$$\int e^{c(\theta)T(x) - B(\theta)} h(x) d\mu(x) = 1 \Rightarrow \int [c'(\theta)T(x) - B'(\theta)] p_{\theta}(x) d\mu(x) = 0 \quad \text{زیرا}$$

$$\Rightarrow c'(\theta)E_{\theta}(T(x)) - B'(\theta) = 0 \Rightarrow E_{\theta}(T(x)) = \frac{B'(\theta)}{c'(\theta)}$$

فرم متعارف خانواده نمایی^۱

با تغییر پارامتر $\eta_j = c_j(\theta)$ ، $j = 1, 2, \dots, k$ در خانواده توزیع‌های نمایی (۱) داریم که

$$p_{\underline{\eta}}(x) = \left\{ e^{\sum_{j=1}^k \eta_j T_j(x) - A(\underline{\eta})} \right\} h(x) \quad (2)$$

که این فرم را فرم متعارف خانواده توزیع‌های نمایی گویند. در این حالت مجموعه‌ی

$$\{\underline{\eta} = (\eta_1, \dots, \eta_k) \mid \int_{-\infty}^{+\infty} h(x) e^{\sum_{j=1}^k \eta_j T_j(x)} d\mu(x) < \infty\}$$

را فضای پارامتری طبیعی و $\underline{\eta}$ را پارامتر طبیعی گویند. خواص زیر را برای این خانواده می‌توان اثبات کرد.

^۱ Canonical Form

۱- (کتاب TSH فصل ۲، قضیه ۹) برای هر تابع انتگرال پذیر f و هر η در فضای پارامتری طبیعی،

انتگرال $\int f(x) e^{\sum_{j=1}^k \eta_j T_j(x)} h(x) d\mu(x)$ پیوسته و دارای مشتقات تمام مراتب نسبت به η_j ها است و این مشتقات می توانند به درون انتگرال برده شوند.

۲- (کتاب TSH بخش ۲.۷ لم ۸) اگر X دارای توزیعی از خانواده نمایی با تابع چگالی بفرم (۲) نسبت به اندازه μ باشد آنگاه $T = (T_1, \dots, T_k)$ دارای توزیعی از خانواده نمایی با چگالی

$$p_{\eta}(t_1, \dots, t_k) = \left\{ e^{\sum_{j=1}^k \eta_j t_j - A(\eta)} \right\} k(t)$$

می گردد).

۳- فرض کنید $X_1, X_2, \dots, X_n$ متغیرهای تصادفی مستقل و هر کدام دارای توزیعی از خانواده نمایی با تابع چگالی زیر باشند

$$p_{\eta}(x_i) = \left\{ e^{\eta T_j(x_i) - A_i(\eta)} \right\} h_i(x_i)$$

در این صورت $\sum_{i=1}^n T_i(x_i)$ دارای توزیعی از خانواده نمایی یک پارامتری می باشد.

اگر دو خانواده نمایی k -پارامتری (با پارامترهای یکسان) مستقل داشته باشیم آنگاه $(T_1, \dots, T_k)$ و $(U_1, \dots, U_k)$ هر کدام دارای توزیعی از خانواده نمایی k پارامتری می باشند و طبق قضیه ۲، $(T_1 + U_1, \dots, T_k + U_k)$ نیز دارای توزیعی از خانواده نمایی k -پارامتری است. به وسیله استقرا قضیه ۳ ثابت می شود.

۴- اگر X دارای توزیعی از خانواده نمایی k -پارامتری با تابع چگالی بفرم (۲) باشد آنگاه

$$E_{\eta}(T_j(x)) = \frac{\partial}{\partial \eta_j} A(\eta)$$

$$\text{cov}_{\eta}(T_i(x), T_j(x)) = \frac{\partial^2}{\partial \eta_i \partial \eta_j} A(\eta)$$

مثال ۳-۵: فرض کنید $X = (X_1, \dots, X_k)$ دارای توزیع چند جمله ای با پارامترهای n و $p_1, \dots, p_k$

باشد. در این صورت:

$$\begin{aligned}
p(\underline{x}) &= P(X_1 = x_1, \dots, X_k = x_k) = \frac{n!}{x_1! x_2! \dots x_k!} p_1^{x_1} \dots p_k^{x_k} \\
&= h(\underline{x}) e^{x_1 \ln p_1 + \dots + x_k \ln p_k} \\
&= h(\underline{x}) e^{x_1 \ln \frac{p_1}{p_k} + x_2 \ln \frac{p_2}{p_k} + \dots + x_{k-1} \ln \frac{p_{k-1}}{p_k}} \times e^{x_k \ln p_k + \sum_{i=1}^{k-1} x_i \ln p_k} \\
&= h(\underline{x}) e^{\sum_{i=1}^{k-1} x_i \ln \frac{p_i}{p_k} + n \ln p_k}
\end{aligned}$$

قرار می دهیم $\eta_i = \ln \frac{p_i}{p_k}$ در این صورت

$$\begin{aligned}
A(\underline{\eta}) &= -n \ln p_k = n \ln \left(1 + \sum_{i=1}^{k-1} e^{\eta_i} \right) \\
(\text{since } e^{\eta_i} &= \frac{p_i}{p_k} \Rightarrow \sum_{i=1}^{k-1} e^{\eta_i} = \frac{1 - p_k}{p_k} = \frac{1}{p_k} - 1 \Rightarrow p_k = \left(1 + \sum_{i=1}^{k-1} e^{\eta_i} \right)^{-1}) \\
\therefore p_{\underline{\eta}}(\underline{x}) &= \left\{ e^{\sum_{i=1}^{k-1} \eta_i x_i - A(\underline{\eta})} \right\} h(\underline{x})
\end{aligned}$$

که $\underline{\eta} = (\eta_1, \dots, \eta_{k-1})$ پارامتر طبیعی است و

$$\begin{aligned}
E(X_i) &= \frac{\partial}{\partial \eta_i} A(\underline{\eta}) = \frac{n e^{\eta_i}}{1 + \sum_{i=1}^{k-1} e^{\eta_i}} = \frac{n \frac{p_i}{p_k}}{\frac{1}{p_k}} = n p_i \\
\text{cov}(X_i, X_j) &= \frac{\partial^2}{\partial \eta_i \partial \eta_j} A(\underline{\eta}) = \dots = -n p_i p_j \quad \blacksquare
\end{aligned}$$

تابع مولد گشتاور و تابع مولد انباشته^۱

اگر $\underline{T} = (T_1, \dots, T_k)$ آنگاه $\alpha_{r_1, \dots, r_k} = E(T_1^{r_1} \dots T_k^{r_k})$ را گشتاور توأم $T_1, \dots, T_k$ گویند و تابع مولد گشتاور $\underline{T} = (T_1, \dots, T_k)$ به صورت زیر تعریف می شود

$$M_{\underline{T}}(u_1, \dots, u_k) = E(e^{u_1 T_1 + \dots + u_k T_k})$$

^۱ Moment Generating fun. and Cumulant Generating fun.

اگر M_T در یک همسایگی حول مبدأ موجود باشد و گشتاورهای $\alpha_{r_1, \dots, r_s}$ موجود باشند آنگاه بسط تابع M_T به صورت زیر خواهد بود

$$M_T(u_1, \dots, u_k) = \sum_{r_1, \dots, r_k} \dots \sum \alpha_{r_1, \dots, r_k} \frac{u_1^{r_1} \dots u_k^{r_k}}{r_1! \dots r_k!}$$

همچنین تابع مولد انباشته T به صورت زیر تعریف می‌شود

$$K_T(u_1, \dots, u_k) = \log M_T(u_1, \dots, u_k) = \sum_{r_1, \dots, r_k} \dots \sum k_{r_1, \dots, r_k} \frac{u_1^{r_1} \dots u_k^{r_k}}{r_1! \dots r_k!}$$

که در آن $k_{r_1, \dots, r_k}$ را انباشته توأم $T = (T_1, \dots, T_k)$ گویند. اگر تنها یک T داشته باشیم آنگاه

$$M_T(u) = E(e^{uT}) = \sum_r \alpha_r \frac{u^r}{r!}$$

$$K_T(u) = \ln M_T(u) = \sum_r k_r \frac{u^r}{r!}$$

و می‌توان نشان داد که (تمرین)

$$\alpha_1 = k_1 \quad \alpha_2 = k_2 + k_1^2 \quad \alpha_3 = k_3 + 3k_1k_2 + k_1^3$$

در خانواده‌ی نمایی از توزیع‌ها توابع $M_T(u)$ و $K_T(u)$ در یک همسایگی نزدیک صفر موجود هستند و:

$$K_T(u) = A(\eta + u) - A(\eta) \quad M_T(u) = \frac{e^{A(\eta + u)}}{e^{A(\eta)}}$$

(اثبات: تمرین)

حال اگر $X_1, X_2, \dots, X_n$ از یکدیگر مستقل باشند و $X = X_1 + X_2 + \dots + X_n$ آنگاه

$$M_X(u) = E[e^{(X_1 + \dots + X_n)u}] = \prod_{i=1}^n M_{X_i}(u) \quad k_X(u) = \sum_{i=1}^n k_{X_i}(u)$$

مثال ۳-۶: فرض کنید $X \sim B(n, p)$ در این صورت

$$p(x) = P(X = x) = \binom{n}{x} p^x q^{n-x} = \binom{n}{x} e^{x \ln \frac{p}{q} + n \ln q}$$

قرار می دهیم $\eta = \ln \frac{p}{q}$ در این صورت

$$A(\eta) = -n \ln q = n \ln(1 + e^\eta) \quad p(x) = \binom{n}{x} e^{x\eta - A(\eta)}$$

$$M_X(u) = \frac{e^{A(\eta+u)}}{e^{A(\eta)}} = \frac{e^{n \ln(1+e^{\eta+u})}}{e^{n \ln(1+e^\eta)}} = \left(\frac{1+e^{\eta+u}}{1+e^\eta} \right)^n = \left(\frac{1+\frac{P}{q}e^u}{1+\frac{P}{q}} \right)^n = (q + pe^u)^n$$

مثال ۳-۷: فرض کنید $X \sim \Gamma(\alpha, \beta)$ در این صورت (α معلوم)

$$p(x) = \frac{1}{\Gamma(\alpha)\beta^\alpha} x^{\alpha-1} e^{-\frac{x}{\beta}} = \left[e^{-\frac{1}{\beta}x - \alpha \ln \beta} \right] \frac{x^{\alpha-1}}{\Gamma(\alpha)}$$

$$\eta = -\frac{1}{\beta} \quad A(\eta) = \alpha \ln \beta = -\alpha \ln \left(\frac{1}{\beta} \right) = -\alpha \ln(-\eta)$$

$$M_X(u) = \frac{e^{A(\eta+u)}}{e^{A(\eta)}} = \frac{e^{-\alpha \ln(-\eta-u)}}{e^{-\alpha \ln(-\eta)}} = \frac{\left(\frac{1}{\beta} - u \right)^{-\alpha}}{\left(\frac{1}{\beta} \right)^{-\alpha}} = \frac{1}{(1 - \beta u)^\alpha}$$

$$k_X(u) = -\alpha \ln(1 - \beta u) \quad u < \frac{1}{\beta}$$

بنابراین

$$E(X^r) = \frac{\partial^r}{\partial u^r} M_X(u) \Big|_{u=0} = \alpha(\alpha+1)\dots(\alpha+r-1)\beta^r = \frac{\Gamma(\alpha+r)}{\Gamma(\alpha)} \beta^r$$

■

بخش ۴: آماره‌های بسنده^۱

تعریف ۴-۱: فرض $X_1, \dots, X_n$ یک نمونه تصادفی از یک خانواده از توزیع‌های $P = \{P_\theta : \theta \in \Theta\}$ باشند (θ ممکن است یک بردار باشد) آماره $T(X_1, \dots, X_n) = T(\underline{X})$ را یک آماره‌ی بسنده برای θ و یا به عبارت دقیق‌تر یک آماره بسنده برای خانواده توزیع‌های P گویند هر گاه توزیع شرطی نمونه تصادفی $X_1, \dots, X_n$ به شرط $T(\underline{X}) = t$ برای هر مقدار t بستگی به θ نداشته باشد.

توجه کنید که آماره‌های بسنده داده‌ها را در خود خلاصه می‌کنند بدون این که اطلاعاتی در مورد پارامتر مجهول را از دست بدهیم. یعنی داده‌ها را می‌توان در آماره‌ی بسنده T خلاصه کرد و با داشتن مقدار T نیازی به داشتن خود داده‌ها (تا آن جا که مربوط به استنباط درباره پارامترهای مجهول می‌شود) نداریم.

قضیه ۴-۱: اگر $X | T = Y | T \stackrel{D}{=}$ آنگاه $X \stackrel{D}{=} Y$

اثبات:

$$\begin{aligned} \text{چگالی } X &= p_X(x) = \int p_{X|T}(x|t)q(t)d\mu(t) = \int p_{Y|T}(x|t)q(t)d\mu(t) \\ &= p_Y(x) = Y \text{ چگالی} \end{aligned}$$

عکس این قضیه برقرار نیست زیرا اگر در پرتاب سه سکه

X = تعداد شیرها Y = تعداد خط‌ها T = تعداد شیر - تعداد خط

آنگاه $X \stackrel{D}{=} Y$ اما

$$\begin{cases} P(X=1|T=1) = \frac{P(X=1, T=1)}{P(T=1)} = \frac{3/8}{3/8} = 1 \\ P(X \neq 1|T=1) = 0 \end{cases}$$

$$\begin{cases} P(Y \neq 2|T=1) = \frac{P(Y=1, T=1)}{P(T=1)} = 0 \\ P(Y=2|T=1) = 1 \end{cases}$$

$$\Rightarrow X | T \neq Y | T \stackrel{D}{=}$$

^۱ Sufficient Statistics

حال فرض کنید که مشاهدات X_1 را کنار گذاشته و مقدار $T = T(X_1)$ را نگه داریم. اکنون X_1' را از توزیع شرطی X_1 با فرض $T = t$ می‌سازیم (به وسیله شبیه‌سازی) چون $X_1' | T = t \stackrel{D}{=} X_1$ پس X_1' و X_1 یک توزیع دارند. در نتیجه با نگه داشتن مقدار T می‌توانیم مشاهداتی همانند (هم‌توزیع) X_1 بسازیم. این کار را با آماره‌های غیربسنده نمی‌توان انجام داد، چون توزیع شرطی X_1 با فرض $T = t$ به پارامتر مجهول بستگی دارد و نمی‌توان X_1' را بازیافت نمود.

مثال ۴-۱: فرض کنید X_1, X_2 یک نمونه تصادفی از توزیع $P(\theta)$ باشد. نشان دهید که $T(X_1) = X_1 + X_2$ یک آماره بسنده برای θ است و از روی آن نمونه‌ی (X_1', X_2') هم توزیع با (X_1, X_2) را بسازید.

$$\begin{aligned}
 P(X_1 = x_1, X_2 = x_2 | T(X_1) = t) &= \begin{cases} \frac{P(X_1 = x_1, X_2 = x_2)}{P(T = t)} & x_1 + x_2 = t \\ \cdot & x_1 + x_2 \neq t \end{cases} \\
 &= \begin{cases} \frac{e^{-\theta} \theta^{x_1}}{x_1!} \cdot \frac{e^{-\theta} \theta^{x_2}}{x_2!} & x_1 + x_2 = t \\ \frac{e^{-2\theta} (2\theta)^t}{t!} & x_1 + x_2 \neq t \end{cases} = \begin{cases} \frac{t!}{x_1! x_2! 2^t} & x_1 + x_2 = t \\ \cdot & x_1 + x_2 \neq t \end{cases} \\
 &= \begin{cases} \binom{t}{x_1} \left(\frac{1}{2}\right)^{x_1} \left(\frac{1}{2}\right)^{t-x_1} & x_1 + x_2 = t \\ \cdot & x_1 + x_2 \neq t \end{cases} \\
 &\Rightarrow X_1' | T = t \sim B\left(t, \frac{1}{2}\right)
 \end{aligned}$$

چون توزیع شرطی به θ بستگی ندارد، پس $T(X_1) = X_1 + X_2$ یک آماره‌ی بسنده برای θ است. حال اگر $X_1' = t - X_2'$ و $X_2' = t - X_1'$ را به ترتیب تعداد شیرها و تعداد خط‌ها در t پرتاب یک سکه در نظر بگیریم آن گاه $X_1' | t \sim B\left(t, \frac{1}{2}\right)$ که باتوجه به قضیه ۴-۱، $X_1' \stackrel{D}{=} X_1$ است. ■

استفاده از تعریف ۴-۱ برای بررسی بسندگی یک آماره، نیاز به در دست داشتن آماره بسنده $T(X_1)$ و همچنین بررسی توزیع شرطی نمونه به شرط T دارد. یک راه ساده برای بدست آوردن یک آماره‌ی

بسندگی و بررسی بسندگی آن در قضیه دسته‌بندی نیمن که به قضیه تجزیه به عوامل مشهور است، آمده است.

قضیه ۴-۲: (قضیه تجزیه به عوامل)^۱

فرض کنید $X_1, X_2, \dots, X_n$ دارای توزیع توأم متعلق به خانواده‌ی $P = \{P_\theta : \theta \in \Theta\}$ که به وسیله اندازه‌ی μ مغلوب شده است، باشند. شرط لازم و کافی برای آن که آماره $T(X)$ برای خانواده‌ی P بسنده باشد این است که توابع نامنفی g و h وجود داشته باشند بطوری که تابع چگالی p_θ مربوط به P_θ در رابطه زیر صدق کند.

$$p_\theta(x) = g(\theta, T(x))h(x) \quad (a.e.\mu)$$

■

اثبات: (کتاب TSH بخش ۶.۲ قضیه ۸ و نتیجه ۱)

مثال ۴-۲: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از $Be(\theta, 1)$ باشند. یک آماره‌ی بسنده برای θ پیدا کنید.

$$f_\theta(x) = \prod_{i=1}^n \theta x_i^{\theta-1} = \begin{cases} \theta^n (\prod_{i=1}^n x_i)^{\theta-1} & = g(\theta, \prod_{i=1}^n x_i) \times 1 \\ \theta^n e^{(\theta-1) \sum_{i=1}^n \ln x_i} & = g(\theta, \sum_{i=1}^n \ln x_i) \times 1 \end{cases}$$

پس $T(X) = \sum_{i=1}^n \ln X_i$ و $S(X) = \prod_{i=1}^n X_i$ برای θ بسنده می‌باشند.

قضیه ۳: اگر T یک آماره‌ی بسنده برای θ و $T = k(U)$ آنگاه U نیز برای θ بسنده است.

اثبات: T بسنده است، بنابراین

$$p_\theta(x) = g(\theta, T(x))h(x) = g(\theta, k(U(x)))h(x) = g^*(\theta, U(x))h(x)$$

■

پس طبق قضیه تجزیه به عوامل $U(X)$ یک آماره بسنده برای θ است.

بنابراین در مثال ۴-۲ $T = \ln(S)$ و $S = e^T$ هر دو برای θ بسنده هستند.

^۱ Factorization Theorem.

مثال ۳-۴: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از $N(\mu, \sigma^2)$ باشند. نشان دهید که $(\bar{X}, S^2)$ یک آماره بسنده برای θ است.

مثال ۴-۴: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از $U(0, \theta)$ باشند. نشان دهید که $X_{(n)}$ یک آماره ی بسنده برای θ است.

مثال ۵-۴: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از یک توزیع P_θ باشند و قرار می‌دهیم $T = (X_{(1)}, \dots, X_{(n)})$ چون کلیه مقادیر $\tilde{X}$ برابر $n!$ حالت $(X_{(1)}, X_{(2)}, \dots, X_{(n)})$ می‌باشد که توزیع شرطی $\tilde{X}$ به شرط هر کدام از این مقادیر برابر $\frac{1}{n!}$ است و به θ بستگی ندارد پس $T(\tilde{X})$ یک آماره بسنده برای θ است.

مثال ۶-۴: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از خانواده ی نمایی k پارامتری باشند در این صورت

$$p_\theta(\tilde{x}) = \left[e^{\sum_{j=1}^k C_j(\theta) \sum_{i=1}^n T_j(x_i) - nB(\theta)} \right] \prod_{i=1}^n h(x_i) = g(\theta, \sum_{i=1}^n T_1(x_i), \dots, \sum_{i=1}^n T_k(x_i))$$

پس $T = (\sum_{i=1}^n T_1(X_i), \dots, \sum_{i=1}^n T_k(X_i))$ یک آماره ی بسنده برای θ است.

آماره بسنده ی مینیمال:

فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع $N(\theta, 1)$ باشند و قرار دهید

$$T_1(\tilde{X}) = (X_1, X_2, \dots, X_n) \quad T_2(\tilde{X}) = \left(\sum_{i=1}^m X_i, \sum_{i=m+1}^n X_i \right) \quad T_3(\tilde{X}) = \sum_{i=1}^n X_i$$

طبق قضیه ی تجزیه به عوامل $T_3(\tilde{X})$ یک آماره ی بسنده برای θ است $T_3 = k_3(T_2)$ و $T_2 = k_2(T_1)$ پس طبق قضیه ۳-۴، T_1 و T_2 نیز برای θ بسنده هستند. همچنین T_3 بیشتر از T_2 و T_2 بیشتر از T_1 داده‌ها را خلاصه می‌کنند.

حال این سوال مطرح است که داده‌ها را تا چه اندازه می‌توان خلاصه کرد. در پاسخ به این سوال بایستی گفت تا جایی می‌توان داده‌ها را خلاصه کرد که اطلاعات داده‌ها در مورد پارامترها از دست نرود. آماره‌ی بسنده‌ای که تا حد ممکن داده‌ها را خلاصه کند آماره‌ی بسنده‌ی مینیمال نامیده می‌شود.

تعریف ۴-۲: آماره‌ی $T(\underline{X})$ یک آماره‌ی بسنده‌ی مینیمال است اگر اولاً این آماره بسنده باشد و ثانیاً برای هر آماره‌ی بسنده‌ی دیگر U آماره‌ی T تابعی از آن باشد یعنی $T = K(U)$.

نکات: ۱- آماره‌های بسنده مینیمال یکتا نیستند اما اگر T_1 و T_2 بسنده‌ی مینیمال باشند یک رابطه یک به یک بین آنها برقرار است.

۲- اگر T یک آماره‌ی بسنده‌ی مینیمال باشد و $U = K(T)$ و K یک تابع یک به یک نباشد آن گاه U یک آماره‌ی بسنده نیست.

به وسیله قضایای زیر روش بدست آوردن آماره‌ی بسنده‌ی مینیمال را بررسی می‌کنیم.

قضیه ۴-۳: آماره‌ی U برای θ بسنده است اگر و فقط اگر به ازای هر $\theta_1, \theta_2 \in \Theta$ ، $\frac{p(x; \theta_1)}{p(x; \theta_2)}$ تابعی از U باشد.

اثبات: (کفایت)

$$\frac{p(x; \theta_1)}{p(x; \theta_2)} = \ell(u; \theta_1, \theta_2) \Rightarrow p(x; \theta_1) = \ell(u; \theta_1, \theta_2) p(x; \theta_2)$$

اگر θ_2 را برابر مقدار خاص θ بگیریم آن گاه

$$\begin{aligned} p(x; \theta_1) &= \ell(u; \theta_1, \theta) p(x; \theta) \quad \forall \theta_1 \\ &= g(U, \theta_1) h(x) \quad \forall \theta_1 \end{aligned}$$

پس طبق قضیه تجزیه به عوامل U برای θ بسنده است.

(لزوم): اگر U برای θ بسنده باشد آن گاه طبق قضیه تجزیه به عوامل داریم که

$$\blacksquare \quad \left. \begin{aligned} p(\underline{x}; \theta_1) &= g(\theta_1, u) h(\underline{x}) \\ p(\underline{x}; \theta_2) &= g(\theta_2, u) h(\underline{x}) \end{aligned} \right\} \Rightarrow \frac{p(\underline{x}; \theta_1)}{p(\underline{x}; \theta_2)} = \frac{g(\theta_1, u)}{g(\theta_2, u)} = \ell(u; \theta_1, \theta_2)$$

قضیه ۴-۵: فرض کنید $P = \{p_., p_1, \dots, p_k\}$ خانواده‌ای متناهی متشکل از $k+1$ توزیع باشد که

همگی دارای تکیه‌گاه یکسان هستند. در این صورت $T(x) = \left(\frac{p_1(x)}{p_.(x)}, \dots, \frac{p_k(x)}{p_.(x)} \right)$ یک آماره

بسندۀ مینیمال برای θ است.

اثبات:

$$p_j(x) = \frac{p_j(x)}{p_.(x)} p_.(x) = g(T(x), j) h(x) \quad j = 1, \dots, k$$

$$p_.(x) = 1 \times p_.(x) = 1 \times h(x)$$

لذا طبق قضیه تجزیه به عوامل $T(\tilde{X})$ یک آماره بسندۀ است. حال فرض کنید U یک آماره‌ی بسندۀ

برای P باشد. طبق قضیه ۴-۴ بایستی T تابعی از U باشد پس T یک آماره‌ی بسندۀ مینیمال

است. ■

قضیه ۴-۶: فرض کنید P خانواده‌ای از توزیع‌ها با تکیه‌گاه یکسان و $P_0 \subset P$ باشد. اگر T

برای P_0 بسندۀ مینیمال و برای P بسندۀ باشد، آنگاه برای P بسندۀ مینیمال نیز هست.

اثبات: فرض کنید U یک آماره‌ی بسندۀ برای P است. پس U برای P_0 نیز بسندۀ است و لذا T

تابعی از U است، پس T برای P بسندۀ مینیمال است. ■

مثال ۴-۷: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع $N(\theta, 1)$ باشد. یک آماره‌ی

بسندۀ مینیمال برای θ پیدا کنید.

$$P_0 = \{N(\theta_0, 1), N(\theta_1, 1)\}$$

$$\frac{p_{\theta_1}(x)}{p_{\theta_0}(x)} = \frac{e^{-\frac{1}{2} \sum_{i=1}^n (x_i - \theta_1)^2}}{e^{-\frac{1}{2} \sum_{i=1}^n (x_i - \theta_0)^2}} = e^{(\theta_1 - \theta_0) \sum_{i=1}^n x_i - \frac{1}{2}(\theta_1^2 - \theta_0^2)} = h\left(\sum_{i=1}^n x_i\right)$$

که h یک تابع یک به یک است. چون $\sum_{i=1}^n x_i$ برای P بسندۀ و برای P_0 بسندۀ مینیمال است پس

برای خانواده‌ی $N(\theta, 1)$ بسندۀ مینیمال است.

مثال ۴-۸: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع $U(\theta - \frac{1}{4}, \theta + \frac{1}{4})$ باشد. یک

$$T(\underline{X}) = (X_{(1)}, X_{(n)})$$

آماره‌ی بسنده‌ی مینیمال برای θ پیدا کنید.

قضیه ۴-۷: فرض کنید X دارای توزیعی از خانواده‌ی نمایی k -پارامتری باشد. در این صورت $T = (T_1, \dots, T_k)$ یک آماره‌ی بسنده‌ی مینیمال برای این خانواده است به شرط آنکه یکی از شرایط زیر برای این خانواده برقرار باشد:

الف- این خانواده پررتبه باشد.

ب- فضای پارامتری شامل $k+1$ نقطه $\eta^{(1)}, \eta^{(2)}, \dots, \eta^{(k)}$ باشد که فضای E_k را پدید می‌آورند. به این معنی که آنها به یک زیر فضای محض E_k تعلق ندارند.

اثبات: اگر شرایط (الف) یا (ب) برقرار باشد آنگاه می‌توان $\eta^{(1)}, \dots, \eta^{(k)}$ را در داخل فضای پارامتری چنان یافت که مستقل خطی باشند. حال خانواده توزیع‌های $\{p_{\eta^{(1)}}, \dots, p_{\eta^{(k)}}\}$ را در P نظر می‌گیریم که در آن $\eta^{(i)}$ یک ترکیب خطی غیرمحدب از $\eta^{(1)}, \dots, \eta^{(k)}$ است. بنابراین

$$T(\underline{X}) = \left(\frac{p(\underline{x}, \eta^{(1)})}{p(\underline{x}, \eta^{(i)})}, \dots, \frac{p(\underline{x}, \eta^{(k)})}{p(\underline{x}, \eta^{(i)})} \right) = \left(e^{\sum_{r=1}^k (\eta_r^{(1)} - \eta_r^{(i)}) t_r}, \dots, e^{\sum_{r=1}^k (\eta_r^{(k)} - \eta_r^{(i)}) t_r} \right)$$

و یا معادلاً

$$W(\underline{x}) = \left(\sum_{r=1}^k (\eta_r^{(1)} - \eta_r^{(i)}) t_r, \dots, \sum_{r=1}^k (\eta_r^{(k)} - \eta_r^{(i)}) t_r \right) \\ = [t_1, \dots, t_k] \begin{bmatrix} \eta_1^{(1)} - \eta_1^{(i)} & \dots & \eta_1^{(k)} - \eta_1^{(i)} \\ \vdots & & \vdots \\ \eta_k^{(1)} - \eta_k^{(i)} & \dots & \eta_k^{(k)} - \eta_k^{(i)} \end{bmatrix} = [t_1, \dots, t_k] B$$

یک آماره‌ی بسنده‌ی مینیمال برای P_0 است. چون $\eta^{(i)}$ ها مستقل خطی هستند پس ستون‌های B مستقل خطی و در نتیجه B وارون پذیر است. بنابراین $(T_1, \dots, T_k)$ برای P_0 بسنده‌ی مینیمال است و چون برای P بسنده است پس برای P بسنده مینیمال نیز هست. ■

مثال ۴-۹: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از $N(\theta, \theta^2)$ باشند. در این صورت

$$P_{\theta}(x) = \{e^{-\frac{1}{2\theta^2} \sum x_i^2 + \frac{1}{\theta} \sum x_i - n \ln(2\pi\theta^2)}\} e^{-\frac{1}{n}}$$

بنابراین آماره $T = (\sum X_i, \sum X_i^2)$ برای θ بسنده است و این آماره برای θ بسنده مینیمال است

زیرا در اینجا فضای پارامتری عبارت از $\eta = (\frac{1}{\theta}, -\frac{1}{2\theta^2})$ است و اگر نقاط

$$\eta^{(1)} = (1, -\frac{1}{2}) \quad \eta^{(2)} = (\frac{1}{2}, -\frac{1}{8}) \quad \eta^{(3)} = (\frac{1}{3}, -\frac{1}{18})$$

را در نظر بگیریم آن گاه ماتریس B فوق عبارت است از

$$\begin{bmatrix} \frac{1}{2} - 1 & \frac{1}{3} - 1 \\ -\frac{1}{8} + \frac{1}{2} & -\frac{1}{18} + \frac{1}{2} \end{bmatrix}$$

که وارون پذیر است

و طبق قضیه قبل T برای θ بسنده مینیمال است. ■

روش لهن - شفه برای بدست آوردن آماره بسنده مینیمال

قضیه ۴-۸: خانواده توزیع های $\{P_{\theta} : \theta \in \Theta\}$ با چگالی p_{θ} را در نظر بگیرید. اگر $T(X)$ آماره ای

باشد که برای هر دو نقطه x و y در تکیه گاه S_X^* داشته باشیم که

$$\frac{p_{\theta}(x)}{p_{\theta}(y)} = \theta \text{ مستقل از } x, y \Leftrightarrow T(x) = T(y)$$

آن گاه T برای این خانواده بسنده مینیمال است.

اثبات: ابتدا ثابت می کنیم که $T(X)$ یک آماره بسنده است. افرازهای ایجاد شده توسط $T(X)$ را به

صورت $A_t = \{x | T(x) = t\}$ نشان می دهیم. فرض کنید U_t یک نقطه به خصوص در A_t باشد. حال

اگر $x \in A_t$ باشد در این صورت

$$\left. \begin{matrix} x \in A_t \\ U_t \in A_t \end{matrix} \right\} \Rightarrow T(x) = T(U_t) \Rightarrow p_{\theta}(x) = K(x, U_t) p_{\theta}(U_t)$$

$$\Rightarrow p_{\theta}(x) = K(x, U_{T(x)}) f_{\theta}(U_{T(x)}) = h(x) g(\theta, T(x))$$

پس بنا به قضیه تجزیه به عوامل، $T(X)$ یک آماره بسنده است.

حال فرض کنید U یک آماره بسنده دیگر باشد در این صورت

$$U(\underline{x}) = U(\underline{y}) \Rightarrow \frac{P_{\theta}(\underline{x})}{P_{\theta}(\underline{y})} = \frac{g(\theta, U(\underline{x}))}{g(\theta, U(\underline{y}))} = \frac{h(\underline{x})}{h(\underline{y})} = K(\underline{x}, \underline{y}) \quad \theta \text{ مستقل از}$$

لذا بنا به فرض $T(\underline{x}) = T(\underline{y})$ یعنی T تابعی از U است پس T بسنده مینیمال است. ■

مثال ۴-۱۰: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع $Pa(\alpha, \theta)$ باشند یک آماره بسنده مینیمال برای (α, θ) پیدا کنید.

$$p_{\alpha, \theta}(\underline{x}) = \prod_{i=1}^n \frac{\theta \alpha^{\theta}}{x_i^{\theta+1}} I_{(\alpha, +\infty)}(x_i) = \frac{\theta^n \alpha^{n\theta}}{(\prod_{i=1}^n x_i)^{\theta+1}} I_{(\cdot, x_{(1)})}(\alpha)$$

$$\frac{p_{\alpha, \theta}(\underline{x})}{p_{\alpha, \theta}(\underline{y})} = \frac{(\prod_{i=1}^n y_i)^{\theta+1}}{(\prod_{i=1}^n x_i)^{\theta+1}} \cdot \frac{I_{(\cdot, x_{(1)})}(\alpha)}{I_{(\cdot, y_{(1)})}(\alpha)} = 1 \Leftrightarrow (\prod_{i=1}^n x_i, x_{(1)}) = (\prod_{i=1}^n y_i, y_{(1)})$$

پس $T(\underline{X}) = (\prod_{i=1}^n X_i, X_{(1)})$ یک آماره بسنده مینیمال برای (α, θ) است. ■

توجه: با توجه به رابطه آماره‌های بسنده و آماره‌های بسنده مینیمال، روش تشخیص بسندگی یا عدم بسندگی یک آماره به صورت زیر می باشد. فرض کنید T یک آماره دلخواه باشد. برای تشخیص بسندگی یا عدم بسندگی آماره T ، ابتدا توسط روش لهن - شفه یک آماره بسنده مینیمال S پیدا می‌کنیم اگر S تابعی از آماره T باشد آنگاه T بسنده است و اگر S تابعی از T نباشد. آن گاه T یک آماره بسنده نخواهد بود (اثبات: برهان خلف).

آماره‌های فرعی^۱:

تعریف ۴-۳: یک آماره $T(\underline{X})$ را فرعی گوئیم هرگاه توزیع آن به پارامتر مجهول بستگی نداشته باشد و یک آماره $T(\underline{X})$ را فرعی از مرتبه اول گوئیم هرگاه $E_{\theta}[T(\underline{X})] = c$ و به θ بستگی نداشته باشد.

مثال ۴-۱۱: اگر $X_1, \dots, X_n$ یک نمونه تصادفی از یک خانواده مکان $f(x - \theta)$ باشند آنگاه $S(\underline{X}) = X_{(n)} - X_{(1)}$ یک آماره فرعی است زیرا

^۱ Ancillary statistics

$$X_i = Z_i + \theta \quad i = 1, \dots, n$$

Z_i دارای تابع چگالی $f(x)$ است که به θ بستگی ندارد.

$$\begin{aligned} R = X_{(n)} - X_{(1)} &\Rightarrow F_R(r) = P(R \leq r) = P(X_{(n)} - X_{(1)} \leq r) \\ &= P(\max_{1 \leq i \leq n} X_i - \min_{1 \leq i \leq n} X_i \leq r) = P(\max_{1 \leq i \leq n} (Z_i + \theta) - \min_{1 \leq i \leq n} (Z_i + \theta) \leq r) \\ &= P(Z_{(n)} - Z_{(1)} \leq r) \end{aligned}$$

که این احتمال به θ بستگی ندارد پس R یک آماره فرعی است. ■

مثال ۴-۱۲: اگر $X_1, \dots, X_n$ یک نمونه تصادفی از یک خانواده مقیاس $\frac{1}{\sigma} f\left(\frac{x}{\sigma}\right)$ باشند آنگاه

$$S(\tilde{X}) = (Y_1, \dots, Y_{n-1}) = \left(\frac{X_1}{X_n}, \dots, \frac{X_{n-1}}{X_n}\right)$$

یک آماره فرعی برای σ است زیرا

$$\begin{aligned} F_{Y_1, \dots, Y_{n-1}}(y_1, \dots, y_{n-1}) &= P\left(\frac{X_1}{X_n} \leq y_1, \dots, \frac{X_{n-1}}{X_n} \leq y_{n-1}\right) \\ &= P\left(\frac{X_1/\sigma}{X_n/\sigma} \leq y_1, \dots, \frac{X_{n-1}/\sigma}{X_n/\sigma} \leq y_{n-1}\right) \\ &= P\left(\frac{Z_1}{Z_n} \leq y_1, \dots, \frac{Z_{n-1}}{Z_n} \leq y_{n-1}\right) \end{aligned}$$

که در آن $Z_i = \frac{X_i}{\sigma}$ و $f_{Z_i}(z) = \sigma f_X(\sigma z) = f(z)$ بنابراین احتمال فوق به σ بستگی ندارد

پس $S(\tilde{X})$ یک آماره فرعی است. ■

توجه: همانگونه که می‌دانیم آماره‌های بسنده جهت کاهش بُعد داده‌ها مورد استفاده قرار می‌گیرند. آن آماره بسنده‌ای می‌تواند این کاهش بُعد را تا حد ممکن انجام دهد که هیچ تابع غیرثابت آن یک آماره فرعی نباشد یعنی اگر $E[f(T)] = c$ نتیجه دهد که $f(T) = c$ a.e. که با کم کردن c از طرفین، تعریف آماره کامل حاصل می‌شود.

بخش ۵: کامل بودن^۱

فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از خانواده توزیع‌های $\{P_\theta : \theta \in \Theta\}$ باشند P باشد اگر $T(X)$ یک آماره باشد آنگاه گویند $P^T = \{P_\theta^T : \theta \in \Theta\}$ خانواده توزیع‌های تولید شده توسط آماره T است.

تعریف ۵-۱: الف- خانواده P_θ^T را کامل گویند اگر برای هر تابع $h(T)$ داشته باشیم که

$$E_\theta[h(T)] = 0 \quad \forall \theta \in \Theta \Rightarrow h(T) = 0 \quad a.e. \quad P^T \quad (۱)$$

و آماره $T(X)$ را کامل گویند اگر خانواده توزیع‌های به وجود آمده توسط T کامل باشد (یعنی خانواده‌ای کامل است که در آن تنها برآوردگر نااریب صفر خود برآوردگر صفر باشد).

ب- خانواده P_θ^T را کامل کراندار گویند اگر برای هر تابع حقیقی کراندار $h(T)$ رابطه (۱) برقرار باشد.

توجه: توجه کنید که اگر یک آماره T کامل باشد آن گاه کامل کراندار نیز است اما عکس این مطلب برقرار نیست. و همچنین اگر T یک آماره کامل باشد آنگاه هر تابع یک به یک آن نیز کامل است.

مثال ۱: فرض کنید $X \sim B(\theta, \frac{1}{p})$ ، نشان دهید که

الف- اگر $\theta \in \Theta = \{0, 1, 2, \dots\}$ آنگاه $\{B(\theta, \frac{1}{p}) \mid \theta \in \Theta\}$ کامل است.

ب- اگر $\theta \in \Theta^* = \{1, 2, \dots\}$ آنگاه $\{B(\theta, \frac{1}{p}) \mid \theta \in \Theta^*\}$ کامل نیست.

حل:

$$E_\theta[h(X)] = 0 \Rightarrow \sum_{x=0}^{\theta} h(x) \binom{\theta}{x} \left(\frac{1}{p}\right)^x = 0 \Rightarrow \sum_{x=0}^{\theta} h(x) \binom{\theta}{x} = 0 \quad \forall \theta \in \Theta$$

الف-

$$\theta = 0 \Rightarrow h(0) = 0$$

$$\theta = 1 \Rightarrow h(0) + h(1) = 0 \Rightarrow h(1) = 0$$

$$\theta = k : h(0) = h(1) = \dots = h(k) = 0 \quad \text{فرض}$$

^۱ Completeness

$$\theta = k + 1: h(k + 1) = \cdot \quad \text{حکم}$$

$$h(\cdot) + h(1) \binom{k+1}{1} + \dots + h(k) \binom{k+1}{k} + h(k+1) = \cdot \stackrel{(2)}{\Rightarrow} h(k+1) = \cdot$$

$$\therefore h(x) = \cdot \quad a.e. P$$

که نتیجه (۲) از فرض حاصل می گردد

ب-

$$\theta = 1 \Rightarrow h(\cdot) + h(1) = \cdot \Rightarrow h(1) = -h(\cdot) = -a$$

$$\theta = 2 \Rightarrow h(\cdot) + 2h(1) + h(2) = \cdot \Rightarrow h(2) = -a$$

⋮

$$h(x) = \begin{cases} a & x = \cdot, 2, \dots \\ -a & x = 1, 3, \dots \end{cases} = (1)^x a, \quad a = h(\cdot)$$

پس اگر $a \neq \cdot$ آنگاه P^* یک خانواده کامل نیست. ■

توجه: مثال قبل نشان می دهد که کامل بودن به خانواده توزیع ها بستگی دارد و نه به آماره $T(X)$.

مثال ۵-۲: فرض کنید X یک متغیر تصادفی گسسته با تابع احتمال زیر باشد.

$$p_\theta(x) = \begin{cases} \theta & x = 1 \\ \theta^{x-2}(1-\theta)^2 & x = 2, 3, \dots \end{cases} \quad 0 < \theta < 1$$

نشان دهید که X یک آماره کامل نیست اما کامل کراندار است.

حل:

$$\cdot = E_\theta[h(X)] = \sum_{x=1}^{+\infty} h(x) p_\theta(x) = \theta h(1) + \sum_{x=2}^{+\infty} h(x) \theta^{x-2} (1-\theta)^2$$

$$\Rightarrow \sum_{x=2}^{+\infty} h(x) \theta^{x-2} = -\theta h(1) (1-\theta)^{-2}$$

$$\text{با توجه به اینکه } \frac{1}{(1-\theta)^2} = \sum_{x=\cdot}^{\infty} (x+1) \theta^x \text{ داریم:}$$

$$\Rightarrow \sum_{y=\cdot}^{y=x-2} h(y+2) \theta^y = -h(1) \sum_{x=\cdot}^{+\infty} (x+1) \theta^{x+1} = -h(1) \sum_{y=1}^{+\infty} y \theta^y$$

$$\Rightarrow \begin{cases} h(2) = 0 \\ h(y+2) = -yh(1) \end{cases} \quad y = 1, 2, 3, \dots \Rightarrow h(x) = -(x-2)h(1) \quad x = 1, 2, 3, \dots$$

بنابراین X کامل نیست زیرا برای $E_\theta[h(X)] = 0$ نتیجه گرفتیم که $h(x) \neq 0$ (با قرار دادن $h(1) \neq 0$). اما X یک آماره کامل کراندار است زیرا اگر $h(x)$ یک تابع کراندار باشد و $h(1) \neq 0$ آنگاه $\lim_{x \rightarrow \pm\infty} h(x) = \pm\infty$. بنابراین برای اینکه $h(x)$ یک تابع حقیقی کراندار باشد بایستی $h(1) = 0$.

باشد پس $h(x) = 0, \quad \forall x = 1, 2, \dots$ ■

قضیه ۵-۱: اگر X دارای توزیعی از خانواده نمایی k -پارامتری پررتبه باشد آنگاه آماره $T_k = (T_1, \dots, T_k)$ برای این خانواده کامل است.

اثبات: (کتاب TSH بخش ۳. ۴ قضیه ۱) ■

با توجه به قضیه فوق می‌توان نتیجه گرفت که در توزیع $N(\mu, \sigma^2)$ آماره $T(X) = (\bar{X}, S^2)$ و در توزیع $B(n, \theta)$ آماره $T(X) = \bar{X}$ و در توزیع $Be(\theta, 1)$ آماره $T(X) = \sum \ln X_i$ آماره‌های کامل (و بسنده مینیمال) می‌باشند. مثالهای دیگر در غیر از خانواده نمایی در کتاب خوانده شود.

همچنین توجه کنید که اگر خانواده نمایی پررتبه نباشد آنگاه آماره $T_k = (T_1, \dots, T_k)$ کامل نیست. برای مثال $N(\theta, \theta^2)$.

مثال ۵-۳: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از $U(\theta, \theta+1)$ باشد. در این صورت به راحتی

دیده می‌شود که $U(X) = (X_{(1)}, X_{(n)})$ و یا معادلاً $T(X) = (X_{(n)} - X_{(1)}, \frac{X_{(1)} + X_{(n)}}{2})$ یک

آماره بسنده مینیمال توأم برای θ است و $S(X) = X_{(n)} - X_{(1)}$ یک آماره فرعی است و چون S قسمتی از آماره T است پس S و T از یکدیگر مستقل نمی‌باشند.

توجه: توجه کنید که آماره بسنده مینیمال تمام اطلاعات را در مورد پارامتر مجهول θ دارد اما آماره فرعی هیچ اطلاعی در مورد θ ندارد و به نظر می‌رسد که بایستی از یکدیگر مستقل باشند. اما در مثال بالا چنین نبود.

قضیه ۵-۲ (قضیه باسو^۱): اگر T یک آماره بسنده و کامل کراندار برای خانواده $\mathcal{P} = \{p_\theta : \theta \in \Theta\}$ باشد و V یک آماره فرعی برای $\mathcal{P}$ باشد آنگاه T و V از یکدیگر مستقل هستند.

اثبات: فرض کنید E یک پیشامد باشد و قرار دهیم $I_E(w) = \begin{cases} 1 & w \in E \\ 0 & w \notin E \end{cases}$ در این صورت I_E

یک متغیر تصادفی کراندار است و $E(I_E) = P(E)$. حال نشان می‌دهیم که برای هر مجموعه برل A

$$P(V \in A | T = t) = P(V \in A) \quad \text{داریم که}$$

قرار می‌دهیم $E = (V \in A)$ در این صورت

$$P(V \in A) = P(E) = E(I_E) \sim \theta \quad \text{مستقل از } \theta$$

$$\begin{aligned} P(V \in A | T = t) - P(V \in A) &= P(E | T = t) - P(E) = E[I_E | T = t] - E[I_E] \\ &= h(t) \quad \text{مستقل از } \theta \end{aligned}$$

$$h(T) = E[I_E | T] - E[I_E]$$

$$E_\theta[h(T)] = E\{E[I_E | T = t]\} - E\{E(I_E)\} \quad h(T) \text{ تابعی کراندار است و}$$

$$= E(I_E) - E(I_E) = 0 \stackrel{(۳)}{\Rightarrow} h(T) = 0 \quad a.e. \mathcal{P}$$

که در آن (۳) از کامل کراندار بودن T نتیجه می‌گردد. پس

$$P(V \in A | T = t) = P(V \in A) \quad a.e. \mathcal{P}$$

■

قضیه ۵-۳ (قضیه بهادر^۲): اگر یک آماره بسنده کامل باشد آنگاه آن آماره بسنده مینیمال است.

اثبات: فرض کنید U یک آماره بسنده کامل و T یک آماره بسنده مینیمال باشد، در این

صورت $T = K(U)$ و بنابراین

$$E_\theta(U) = E_\theta[E(U | T)] = E_\theta[h(T)] = E_\theta[h(K(U))]$$

$$\Rightarrow E_\theta[U - h(K(U))] = 0 \stackrel{(۴)}{\Rightarrow} U = h(K(U)) \quad a.e. \mathcal{P}$$

که در آن (۴) از کامل بودن U نتیجه می‌گردد.

^۱ Basu's Theorem

^۲ Bahadour's Theorem

■ U بسنده می‌نیمال است. $\Rightarrow a.e. P \Rightarrow U = h(T)$

توجه: توجه کنید که عکس قضیه بالا درست نیست یعنی اگر یک آماره بسنده مینیمال باشد آنگاه آن آماره ممکن است کامل نباشد. برای مثال توزیع‌های $N(\theta, \theta^2)$ و $U(\theta, \theta+1)$.

سوال: آیا آماره کاملی سراغ دارید که بسنده نباشد.

جواب: $P(U=c)=1$ در این صورت

$$E[h(U)] = E[h(c)] = 0 \Rightarrow P(h(c)=0) = 1$$

■ پس $U=c$ یک آماره کامل است ولی بسنده نیست.

توجه: با توجه به قضیه بهادر، اگر بخواهیم یک آماره کامل در خانواده P پیدا کنیم ابتدا در این خانواده طبق روش لهن شفه یک آماره بسنده مینیمال پیدا می‌کنیم. اگر این آماره کامل بود که مسئله حل شده است و اگر نبود آنگاه طبق قضیه بهادر می‌توان نتیجه گرفت که در این خانواده آماره بسنده کامل وجود ندارد.

بخش ۶: توابع زیان محدب^۱

فرض کنید در برآورد پارامتر مجهول $\theta \in \Theta$ برآوردگر $\delta \in D$ (کلاس کلیه برآوردگرها) را در نظر بگیریم. مهمترین خصوصیت یک برآوردگر خوب آن است که برآوردی را برای θ بدست دهد که تا حد ممکن به مقدار واقعی θ نزدیک باشد. میزان این نزدیکی را با تابع زیان اندازه‌گیری می‌کنیم که با نماد $L(\theta, \delta(X))$ نمایش داده می‌شود و

$$L : \Theta \times D \rightarrow [0, +\infty)$$

$$L(\theta, \delta) \geq 0 \quad \forall \theta, \delta$$

$$L(\theta, \delta) = 0 \quad \text{if} \quad \delta = \theta$$

چند تابع زیان مشهور عبارتند از :

$$(i) \quad L(\theta, T) = (\delta - \theta)^2 \quad \text{تابع زیان درجه دوم یا مربع خطا}$$

$$(ii) \quad L(\theta, T) = |\delta - \theta| \quad \text{تابع زیان قدر مطلق}$$

^۱ Convex Loss Function

$$(iii) \quad L(\theta, T) = b \{e^{a(\delta - \theta)} - a(\delta - \theta) - 1\} \quad \text{تابع زیان LINEX}$$

چون تابع زیان بستگی به $\delta(X)$ دارد پس مقدار آن ثابت نیست و یک متغیر تصادفی است. برای مقایسه دو برآوردگر، امید ریاضی تابع زیان این دو برآوردگر را که تابعی از θ می باشد با یکدیگر مقایسه می کنند که به آن تابع مخاطره^۱ گویند. پس

$$R(\theta, \delta) = E_{\theta}[L(\theta, \delta(X))]$$

در خصوص توابع زیان محدب قضایای مهمی وجود دارد که در این بخش ابتدا توابع محدب را معرفی و سپس این قضایا را بیان می کنیم.

تعریف ۱-۶: تابع حقیقی φ را روی فاصله $I = (a, b)$ که $-\infty \leq a < b \leq \infty$ محدب گوئیم هرگاه برای هر $a < x < y < b$ و هر $0 < \delta < 1$ داشته باشیم

$$\varphi(\gamma x + (1 - \gamma)y) \leq \gamma \varphi(x) + (1 - \gamma)\varphi(y)$$

و آن را اکیداً محدب گوئیم هرگاه نامساوی اکید باشد.

می توان نشان داد که توابع محدب، توابعی پیوسته روی (a, b) هستند.

قضیه ۱-۶: الف - اگر φ روی (a, b) مشتق پذیر باشد آنگاه شرط لازم و کافی برای آنکه φ محدب باشد آن است که

$$\varphi'(x) \leq \varphi'(y) \quad \forall \quad a < x < y < b$$

و اکیداً محدب است اگر نامساوی اکید باشد.

ب - اگر φ روی (a, b) دوبار مشتق پذیر باشد آنگاه شرط لازم و کافی برای آن که φ محدب باشد آن است که

$$\varphi''(x) \geq 0 \quad \forall \quad a < x < y < b$$

و اکیداً محدب است اگر نامساوی اکید باشد. ■

^۱ Risk Function

قضیه ۶-۲: اگر φ یک تابع محدب روی یک فاصله باز I باشد و X یک متغیر تصادفی با $P(X \in I) = 1$ و امید ریاضی متناهی باشد آنگاه $E[\varphi(X)] \leq \varphi(E(X))$ و اگر φ اکیداً محدب باشد آنگاه نامساوی اکید است، بجز اینکه X با احتمال یک روی I ثابت باشد.

اثبات: (تمرین)

نتیجه: اگر X یک متغیر تصادفی نااثبات مثبت با امید متناهی باشد آنگاه

$$\frac{1}{E(X)} \leq \left(\frac{1}{X}\right), \quad E(\log X) < \log(E(X))$$

برای توابع زیان محدب قضایای زیر را داریم.

قضیه ۶-۳ (قضیه رائوبلاکول): فرض کنید X یک متغیر تصادفی با توزیعی از خانواده $D = \{p_\theta : \theta \in \Theta\}$ باشد و T یک آماره بسنده برای D باشد. فرض کنید δ برآوردگری برای تابع پارامتری $g(\theta)$ باشد و تابع زیان $L(\theta, d)$ یک تابع اکیداً محدب بر حسب d باشد. اگر δ دارای امید

$$R(\theta, \delta) = E\{L(\theta, \delta(X))\} < \infty$$

و مخاطره متناهی

باشد و قرار دهیم $g(T) = E[\delta(X) | T]$ آنگاه

$$R(\theta, g) < R(\theta, \delta)$$

بجز اینکه با احتمال یک $g(T) = \delta(X)$ باشد.

اثبات: قرار می‌دهیم $\varphi(d) = L(\theta, d)$ و $\delta = \delta(X)$ و فرض کنید $P^{x|t}$ توزیع شرطی X به شرط

$T = t$ باشد. در این صورت طبق نامساوی جنسن داریم که

$$L(\theta, E[\delta(X) | T]) < E\{L(\theta, \delta(X)) | T\} \Rightarrow L(\theta, g(T)) < E\{L(\theta, \delta(X)) | T\}$$

بجز اینکه با احتمال یک $\delta(X) = g(T)$ باشد. حال اگر از طرفین امید ریاضی بگیریم آنگاه

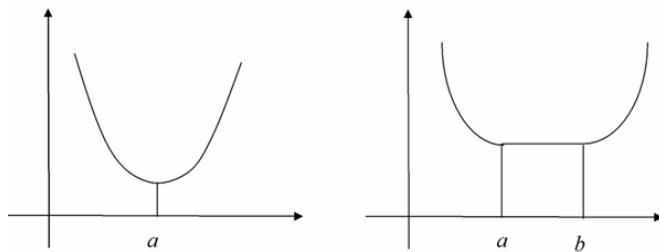
$$E[L(\theta, g(T))] < E\{E[L(\theta, \delta(X)) | T]\} = E[L(\theta, \delta(X))] \Rightarrow R(\theta, g) < R(\theta, \delta)$$

بجز اینکه با احتمال یک $\delta(X) = g(T)$ باشد.

توجه: قضیه رائوبلاکول می‌گوید که با شرطی کردن یک برآوردگر روی یک آماره بسنده، برآوردگری حاصل می‌گردد که دارای مخاطره کمتری نسبت به برآوردگر اولیه است و این یک توجیه مناسب برای

استفاده از آماره‌های بسنده در برآوردیابی می‌باشد. همچنین توجه کنید که از بسندگی تنها برای این استفاده شد که بگوئیم $g(T) = E[\delta(X)|T]$ یک برآوردگر است و به پارامتر θ بستگی ندارد. اگر در قضیه تابع زیان تنها محدب باشد آنگاه نامساوی به صورت $R(\theta, g) \leq R(\theta, \delta)$ تبدیل می‌شود.

لم ۱-۶: فرض کنید φ یک تابع محدب روی $(-\infty, +\infty)$ باشد که از پائین کراندار و همچنین غیر یکنوا است. در این صورت φ دارای یک مینیمم مقدار است و مجموعه مقادیر S که φ مقدار مینیمم خود را می‌گیرد یک فاصله بسته است و اگر φ اکیداً محدب باشد این مجموعه تنها یک نقطه است.



اثبات: چون φ محدب و غیر یکنواست

پس $\lim_{x \rightarrow \pm\infty} \varphi(x) = +\infty$ و چون φ پیوسته

است مقدار مینیمم خود را می‌گیرد. چون

φ محدب است، مجموعه S یک فاصله

است و چون φ پیوسته است، پس این مجموعه بسته است. ■

قضیه ۴-۶: فرض کنید ρ یک تابع محدب باشد که روی $(-\infty, +\infty)$ تعریف شده و X یک متغیر تصادفی باشد که $\varphi(a) = E[\rho(X - a)]$ برای بعضی مقادیر a متناهی باشد. اگر ρ یکنوا نباشد آنگاه $\varphi(a)$ مقدار مینیمم خود را روی یک فاصله بسته می‌گیرد. اگر ρ اکیداً محدب باشد مقدار مینیمم یکتا خواهد بود.

اثبات: چون $a \rightarrow \pm\infty$ as $X - a \xrightarrow{P} \pm\infty$, $\lim_{t \rightarrow \pm\infty} \rho(t) = +\infty$ پس $\lim_{a \rightarrow \pm\infty} \varphi(a) = +\infty$ پس φ

غیر یکنواست و چون ρ (اکیداً) محدب است پس φ نیز (اکیداً) یکنواست. حال که شرایط لم ۱-۶ برای φ برقرار است پس نتیجه از لم ۱-۶ حاصل می‌شود. ■

مثال ۱-۶: (تابع زیان درجه دوم): اگر $\rho(t) = t^2$ و $E(X^2) < \infty$ ، چون ρ تابعی اکیداً محدب است پس $\varphi(a) = E[(X - a)^2]$ مقدار مینیمم خود را در نقطه $a = \mu$ اختیار می‌کند.

مثال ۶-۲: (تابع زیان قدرمطلق خطا): اگر $\rho(t) = |t|$ و $E(|X|) < \infty$ ، چون ρ تابعی محدب است پس $\varphi(a) = E[|X - a|]$ مقدار مینیمم خود را در نقاط یک فاصله بسته بدست می‌آورد که همان میانه X است.

نتیجه ۶-۱: تحت شرایط قضیه ۶-۴ اگر ρ زوج باشد و X حول μ متقارن باشد آنگاه $\varphi(a)$ مینیمم خود را در $a = \mu$ بدست می‌آورد.

اثبات: بر طبق قضیه ۶-۴، φ مقدار مینیمم خود را بدست می‌آورد. اگر $\mu + c$ یک مقدار مینیمم باشد آنگاه $\mu - c$ نیز مینیمم می‌باشد و در نتیجه کلیه نقاط در فاصله $\mu - c$ تا $\mu + c$ مینیمم می‌باشند پس $a = \mu$ نیز مینیمم است. ■

(تعمیم توابع محدب به فضای k بُعدی $\mathbb{R}^k$ به عهده دانشجو)

(مبحث همگرایی‌ها بخش ۸ به عهده دانشجو)

فصل ۲

ناریبی

در برآورد نقطه‌ای می‌خواهیم برآوردگری را بدست آوریم که برای تمام مقادیر ممکن پارامتر، تابع مخاطره را مینیمم کند. اما در عمل با توجه به بزرگی کلاس برآوردگرها چنین امکانی وجود ندارد، یعنی در آمار بهترین به معنای مطلق وجود ندارد. بنابراین در آمار همیشه به دنبال بدست آوردن برآوردگرهایی هستیم که دارای خاصیت بهینه باشند. برای بدست آوردن برآوردگر بهینه دو روش زیر بیشتر مورد توجه است:

۱- محدود کردن^۱ کلاس برآوردگرها به یک کلاس کوچکتر

الف- کلاس برآوردگرهای ناریب ب- کلاس برآوردگرهای پایا

۲- برقراری یک نوع رابطه ترتیب^۲ بین برآوردگرها

الف- برآوردگرهای بیز ب- برآوردگرهای مینیماکس

بخش ۱: برآوردگرهای ناریب^۳

تعریف ۱-۱: برآوردگر $T(\tilde{X})$ را برای $\gamma(\theta)$ ناریب گویند هرگاه $E_{\theta}[T(\tilde{X})] = \gamma(\theta), \forall \theta \in \Theta$ و اگر برآوردگر $T(\tilde{X})$ برای $\gamma(\theta)$ ناریب نباشد مقدار $b(\theta) = E_{\theta}[T(\tilde{X})] - \gamma(\theta)$ را اریبی برآوردگر $T(\tilde{X})$ در برآورد $\gamma(\theta)$ گویند. ■

توجه کنید که در بعضی موارد ممکن است برآوردگر ناریب موجود نباشد.

مثال ۱-۱: فرض کنید $X \sim B(n, \theta)$ نشان دهید که $\gamma(\theta) = \frac{1}{\theta}$ دارای هیچ برآوردگر ناریبی نیست.

حل: فرض کنید $T(\tilde{X})$ یک برآوردگر ناریب $\frac{1}{\theta}$ باشد، در اینصورت

$$E(T(\tilde{X})) = \frac{1}{\theta} \Rightarrow \sum_{x=0}^n T(x) \binom{n}{x} \theta^x (1-\theta)^{n-x} = \frac{1}{\theta}$$

$$\Rightarrow T(\cdot)(1-\theta) + \sum_{x=1}^n T(x) \binom{n}{x} \theta^x (1-\theta)^{n-x} = \frac{1}{\theta}$$

^۱ Restriction

^۲ Ordering

^۳ Unbiased Estimators

حال اگر $\theta \rightarrow 0^+$ آنگاه سمت چپ معادله فوق به سمت $T(0)$ و سمت راست به سمت $+\infty$ میل می‌کند و نشان می‌دهد که $T(\tilde{X})$ وجود ندارد که در رابطه فوق صدق کند. ■

تعریف ۱-۲: الف - در خانواده توزیع‌های $\{p_\theta : \theta \in \Theta\}$ تابع حقیقی $\gamma(\theta)$ را برآورد

پذیر^۱ گویند اگر آماره $h(\tilde{X})$ وجود داشته باشد بطوری که $E_\theta[h(\tilde{X})] = \gamma(\theta), \forall \theta \in \Theta$

ب- آماره $\delta(\tilde{X})$ را برآوردگر ناریب با واریانس بطور یکنواخت مینیمم^۲ (UMVU) برای $\gamma(\theta)$

گویند اگر $E_\theta[h(\tilde{X})] = \gamma(\theta), \forall \theta \in \Theta$ و برای هر برآوردگر ناریب $h(\tilde{X})$ پارامتر $\gamma(\theta)$ داشته

$$V_\theta(\delta(\tilde{X})) \leq V_\theta(h(\tilde{X})) \quad \forall \theta \in \Theta$$

باشیم که

ج- آماره $\delta(\tilde{X})$ را برآوردگر ناریب با واریانس بطور موضعی مینیمم^۳ (LMVU) برای $\gamma(\theta)$ در

$\theta = \theta_0$ گویند اگر این برآوردگر ناریب بوده و برای هر برآوردگر ناریب $h(\tilde{X})$ پارامتر $\gamma(\theta)$

$$V_{\theta_0}(\delta(\tilde{X})) \leq V_{\theta_0}(h(\tilde{X}))$$

داشته باشیم که

یک روش برای بدست آوردن برآوردگرهای UMVU یا LMVU یافتن کلاس کلیه برآوردگرهای

ناریب $\gamma(\theta)$ یعنی Δ_γ است و سپس در این کلاس برآوردگرهای UMVU یا LMVU را بدست

می‌آوریم. برای این منظور از لم زیر استفاده می‌کنیم.

لم ۱-۱: اگر δ یک برآوردگر ناریب $\gamma(\theta)$ باشد آنگاه کلیه برآوردگرهای ناریب $\gamma(\theta)$ بصورت

$\delta - U$ می‌باشد که در آن U هر برآوردگر ناریب صفر است (یعنی $\forall \theta \in \Theta$)

$$E_\theta(U) = 0 \text{ و بعلاوه } Var(\delta) = E[(\delta - U)^2] - [\gamma(\theta)]^2$$

اثبات : واضح است. ■

قرار می‌دهیم

$$\Delta_\gamma = \{\delta | E_\theta(\delta) = \gamma(\theta), \forall \theta \in \Theta\} \text{ و } \Delta = \{U | E_\theta(U) = 0, \forall \theta \in \Theta\}$$

^۱ Estimable

^۲ Uniformly Minimum Variance Unbiased

^۳ Locally Minimum Variance Unbiased

برای استفاده از لم فوق در حل یک مسئله، ابتدا کلاس کلیه برآوردگرهای ناریب $\gamma(\theta)$ را به دست می‌آوریم و واریانس آنها را در یک نقطه ثابت θ مینیم می‌کنیم. اگر این مقدار مینیم به مقدار θ بستگی نداشته باشد (و در نتیجه حل مسئله به مقدار مشخص θ بستگی نداشته باشد) آنگاه برآوردگر حاصل UMVU پارامتر مورد نظر است. در غیر این صورت اگر برآوردگر حاصل به θ بستگی داشته باشد (و در نتیجه کمترین واریانس را در بعضی از نقاط θ بدست آورد) آنگاه این برآوردگر LMVU نامیده می‌شود و در این حالت UMVU پارامتر مورد نظر وجود ندارد.

مثال ۱-۲: فرض کنید X یک متغیر تصادفی گسسته با تابع احتمال زیر باشد

$$p_{\theta}(x) = \begin{cases} \theta & x = 1 \\ \theta^{x-2}(1-\theta)^2 & x = 2, 3, \dots \end{cases} \quad \langle \theta \rangle$$

بررسی کنید که آیا پارامترهای الف) $\gamma_1(\theta) = \theta$ ب) $\gamma_2(\theta) = (1-\theta)^2$ دارای برآوردگر UMVU هستند؟

حل: در یکی از مثال‌های فصل قبل نشان دادیم که

$$E_{\theta}[U(X)] = 0 \Rightarrow U(x) = (x-2)(-h(1)) = (x-2)a \quad x = 1, 2, \dots$$

و نتیجه گرفتیم که X یک آماره کامل نیست اما کامل کراندار است. حال قرار می‌دهیم:

$$\delta_1(X) = I_{\{1\}}(X) \Rightarrow E_{\theta}[\delta_1(X)] = \theta$$

$$\delta_2(X) = I_{\{2\}}(X) \Rightarrow E_{\theta}[\delta_2(X)] = (1-\theta)^2$$

پس δ_1 و δ_2 به ترتیب برآوردگرهای ناریب برای $\gamma_1(\theta)$ و $\gamma_2(\theta)$ می‌باشند. در هریک از حالات فوق، فرم کلی برآوردگرهای ناریب به فرم زیر می‌باشد:

$$\delta_i^*(X) = \delta_i(X) - U(X) \quad i = 1, 2, \dots$$

و δ_i^* موقعی کمترین واریانس خود را در θ بدست می‌آورد که

$$V_{\theta}(\delta_i^*) = E_{\theta}[(\delta_i - U)^2] - [\gamma_i(\theta)]^2$$

و یا معادلاً $E_{\theta}[(\delta_i - U)^2]$ مینیم شود. اما

$$E_{\theta}[(\delta_i(X) - U(X))^2] = \sum_{x=1}^{+\infty} (\delta_i(x) - U(x))^2 p_{\theta}(x)$$

$$= \sum_{x=1}^{+\infty} (\delta_i(x) - a(x-1))^\gamma p_\theta(x)$$

بنابراین مقداری از a را می‌خواهیم پیدا کنیم که مجموع بالا را مینیمم کند.

$$i = 1 \Rightarrow g_1(a) = E_\theta[(T_1(x) - h(x))^\gamma] = (1-a)^\gamma \theta + \sum_{x=2}^{+\infty} a^\gamma (x-1)^\gamma \theta^{x-1} (1-\theta)^\gamma \quad \text{الف}$$

$$= (1+a)^\gamma \theta + \sum_{y=1}^{+\infty} a^\gamma y^\gamma \theta^y (1-\theta)^\gamma$$

$$g'_1(a) = \gamma(1+a)\theta + \gamma a(1-\theta)^\gamma \sum_{y=1}^{+\infty} y^\gamma \theta^y = \gamma a \left(\theta + (1-\theta)^\gamma \sum_{y=1}^{+\infty} y^\gamma \theta^y \right) + \gamma \theta.$$

$$g''_1(a) = \text{constant} > 0 \Rightarrow g'_1(a) = 0 \Rightarrow a_1^* = \frac{-\theta}{\theta + (1-\theta)^\gamma \sum_{y=1}^{+\infty} y^\gamma \theta^y}$$

بنابراین $\delta_1^*(X) = \delta_1(X) - a_1^*(X-1)$ به دلیل آنکه a_1^* به مقدار θ بستگی دارد یک برآوردگر LMVU برای θ است.

ب-

$$\begin{aligned} i = 2 \Rightarrow g_2(\theta) &= E_\theta[(\delta_2(X) - U(X))^\gamma] \\ &= a^\gamma \theta + (1-\theta)^\gamma + \sum_{x=3}^{+\infty} a^\gamma (x-2)^\gamma \theta^{x-2} (1-\theta)^\gamma \\ &= a^\gamma \left[\theta + \sum_{y=1}^{+\infty} y^\gamma \theta^y (1-\theta)^\gamma \right] + (1-\theta)^\gamma \end{aligned}$$

$$g'_2(a) = \gamma a \left[\theta + (1-\theta)^\gamma \sum_{y=1}^{+\infty} y^\gamma \theta^y \right], \quad g''_2(a) = \text{constant} > 0.$$

$$\Rightarrow g'_2(a) = 0 \Rightarrow a_2^* = 0.$$

چون a_2^* بستگی به θ ندارد بنابراین برآوردگر

$$\delta_2^*(X) = \delta_2(X) - a_2^*(X-2) = \delta_2(X)$$

برآوردگری نااریب برای $(1-\theta)^2$ است که واریانس را نه تنها برای θ بلکه برای کلیه $\theta \in \Theta$ مینیمم می‌کند پس $T_{\gamma}^*(X)$ برآوردگر UMVU برای $(1-\theta)^2$ است. ■

با توجه به مثال بالا این سوال مطرح می‌شود که چه برآوردگرهایی برای پارامتر $\gamma(\theta)$ برآوردگرهای UMVU هستند و همچنین چه توابع برآوردپذیر $\gamma(\theta)$ دارای برآوردگرهای UMVU هستند. جواب به این سوال از طریق قضیه زیر داده می‌شود. در این قضیه $\Delta = \{\delta | E(\delta^2) < \infty\}$ است.

قضیه ۱-۱ (قضیه لهن-شفه): فرض کنید X یک متغیر تصادفی با توزیعی از خانواده $\mathcal{P} = \{P_{\theta} : \theta \in \Theta\}$ باشد و $\delta \in \Delta$ و Δ کلاس کلیه برآوردگرهای نااریب صفر باشد. در این صورت شرط لازم و کافی برای آنکه δ یک برآوردگر UMVU برای امید خود $\gamma(\theta) = E_{\theta}[\delta(X)]$ باشد آن است که :

$$Cov(\delta(X), U(X)) = 0 \quad \forall U \in \Delta, \forall \theta \in \Theta$$

توجه: اگر به جای θ مقدار θ را در قضیه قرار دهیم، آنگاه T یک برآوردگر LMVU خواهد بود.

$$Cov(\delta(X), U(X)) = 0 \Leftrightarrow E[\delta U] = 0 \quad \text{همچنین چون } E_{\theta}(U) = 0 \text{ پس}$$

اثبات: (شرط کافی): فرض کنید

$$Cov(\delta(X), U(X)) = 0 \quad \forall U \in \Delta, \forall \theta \in \Theta$$

و فرض کنید δ^* برآوردگر نااریب دیگری برای $\gamma(\theta)$ باشد. یعنی

$$E_{\theta}[\delta^*(X)] = E_{\theta}[\delta(X)] = \gamma(\theta) \Rightarrow E_{\theta}[\delta^*(X) - \delta(X)] = 0 \quad \forall \theta \in \Theta$$

بنابراین طبق فرض مساله داریم که

$$Cov(\delta(X), \delta^*(X) - \delta(X)) = 0 \quad \forall \theta \in \Theta \Rightarrow Cov(\delta(X), \delta^*(X)) - V(\delta(X)) = 0$$

$$\Rightarrow Var_{\theta}(\delta(X)) = Cov_{\theta}(\delta(X), \delta^*(X)) \stackrel{(1)}{\leq} \sqrt{Var_{\theta}(\delta(X)) Var_{\theta}(\delta^*(X))}$$

$$\Rightarrow \sqrt{Var_{\theta}(\delta(X))} \leq \sqrt{Var_{\theta}(\delta^*(X))} \Rightarrow V_{\theta}(\delta(X)) \leq V_{\theta}(\delta^*(X)) \quad \forall \theta \in \Theta$$

که در آن نامساوی (۱) از نامساوی شوارتز بدست می‌آید. پس $\delta(X)$ برآوردگر UMVU پارامتر $\gamma(\theta) = E_{\theta}(\delta)$ است.

(شرط لازم): فرض کنید $\delta(X)$ برآوردگر UMVU برای $\gamma(\theta) = E_{\theta}(\delta)$ باشد و

$$Cov_{\theta}(\delta(X), U(X)) = 0 \quad E_{\theta}[U(X)] = 0 \quad \forall \theta \in \Theta$$

می‌خواهیم نشان دهیم که

برای یک مقدار ثابت λ قرار می‌دهیم

$$\delta^*(X) = \delta(X) + \lambda U(X)$$

چون $\delta(X)$ برآوردگر UMVU و $\delta^*(X)$ برآوردگر ناریب $\gamma(\theta)$ می‌باشند پس

$$V_{\theta}(\delta(X)) \leq V_{\theta}(\delta^*(X)) = V_{\theta}(\delta(X) + \lambda U(X))$$

$$= V_{\theta}(\delta(X)) + \lambda^2 V_{\theta}(U(X)) + 2\lambda Cov_{\theta}(\delta(X), U(X))$$

$$\Rightarrow \lambda^2 V_{\theta}(U(X)) + 2\lambda Cov_{\theta}(\delta(X), U(X)) \geq 0 \quad \forall \theta \in \Theta, \quad \forall \lambda \quad (2)$$

معادله $\lambda^2 V_{\theta}(U(X)) + 2\lambda Cov_{\theta}(\delta(X), U(X)) = 0$ یک معادله درجه دوم بر حسب λ است

که دارای دو ریشه $\lambda_1 = 0$ و $\lambda_2 = \frac{-2Cov_{\theta}(\delta(X), U(X))}{V_{\theta}(U(X))}$ است و در نتیجه بین این دو ریشه

علامت معادله منفی خواهد بود. بنابراین اگر نامعادله (۲) بخواهد برای تمام مقادیر λ برقرار باشد

بایستی $\lambda_1 = \lambda_2 = 0$ باشد یعنی بایستی $Cov_{\theta}(\delta(X), U(X)) = 0$ باشد. ■

مثال ۱-۳: در مثال ۱-۲ کلاس کلیه برآوردگرهای UMVU و همچنین کلاس کلیه پارامترهای $\gamma(\theta)$

برآوردپذیر دارای برآوردگر UMVU را بدست آورید و نشان دهید که پارامتر θ دارای هیچ

برآوردگر UMVU نیست.

حل: کلاس کلیه برآوردگرهای ناریب صفر عبارت است از

$$U(X) = a(X - 2) \quad X = 1, 2, \dots, \quad a = -U(1)$$

$$\gamma(\theta) = E_{\theta}(\delta) \text{ برای UMVU برآوردگر } \delta(X) \Leftrightarrow Cov_{\theta}(\delta(X), U(X)) = 0$$

$$\Leftrightarrow E_{\theta}[\delta(X)a(X-2)] = 0 \Leftrightarrow E_{\theta}[\delta(X)(X-2)] = 0$$

$$\Leftrightarrow \delta(X)(X-2) = \{-\delta(1)(1-2)\}(X-2) = \delta(1)(X-2)$$

$$\Rightarrow \delta(x) = \begin{cases} \delta(1) & x = 1, 3, 4, \dots \\ b & x = 2 \end{cases} \Rightarrow \delta(x) = \begin{cases} a & x = 1, 3, 4, \dots \\ b & x = 2 \end{cases}$$

که در آن b یک ثابت دلخواه است. بنابراین فرم کلی برآوردگرهای UMVU بصورت $\delta(X)$

است.

$$\begin{aligned}
 \gamma(\theta) &= E_{\theta}[\delta(X)] = a\theta + b(1-\theta)^2 + \sum_{x=3}^{+\infty} a\theta^{x-2}(1-\theta)^2 = a\theta + b(1-\theta)^2 + a\theta(1-\theta) \\
 &= a(2\theta - \theta^2) + b(1-\theta)^2 = -a(1-\theta)^2 + a + b(1-\theta)^2 \\
 &= a + (b-a)(1-\theta)^2 = a + c(1-\theta)^2
 \end{aligned}$$

بنابراین تنها پارامترهای بفرم فوق دارای برآوردگر UMVU هستند، ولی θ دارای برآوردگر UMVU نیست زیرا

$$\theta = a + c(1-\theta)^2 \Rightarrow c(1-2\theta + \theta^2) - \theta + a = 0 \Rightarrow c\theta^2 - (2c+1)\theta + a = 0$$

$$\forall \theta \in (0,1) \Rightarrow \begin{cases} a = 0 \\ c = 0 \\ 2c+1=0 \Rightarrow c = -\frac{1}{2} \end{cases}$$

■

که به تناقض بر می‌خوریم.

نکاتی در مورد قضیه لهن - شف:

(۱) با توجه به مثال قبل، بوسیله این قضیه می‌توان کلیه برآوردگرهای UMVU و همچنین کلیه پارامترهای برآوردپذیر دارای برآوردگر UMVU را بدست آورد و تنها مشکل آن بدست آوردن برآوردگرهای نااریب صفر است.

(۲) تعبیر هندسی: اگر $\Delta = \{\delta | E(\delta^2) < \infty\}$ آنگاه Δ یک فضای برداری است زیرا

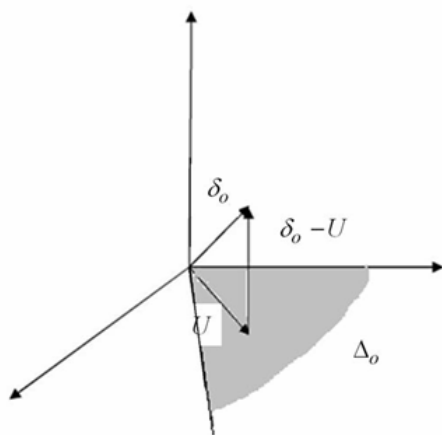
$$\forall a, b \in \mathbb{R}, \delta_1, \delta_2 \in \Delta \Rightarrow a\delta_1 + b\delta_2 \in \Delta$$

و Δ_0 نیز یک زیر فضای Δ است. اگر قرار دهیم $\langle \delta_1, \delta_2 \rangle = Cov(\delta_1, \delta_2)$ آنگاه یک ضرب

داخلی روی Δ خواهیم داشت و با توجه به ضرب داخلی

$$\delta_1, \delta_2 \text{ فاصله} = d(\delta_1, \delta_2) \text{ فوق داریم که}$$

$$= \sqrt{\langle \delta_1 - \delta_2, \delta_1 - \delta_2 \rangle} = \sqrt{Var(\delta_1 - \delta_2)}$$



اگر δ برآوردگر UMVU باشد آنگاه به ازای هر $\delta' \in \Delta$

$$\delta = \delta_0 - U \text{ حال چون } \sqrt{V}(\delta) \leq \sqrt{V}(\delta') \text{ داریم که}$$

می‌خواهیم U را طوری انتخاب کنیم که $V(\delta) = V(\delta_0 - U) = d(\delta_0, U)$ مینیمم شود. اما برای مینیمم کردن $d(\delta_0, U)$ بایستی U را تصویر δ_0 در زیر فضای Δ_0 بگیریم. بنابراین باید $\delta_0 - U$ بر زیر فضای Δ_0 عمود باشد. یعنی

$$\blacksquare \quad \text{Cov}(\delta_0 - U, U) = 0 \Leftrightarrow \text{Cov}(\delta, U) = 0.$$

(۳) از قضیه لهن - شفه برای اثبات عدم UMVU بودن یک برآوردگر می‌توان استفاده کرد.

مثال ۱-۴: فرض کنید $X \sim U(\theta, \theta+1)$ می‌دانیم که $T(X) = X$ یک آماره بسنده مینیمال برای θ است. نشان دهید که $T(X)$ برای θ کامل نیست و $S(X) = X - \frac{1}{2}$ یک برآوردگر نااریب θ است که UMVUE آن نیست.

حل:

$$E[h(X)] = 0 \Rightarrow \int_{\theta}^{\theta+1} h(x) dx = 0 \Rightarrow \frac{d}{d\theta} \int_{\theta}^{\theta+1} h(x) dx = 0$$

$$\Rightarrow h(\theta+1) - h(\theta) = 0 \Rightarrow h(\theta+1) = h(\theta) \quad \forall \theta \in \mathfrak{R}$$

بنابراین توابع متناوب با دوره تناوب ۱ دارای امید صفر هستند. برای مثال اگر $h(x) = \sin(2\pi x)$

آنگاه $E[\sin(2\pi x)] = 0$ اما $P(\sin(2\pi x) = 0) = 0 \neq 1$.

$$E(S(X)) = E\left(X - \frac{1}{2}\right) = \int_{\theta}^{\theta+1} \left(x - \frac{1}{2}\right) dx = \frac{1}{2}x^2 - \frac{1}{2}x \Big|_{\theta}^{\theta+1} = \frac{1}{2}(2\theta) = \theta$$

$$\text{Cov}_{\theta}(S(X), h(X)) = E_{\theta}(S(X)h(X)) = E_{\theta}(X \sin(2\pi X))$$

$$\begin{aligned} &= \int_{\theta}^{\theta+1} x \sin(2\pi x) dx = \left[-\frac{1}{2\pi} x \cos(2\pi x) + \frac{1}{4\pi} \sin(2\pi x) \right]_{\theta}^{\theta+1} \\ &= -\frac{1}{2\pi}(\theta+1) \cos(2\pi\theta) + \frac{1}{4\pi} \sin(2\pi\theta) + \frac{1}{2\pi} \theta \cos(2\pi\theta) - \frac{1}{4\pi} \sin(2\pi\theta) \\ &= -\frac{1}{2\pi} \cos(2\pi\theta) \neq 0 \end{aligned}$$

پس $S(X) = X - \frac{1}{2}$ برآوردگر UMVU پارامتر θ نیست. $\blacksquare$

(۴) توسط قضیه لهن - شفه می‌توان یکتایی برآوردگرهای UMVU را به صورت زیر ثابت کرد.

لم ۱-۲: برای هر دو متغیر تصادفی X و Y با گشتاورهای دوم متنهای داریم که $V(X) \geq [Cov(X, Y)]^2$ و تساوی برقرار است اگر و فقط اگر ثابت‌های a و b موجود باشند بطوری که $P(X = a + bY) = 1$.

قضیه ۱-۲: اگر $\delta(X)$ یک برآوردگر UMVU برای پارامتر $\gamma(\theta)$ باشد آنگاه $\delta(X)$ برآوردگر یکتای UMVU برای $\gamma(\theta)$ است.

اثبات: فرض کنید δ و δ' دو برآوردگر UMVU پارامتر $\gamma(\theta)$ باشند. پس

$$E_{\theta}(\delta) = E_{\theta}(\delta') = \gamma(\theta), \quad V_{\theta}(\delta) = V_{\theta}(\delta') \quad \forall \theta \in \Theta \quad (3)$$

در نتیجه $E_{\theta}[\delta - \delta'] = 0$ پس $\delta - \delta'$ برآوردگر نااریب صفر است. پس طبق قضیه لهن - شفه

$$0 = \text{cov}(\delta, \delta - \delta') = V(\delta) - \text{Cov}(\delta - \delta') \Rightarrow \text{Cov}(\delta - \delta') = V(\delta)$$

$$\Rightarrow [\text{Cov}(\delta, \delta')]^2 = [V(\delta)]^2 = V_{\theta}(\delta)V_{\theta}(\delta')$$

یعنی در نامساوی کواریانس حالت تساوی برقرار است پس طبق لم ۱-۳ داریم که $P_{\theta}(\delta = a\delta' + b) = 1$ و در نتیجه طبق (۳) داریم که $P_{\theta}(\delta = \delta') = 1$ یعنی برآوردگر UMVU یکتا است. ■

توجه: برآوردگر $\delta = c$ برآوردگری UMVU برای امید خود است، زیرا دارای واریانس صفر است. حال اگر برآوردهای ثابت را کنار بگذاریم، برای یک تابع برآوردپذیر $\gamma(\theta)$ سه حالت زیر پیش می‌آید:

حالت اول: هیچ تابع غیر ثابتی از θ برآوردگر UMVU ندارد.

مثال ۱-۵: فرض کنید متغیر تصادفی X دارای تابع احتمال زیر باشد، نشان دهید که هیچ تابع غیر ثابتی از θ برآوردگر UMVU ندارد

x	$\theta - 1$	θ	$\theta + 1$
$p_{\theta}(x)$	$\frac{1}{3}$	$\frac{1}{3}$	$\frac{1}{3}$

ابتدا کلاس برآوردهای نااریب صفر یعنی Δ را بدست می‌آوریم.

^۱ Unique

$$E_{\theta}[U(X)] = 0 \Rightarrow \frac{1}{3}[U(\theta-1) + U(\theta) + U(\theta+1)] = 0.$$

$$\Rightarrow U(\theta-1) + U(\theta) + U(\theta+1) = 0.$$

$$\left. \begin{aligned} \theta = 0 &\Rightarrow U(-1) + U(0) + U(1) = 0 \\ \theta = 1 &\Rightarrow U(0) + U(1) + U(2) = 0 \\ \theta = 2 &\Rightarrow U(1) + U(2) + U(3) = 0 \end{aligned} \right\} \Rightarrow \begin{cases} U(k) = U(k+3) \\ U(k) + U(k+1) + U(k+2) = 0 \end{cases} \quad k \in \mathbb{Z}$$

$$\Rightarrow \Delta = \{U \mid U(3k) = a, U(3k+1) = b, U(3k+2) = -(a+b), k \in \mathbb{Z}, a, b \in \mathbb{R}\}$$

$$UMVU \text{ برآوردگر } \delta \Leftrightarrow \delta U \in \Delta \Rightarrow \begin{cases} \delta(3k)U(3k) = a \\ \delta(3k+1)U(3k+1) = b \\ \delta(3k+2)U(3k+2) = -(a+b) \end{cases}$$

$$\begin{aligned} a \neq 0, b \neq 0, &\begin{cases} \delta(3k) = 1 \\ \delta(3k+1) = 1 \Rightarrow \delta(k) = 1 \quad \forall k \in \mathbb{Z} \\ \Rightarrow \delta(3k+2) = 1 \end{cases} \\ a+b \neq 0. & \end{aligned}$$

پس δ یک برآوردگر ثابت است که نمی‌تواند برای $\gamma(\theta)$ ناریب باشد پس UMVU نیز نیست. ■

حالت دوم: تنها بعضی از توابع برآوردپذیر غیر ثابت دارای برآوردگر UMVU هستند. (مثال ۳۱-)

حالت سوم: هر پارامتر برآوردپذیر $\gamma(\theta)$ دارای برآوردگر UMVU است.

یک روش برای به دست آوردن UMVU در این حالت توسط قضیه راثوبلاکول است. طبق قضیه

راثوبلاکول اگر T یک آماره بسنده برای θ باشد و $\delta(\tilde{X})$ یک برآوردگر ناریب برای $\gamma(\theta)$ باشد

آنگاه $g(T) = E[\delta(\tilde{X})|T]$ یک آماره است و $E_{\theta}[g(T)] = E_{\theta}[\delta(\tilde{X})] = \gamma(\theta)$ و

$$P(g(T) = \delta(\tilde{X})) = 1 \text{ اگر و فقط اگر } V_{\theta}[g(T)] < V_{\theta}[\delta(\tilde{X})]$$

بنابراین هر برآوردگر ناریب $\gamma(\theta)$ را می‌توان با امید شرطی نسبت به یک آماره بسنده بهبود بخشید.

حال این سوال مطرح است که آیا به ازای هر برآوردگر ناریب اولیه $\delta(\tilde{X})$ ، برآوردگر حاصل از امید

شرطی یکتا است؟

قضیه ۳-۱: فرض کنید X دارای توزیعی از خانواده $\mathcal{P} = \{P_\theta : \theta \in \Theta\}$ باشد و T یک آماره بسنده کامل برای $\mathcal{P}$ باشد. در این صورت هر تابع برآوردپذیر $\gamma(\theta)$ دارای یک و تنها یک برآوردگر نااریب است که تابعی از T می باشد.

اثبات: وجود این برآوردگر از قضیه راثوبلاکول بدست می آید و اگر $g_1(T)$ و $g_2(T)$ دو برآوردگر نااریب $\gamma(\theta)$ باشند آنگاه

$$\blacksquare \quad E[g_1(T) - g_2(T)] = 0 \stackrel{\text{completeness}}{\Rightarrow} g_1(T) = g_2(T) \quad \text{با احتمال یک}$$

از قضیه راثوبلاکول و قضیه قبل نتیجه زیر را می توان بدست آورد.

قضیه ۴-۱: فرض کنید X دارای توزیعی از خانواده $\mathcal{P} = \{P_\theta : \theta \in \Theta\}$ باشد و T یک آماره بسنده و کامل برای خانواده $\mathcal{P}$ باشد.

الف) برای هر تابع برآوردپذیر $\gamma(\theta)$ یک برآوردگر نااریب وجود دارد که به طور یکنواخت مخاطره را برای هر تابع زیان محدب $L(\theta, d)$ مینیمم می کند. این برآوردگر در حالت خاص UMVU است.
ب) برآوردگر UMVU حالت قبل همان برآوردگر نااریب $\gamma(\theta)$ است که تابعی از T می باشد. این برآوردگر نااریب با مخاطره مینیمم است به شرط آنکه مخاطره آن متناهی و تابع زیان اکیدا محدب باشد. $\blacksquare$

توجه: اگر تابع زیان به غیر از تابع زیان درجه دوم باشد، برآوردگر حاصل را UMRU گویند.

قضایای قبل نه تنها وجود برآوردگر UMVU را تضمین می کند بلکه روش هایی را برای بدست آوردن برآوردگرهای UMVU ارائه می دهد. با داشتن آماره بسنده کامل T در برآورد پارامتر $\gamma(\theta)$ این روش ها عبارتند از:

الف) ابتدا یک برآوردگر نااریب $\delta(\underline{X})$ برای $\gamma(\theta)$ تعیین می کنیم در این صورت $g(T) = E[\delta(\underline{X})|T]$ برآوردگر UMVU پارامتر $\gamma(\theta)$ است.

مثال ۶-۱: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $U(\cdot, \theta)$ باشند و $\gamma(\theta) = \frac{\theta}{2}$ در این

صورت $T = X_{(n)}$ یک آماره بسنده و کامل برای θ است و $E[X_1] = \frac{\theta}{2}$ پس

$g(T) = E[X_1 | T]$ برآوردگر UMVU پارامتر $\gamma(\theta)$ است. اما اگر $X_{(n)} = T = t$ باشد آنگاه

$X_1 = t$ یا $X_1 < t$ خواهد بود. پس اگر $X_{(n)} = t$ باشد آنگاه با احتمال $\frac{1}{n}$ ، $X_1 = t$ و با احتمال

$X_1, \frac{n-1}{n}$ دارای توزیع یکنواخت در فاصله $(0, t)$ می‌باشد. بنابراین

$$g(t) = E(X_1 | X_{(n)} = t) = t \left(\frac{1}{n} \right) + \frac{t}{2} \left(\frac{n-1}{n} \right) = \frac{n+1}{n} \cdot \frac{t}{2}$$

پس $g(T) = \frac{n+1}{n} \cdot \frac{T}{2}$ برآوردگر UMVU پارامتر $\gamma(\theta) = \frac{\theta}{2}$ است. ■

(ب) چون برآوردگر UMVU یکتا است، پس برآوردگر UMVU تابع برآوردپذیر $\gamma(\theta)$ با حل

$$E_\theta[g(T)] = \gamma(\theta) \quad \forall \theta \in \Theta$$
 معادله زیر بدست می‌آید.

$$\frac{X_{(n)}}{\theta} \sim Be(n, 1) \Rightarrow E\left(\frac{X_{(n)}}{\theta}\right) = \frac{n}{n+1}$$

مثال ۷-۱: در مثال قبل داریم که

$$\Rightarrow E\left(\frac{n+1}{n} \cdot \frac{X_{(n)}}{2}\right) = \frac{\theta}{2} \Rightarrow \frac{n+1}{n} \cdot \frac{X_{(n)}}{2} \text{ is UMVUE of } \frac{\theta}{2} \quad \blacksquare$$

مثال ۸-۱: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $B(1, \theta)$ باشند. برآوردگر UMVU

پارامتر $\gamma(\theta) = \theta(1-\theta)$ را بدست آورید.

حل: می‌دانیم که $T = \sum_{i=1}^n X_i \sim B(n, \theta)$ آماره بسنده و کامل برای θ است.

$$\begin{aligned} E[g(T)] &= \theta(1-\theta) \Rightarrow \sum_{t=0}^n g(t) \binom{n}{t} \theta^t (1-\theta)^{n-t} = \theta(1-\theta) \\ &\Rightarrow \sum_{t=0}^n g(t) \binom{n}{t} \left(\frac{\theta}{1-\theta}\right)^t = \theta(1-\theta)^{1-n} \end{aligned}$$

قرار می‌دهیم $\beta = \frac{\theta}{1-\theta} > 0$ در این صورت $\theta = \frac{\beta}{1+\beta}$ و $1-\theta = \frac{1}{1+\beta}$ بنابراین

$$\sum_{t=0}^n g(t) \binom{n}{t} \beta^t = \frac{\beta}{1+\beta} \left(\frac{1}{1+\beta} \right)^{1-n} = \beta (1+\beta)^{n-2} = \sum_{u=0}^{n-2} \binom{n-2}{u} \beta^{u+1}$$

$$= \sum_{t=1}^{n-1} \binom{n-2}{t-1} \beta^t \quad \forall \beta > 0.$$

$$\Rightarrow g(t) = \begin{cases} 0 & t = 0, n \\ \frac{\binom{n-2}{t-1}}{\binom{n}{t}} & t \neq 0, n \end{cases} = \begin{cases} 0 & t = 0, n \\ \frac{t(n-t)}{n(n-1)} & t \neq 0, n \end{cases} \Rightarrow g(t) = \frac{t(n-t)}{n(n-1)} \quad t = 0, \dots, n$$

پس $g(T) = \frac{T(n-T)}{n(n-1)}$ برآوردگر UMVU پارامتر $\theta(1-\theta)$ است. ■

بخش‌های ۲-۲، ۲-۳ و ۲-۴ به عنوان سمینار برای دانشجویان

بخش ۵: نامساوی اطلاع^۱ (نامساوی کرامر - رائو)

فرض کنید δ یک برآوردگر $g(\theta)$ باشد که $E(\delta) = g(\theta)$ و $\psi(x, \theta)$ تابعی از x و θ با گشتاور دوم متناهی باشد، در این صورت طبق نامساوی کواریانس داریم که

$$Var(\delta) \geq \frac{[Cov(\delta, \psi)]^2}{Var(\psi)} \quad (1)$$

برای یافتن یک کران پایین برای $V(\delta)$ می‌خواهیم سمت راست نامساوی فوق به δ بستگی نداشته باشد و حداقل به $E(\delta) = g(\theta)$ بستگی داشته باشد.

قضیه ۵-۱ (بلیت^۲ ۱۹۷۴): شرط لازم و کافی برای آن که $Cov(\delta, \psi)$ بستگی به δ تنها از

طریق $g(\theta)$ داشته باشد آن است که $Cov(U, \psi) = 0 \quad \forall U \in \Delta$.

که در آن $\Delta = \{U : E_\theta(U) = 0, E_\theta(U^2) < \infty \quad \forall \theta \in \Theta\}$

^۱ Information Inequality

^۲ Blyth

اثبات: (کفایت) فرض کنید δ_1 و δ_2 هر دو برای $g(\theta)$ نااریب باشند در این صورت $\delta_1 - \delta_2 \in \Delta$. پس $Cov(\delta_1 - \delta_2, \psi) = 0$ و در نتیجه $Cov(\delta_1, \psi) = Cov(\delta_2, \psi)$. یعنی از $E(\delta_1) = E(\delta_2)$ نتیجه گرفتیم که $Cov(\delta_1, \psi) = Cov(\delta_2, \psi)$ پس $Cov(\delta, \psi)$ تنها از طریق $E(\delta)$ به δ وابسته است.

(لزوم): فرض کنید $Cov(\delta, \psi)$ تنها از طریق $g(\theta) = E(\delta)$ به δ وابسته باشد. چون $E(\delta + U) = g(\theta)$ پس

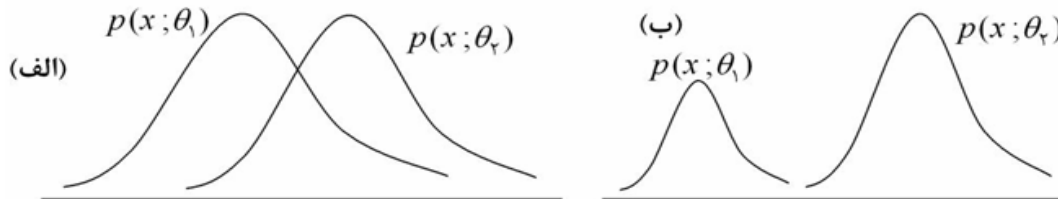
$$Cov(\delta, \psi) = Cov(\delta + U, \psi) \quad \forall \theta \in \Theta \quad \forall U \in \Delta.$$

$$\Rightarrow Cov(\delta, \psi) = Cov(\delta, \psi) + Cov(U, \psi)$$

$$\Rightarrow Cov(U, \psi) = 0 \quad \forall \theta \in \Theta \quad \forall U \in \Delta.$$

اطلاع فشر:

نمودارهای زیر را در نظر بگیرید.



سؤال: مشاهده X در کدام حالت اطلاع بیشتری در مورد پارامتر θ به ما می‌دهد؟

جواب: بدن شک در حالت (ب) مشاهده X به ما اطلاع بیشتری می‌دهد.

پس هر چه تابع چگالی X از یک مقدار θ به مقدار دیگر θ بیشتر تغییر کند، مشاهده X به ما در مورد مقدار پارامتر θ بیشتر اطلاع می‌دهد. حال توجه کنید که

$$\text{مقدار تغییر نسبی در چگالی در نقطه } X \text{ به ازای تغییر پارامتر به اندازه } \Delta = \left[\frac{P(x; \theta + \Delta) - p(x; \theta)}{\Delta} \cdot \frac{1}{p(x; \theta)} \right]$$

$$\begin{aligned} \text{میزان تغییر نسبی}^1 \text{ در چگالی در نقطه } X &= \lim_{\Delta \rightarrow 0} \left[\frac{P(x; \theta + \Delta) - p(x; \theta)}{\Delta} \cdot \frac{1}{p(x; \theta)} \right] \\ &= \frac{\frac{\partial}{\partial \theta} p(x; \theta)}{p(x; \theta)} = \frac{\partial}{\partial \theta} \ln p(x; \theta) \end{aligned}$$

میانگین توان دوم این تغییر نسبی را میزان اطلاعاتی که X در مورد پارامتر θ دارد می‌نامند و به آن اطلاع فیشر گویند یعنی

$$I_X(\theta) = E_\theta \left[\frac{\partial}{\partial \theta} \ln p(X; \theta) \right]^2 = \int \left(\frac{p'_\theta}{p_\theta} \right)^2 p_\theta d\mu$$

توجه کنید که هر چه این امید ریاضی در یک نقطه θ بیشتر شود آنگاه می‌توان θ را از نقاط همسایه θ بهتر تشخیص داد و بنابراین θ را در نقطه $\theta = \theta$ دقیق‌تر برآورد کرد.

حال فرض کنید شرایط زیر برقرار باشند:

(الف) Θ یک زیر فاصله باز از اعداد حقیقی باشد ($\Theta \subseteq R$)

(ب) توزیع‌های P_θ دارای تکیه‌گاه یکسان باشند و بنابراین $S_X^* = \{x \mid p(x; \theta) > 0\}$ به θ بستگی نداشته باشد.

(ج) $\frac{\partial}{\partial \theta} p(x; \theta)$ برای تمام $\theta \in \Theta$ موجود و متناهی باشد.

$$\frac{\partial}{\partial \theta} \int p(x; \theta) d\mu = \int \frac{\partial}{\partial \theta} p(x; \theta) d\mu \quad (\text{د})$$

$$E \left[\frac{\partial}{\partial \theta} \ln p(x; \theta) \right]^2 > 0 \quad \forall \theta \in \Theta \quad (\text{هـ})$$

$$\frac{\partial}{\partial \theta} E_\theta[\delta] = \frac{\partial}{\partial \theta} \int \delta p_\theta d\mu = \int \delta \frac{\partial}{\partial \theta} p_\theta d\mu \quad (\text{و}) \text{ برای هر برآوردگر } \delta \text{ داشته باشیم که:}$$

این شرایط را شرایط نظم^۲ گویند.

حال اگر قرار دهیم $\psi(x; \theta) = \frac{\partial}{\partial \theta} \ln p(x; \theta)$ آنگاه برای برآوردگر U که $E_\theta(U) = 0$ داریم که

^۱ Relative Rate

^۲ Regularity Condition

$$\begin{aligned} \cdot &= \frac{\partial}{\partial \theta} E_{\theta}(U) \stackrel{(۱)}{=} \int U \frac{\partial}{\partial \theta} p_{\theta} d\mu = \int [U \frac{\partial}{\partial \theta} \ln p_{\theta}] p_{\theta} d\mu = E[U \cdot \frac{\partial}{\partial \theta} \ln p_{\theta}] \\ &= \text{Cov}(U, \frac{\partial}{\partial \theta} \ln p_{\theta}) \end{aligned}$$

بنابراین شرایط قضیه ۵-۱ برقرار است. همچنین چون

$$E \left[\frac{\partial}{\partial \theta} \ln p_{\theta} \right] = \int \left(\frac{\partial}{\partial \theta} \ln p_{\theta} \right) p_{\theta} d\mu = \int \frac{\partial}{\partial \theta} p_{\theta} d\mu \stackrel{(۱)}{=} \frac{\partial}{\partial \theta} \int p_{\theta} d\mu = \frac{\partial}{\partial \theta} (۱) = ۰$$

$$\text{Cov}(\delta, \psi) = E(\delta \psi) = \int \delta \frac{\partial}{\partial \theta} p_{\theta} d\mu \stackrel{(۱)}{=} \frac{\partial}{\partial \theta} \int \delta p_{\theta} d\mu = \frac{\partial}{\partial \theta} E_{\theta}[\delta] = g'(\theta) \quad \text{پس}$$

بنابراین نامساوی (۱) به صورت زیر تبدیل می‌شود.

$$\text{Var}(\delta) \geq \frac{[g'(\theta)]^2}{E \left[\frac{\partial}{\partial \theta} \ln p(x; \theta) \right]^2} = \frac{[g'(\theta)]^2}{I_X(\theta)}$$

این نامساوی را نامساوی اطلاع یا نامساوی کرامر - رانو گویند. به کران پایین بدست آمده برای واریانس، کران پایین کرامر - رانو (CRLB) گویند. توجه کنید که اگر مشتق دوم $\ln p(x; \theta)$ موجود

$$\text{باشد و } \frac{\partial^2}{\partial \theta^2} \int p_{\theta} d\mu = \int \frac{\partial^2}{\partial \theta^2} p_{\theta} d\mu \quad \text{آنگاه}$$

$$I_X(\theta) = E \left[\frac{\partial}{\partial \theta} \ln p(x; \theta) \right]^2 = -E \left[\frac{\partial^2}{\partial \theta^2} \ln p(x; \theta) \right]$$

حال اگر $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیعی با چگالی $p_{\theta} : \theta \in \Theta$ باشند آنگاه

$$I(\theta) = E_{\theta} \left[\frac{\partial}{\partial \theta} \ln p(\tilde{X}, \theta) \right]^2 \quad \text{را میزان اطلاع فیشر که نمونه تصادفی } X_1, X_2, \dots, X_n \text{ در مورد } \theta \text{ دارد، گویند.}$$

$$I(\theta) = E_{\theta} \left[\frac{\partial}{\partial \theta} \ln p(\tilde{X}; \theta) \right]^2 = V_{\theta} \left[\frac{\partial}{\partial \theta} \ln p(\tilde{X}; \theta) \right] = V_{\theta} \left[\frac{\partial}{\partial \theta} \ln \prod_{i=1}^n \ln p(X_i; \theta) \right]$$

$$\stackrel{(۲)}{=} \sum_{i=1}^n V_{\theta} \left[\frac{\partial}{\partial \theta} \ln p(X_i; \theta) \right] \stackrel{(۳)}{=} n V_{\theta} \left[\frac{\partial}{\partial \theta} \ln p(X_1; \theta) \right] = n E_{\theta} \left[\frac{\partial}{\partial \theta} \ln p(X_1; \theta) \right]^2 = n I_{X_1}(\theta)$$

که در آن (۲) از استقلال و (۳) از هم‌توزیع بودن X_i ها نتیجه می‌گردد.

مثال ۵-۱: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع $N(\theta, 1), \theta \in R$ باشند، CRLB را برای واریانس برآوردگر نااریب $g(\theta) = \theta^2$ بدست آورده و آن را با واریانس برآوردگر UMVU پارامتر θ^2 مقایسه کنید.

$$g(\theta) = \theta^2 \Rightarrow g'(\theta) = 2\theta$$

$$\ln f_{\theta}(x) = -\frac{1}{2} \ln(2\pi) - \frac{1}{2} \sum_{i=1}^n (x_i - \theta)^2 \Rightarrow \frac{\partial^2}{\partial \theta^2} \ln f_{\theta}(x) = -1$$

$$\Rightarrow I(\theta) = nI_{X_1}(\theta) = n \Rightarrow CRLB = \frac{(2\theta)^2}{n} = \frac{4\theta^2}{n}$$

در این خانواده $S = \sum_{i=1}^n X_i \sim N(n\theta, n)$ آماره بسنده و کامل برای θ است و $T = \frac{S^2 - n}{n^2}$

$$\blacksquare \quad V(T) = \frac{4\theta^2}{n} + \frac{2}{n^2} > \frac{4\theta^2}{n} = CRLB \quad \text{برآوردگر UMVU برای } \theta^2 \text{ است و}$$

مثال ۵-۲: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع نمایی با پارامتر θ باشند.

CRLB را برای واریانس برآوردگرهای نااریب θ و $\frac{1}{\theta}$ بدست آورده و آن را با واریانس برآوردگرهای UMVU مقایسه کنید.

$$\ln f_{\theta}(x) = -\ln \theta - \frac{x}{\theta} \Rightarrow \frac{\partial^2}{\partial \theta^2} \ln f_{\theta}(x) = \frac{1}{\theta^2} - \frac{2x}{\theta^3}$$

$$I_{X_1}(\theta) = -E_{\theta} \left[\frac{\partial^2}{\partial \theta^2} \ln f_{\theta}(X_1) \right] = -\left(\frac{1}{\theta^2} - \frac{2}{\theta^2} \right) = \frac{1}{\theta^2} \Rightarrow I(\theta) = \frac{n}{\theta^2}$$

$$\text{الف) } g_1(\theta) = \theta \Rightarrow g'_1(\theta) = 1 \Rightarrow CRLB = \frac{1}{n/\theta^2} = \frac{\theta^2}{n}$$

$$T = \bar{X} \text{ is UMVUE of } \theta \Rightarrow V(T) = \frac{\theta^2}{n} = CRLB$$

$$\text{ب) } g_2(\theta) = \frac{1}{\theta} \Rightarrow g'_2(\theta) = \frac{-1}{\theta^2} \Rightarrow CRLB = \frac{(-1/\theta^2)^2}{n/\theta^2} = \frac{1}{n\theta^2}$$

$$T = \frac{n-1}{\sum X_i} \text{ is UMVUE of } \frac{1}{\theta} \Rightarrow V(T) = \frac{1}{(n-2)\theta^2} > \frac{1}{n\theta^2} = CRLB$$

یک سوال که در مورد نامساوی کرامر - راثو مطرح است این است که چه موقع نامساوی به مساوی تبدیل می‌شود و یا به عبارتی برای چه برآوردگرهای ناریبی واریانس برآوردگر با CRLB برابر می‌شود.

قضیه ۵-۲: فرض کنید شرایط نظم برقرار باشد و T برآوردگری باشد که $V_{\theta}(T) < \infty, \forall \theta \in \Theta$. در این صورت $V_{\theta}(T) = CRLB$ اگر و فقط اگر تابع چگالی $p_{\theta}(x)$ متعلق به خانواده نمایی یک پارامتری زیر باشد.

$$p_{\theta}(x) = e^{c(\theta)T(x) - B(\theta)} h(x)$$

در این حالت $g(\theta) = E_{\theta}[T(X)] = \frac{B'(\theta)}{c'(\theta)}$ و $I(\theta) = c'(\theta)g'(\theta)$ می‌باشد. (یعنی تنها در خانواده‌های نمایی بفرم فوق و تنها برای توابع پارامتری بفرم فوق واریانس برآوردگر ناریب $g(\theta)$ با CRLB برابر می‌شود.)

اثبات:

$$\begin{aligned} \Leftrightarrow \rho^2 \left(\frac{\partial}{\partial \theta} \ln f_{\theta}(X), T(X) \right) &= 1 \\ \Leftrightarrow \frac{\partial}{\partial \theta} \ln f_{\theta}(X) &= a(\theta)T(X) + b(\theta) \\ \Leftrightarrow \ln f_{\theta}(X) &= T(X) \int a(\theta) d\theta + \int b(\theta) d\theta \\ &= T(X)c(\theta) - B(\theta) + R(X) \\ \Leftrightarrow f_{\theta}(X) &= e^{c(\theta)T(X) - B(\theta) + R(X)} \\ \Leftrightarrow f_{\theta}(x) &= e^{c(\theta)T(x) - B(\theta)} h(x) \\ \text{در این خانواده} \quad g(\theta) &= E_{\theta}(T(X)) = \frac{B'(\theta)}{c'(\theta)} \text{ و} \\ I(\theta) &= -E[c''(\theta)T(X) - B''(\theta)] = -\frac{c''(\theta)B'(\theta) - B''(\theta)c'(\theta)}{c'(\theta)} \end{aligned}$$

$$\blacksquare \quad = \left[\frac{B'(\theta)}{C'(\theta)} \right]' c'(\theta) = g'(\theta) c'(\theta)$$

مثال ۳-۵: نشان دهید که تنها توزیع گسسته‌ای که دارای میانگین و واریانس یکسان است و $V_\theta(X)$ در آن با CRLB برابر می‌شود توزیع پواسون است.

$$g(\theta) = E_\theta(X) = V_\theta(X) = \theta$$

$$V_\theta(X) = CRLB \stackrel{(*)}{\Leftrightarrow} \theta = \frac{(\theta')^2}{I_X(\theta)} = \frac{1}{I_X(\theta)} = \frac{1}{g'(\theta)c'(\theta)} = \frac{1}{c'(\theta)}$$

$$\Rightarrow c'(\theta) = \frac{1}{\theta} \Rightarrow c(\theta) = \ln \theta \stackrel{(*)}{\Rightarrow} p_\theta(x) = e^{(\ln \theta)x - B(\theta)} h(x)$$

که در آن $(*)$ از قضیه ۲-۵ بدست می‌آید. اما $\theta = g(\theta) = E_\theta(X) = \frac{+B'(\theta)}{1/\theta}$ پس $B'(\theta) = 1$ یا

$$\blacksquare \quad p_\theta(x) = \frac{\theta^x e^{-\theta}}{x!} \quad \text{پس } h(x) = x! \quad \text{و در نتیجه } B(\theta) = \theta$$

قضیه ۳-۵: اگر $\theta = h(\zeta)$ و h مشتق پذیر باشد آن گاه $I^*(\zeta) = I[h(\zeta)][h'(\zeta)]^2$ اما CRLB برای واریانس برآوردگرهای θ و ζ یکی است. (یعنی CRLB با تبدیلات روی پارامترها تغییر نمی‌کند.)

$$I(\zeta) = E \left[\frac{\partial}{\partial \zeta} \ln f_\zeta(X) \right]^2 = E \left[\frac{\partial}{\partial \zeta} \ln f_\theta(X) \cdot \frac{\partial \theta}{\partial \zeta} \right]^2 \quad \text{اثبات:}$$

$$= E \left[\frac{\partial}{\partial \theta} \ln f_\theta(X) \right]^2 \left[\frac{\partial \theta}{\partial \zeta} \right]^2 = I(\theta) [h'(\zeta)]^2$$

$$\frac{\partial}{\partial \zeta} E[T] = \frac{\partial}{\partial \theta} E[T] \frac{\partial \theta}{\partial \zeta} = g'(\theta) h'(\zeta)$$

$$\Rightarrow CRLB_\zeta = \frac{[g'(\zeta)]^2}{I(\zeta)} = \frac{[g'(\theta)h'(\zeta)]^2}{I(\theta)[h'(\zeta)]^2} = \frac{[g'(\theta)]^2}{I(\theta)} = CRLB_\theta \quad \blacksquare$$

مثال ۴-۵: فرض کنید X متعلق به یک خانواده نمایی با تابع چگالی

$$p_\theta(x) = e^{C(\theta)T(x) - B(\theta)} h(x) \quad \text{باشد و فرض کنید } g(\theta) = E_\theta[T(X)] \quad \text{در این صورت نشان}$$

$$\text{دهید که } I[g(\theta)] = \frac{1}{\text{Var}(T)}$$

$$I(\theta) = V_{\theta} \left[\frac{\partial}{\partial \theta} \ln p_{\theta}(X) \right] V_{\theta} [c'(\theta)T(X) - B'(\theta)] = [c'(\theta)]^T V_{\theta}(T) \quad \text{حل:}$$

$$I(g(\theta)) = \frac{I(\theta)}{[g'(\theta)]^T} = \left[\frac{c'(\theta)}{g'(\theta)} \right]^T V_{\theta}(T) = \frac{1}{V_{\theta}(T)}$$

$$\left(\text{Since } V_{\theta}(T) = \frac{B''(\theta) - c''(\theta)g(\theta)}{[c'(\theta)]^T} = \frac{g'(\theta)}{c'(\theta)} \right)$$

بخش ۶: نامساوی اطلاع (کرامر - رائو) در حالت چند پارامتری

در این قسمت فرض می‌کنیم $\theta = (\theta_1, \dots, \theta_k)'$ باشد و شرایط نظم را به صورت زیر در نظر می‌گیریم:

الف - $\Theta \subseteq R^k$ یک گوی باز در فضای R^k باشد، $\Theta \subseteq R^k$.

ب - توزیع‌های P_{θ} دارای تکیه گاه یکسان باشند و بنابراین $S_X^* = \{x \mid p_{\theta}(x) > 0\}$ به θ بستگی نداشته باشد.

ج - $\frac{\partial}{\partial \theta_i} f_{\theta}(x)$ برای هر $i = 1, 2, \dots, k$ و برای تمام $\theta \in \Theta$ موجود باشد.

$$\forall i = 1, 2, \dots, k \quad \frac{\partial}{\partial \theta_i} \int f_{\theta}(x) d\mu(x) = \int \frac{\partial}{\partial \theta_i} f_{\theta}(x) d\mu(x) \quad \text{د -}$$

ه - $\frac{\partial}{\partial \theta_i} \ln f_{\theta}(x)$, $i = 1, 2, \dots, k$ مستقل خطی باشند.

و - برای هر برآوردگر T پارامتر $g(\theta)$ داشته باشیم که

$$\forall i = 1, 2, \dots, k \quad \frac{\partial}{\partial \theta_i} E_{\theta}[T(\tilde{X})] = \frac{\partial}{\partial \theta_i} \int T(x) f_{\theta}(x) d\mu(x) = \int T(x) \frac{\partial}{\partial \theta_i} f_{\theta}(x) d\mu(x)$$

تعریف ۶-۱: ماتریس $I(\theta) = (I_{ij}(\theta))_{k \times k}$ که در آن

$$I_{ij}(\theta) = \text{Cov} \left(\frac{\partial}{\partial \theta_i} \ln f_{\theta}(X), \frac{\partial}{\partial \theta_j} \ln f_{\theta}(X) \right) = E \left[\frac{\partial}{\partial \theta_i} \ln f_{\theta}(X) \frac{\partial}{\partial \theta_j} \ln f_{\theta}(X) \right]$$

را ماتریس اطلاع فیشر θ گویند. می توان نشان داد که اگر شرط (ه) برقرار باشد آنگاه ماتریس $I(\theta)$

معین مثبت است (زیرا ماتریس واریانس کواریانس معین نامنفی است) $\left(0 \leq V\left(\sum_{i=1}^n a_i X_i\right) = \underline{a}' \Sigma \underline{a}\right)$

و در صورتی معین مثبت است که متغیرهای تصادفی بردار تصادفی مربوطه وابسته خطی نباشند.

همچنین اگر در شرایط (ج) و (د) بتوان مشتق مرتبه اول را با مشتق مرتبه دوم عوض کرد، آنگاه

$$I_{ij}(\theta) = -E \left[\frac{\partial^2}{\partial \theta_i \partial \theta_j} \ln f_{\theta}(X) \right] \quad \text{می توان نشان داد که}$$

توجه کنید که اگر تغییر پارامتر زیر را داشته باشیم $\theta_i = h_i(\zeta_1, \zeta_2, \dots, \zeta_k) \quad i = 1, 2, \dots, k$ و

$$\text{قراردهیم } J_{k \times k} = \left[\frac{\partial \theta_i}{\partial \zeta_j} \right]_{ij} \text{ آنگاه}$$

$$\begin{aligned} I_{ij}^*(\zeta) &= E \left[\frac{\partial}{\partial \zeta_i} \ln p_{\theta(\zeta)}(X) \frac{\partial}{\partial \zeta_j} \ln p_{\theta(\zeta)}(X) \right] \\ &= \sum_{r=1}^k \sum_{\ell=1}^k E \left[\frac{\partial}{\partial \theta_r} \ln p_{\theta}(X) \frac{\partial}{\partial \theta_{\ell}} \ln p_{\theta}(X) \right] \frac{\partial \theta_r}{\partial \zeta_i} \cdot \frac{\partial \theta_{\ell}}{\partial \zeta_j} \end{aligned}$$

$$\Rightarrow I^*(\zeta) = J(\theta) J'$$

مثال ۶-۱: فرض کنید X دارای چگالی از خانواده نمایی k پارامتری بفرم زیر باشد

$$p_{\theta}(x) = \left[e^{\sum_{j=1}^k c_j(\theta) T_j(x) - B(\theta)} \right] h(x)$$

و قرار دهیم $\tau_i = E(T_i(X)) \quad i = 1, 2, \dots, k$ در این صورت نشان دهید که $I(\tau) = C^{-1}$ که C

ماتریس واریانس کواریانس $(T_1, \dots, T_k)$ است.

حل: اگر تغییر پارامتر $\eta_j = c_j(\theta), \quad j = 1, \dots, k$ را بدهیم آنگاه توزیع به صورت

$$p_{\theta}(x) = \left[e^{\sum_{j=1}^k \eta_j T_j(x) - A(\underline{\eta})} \right] h(x)$$

در می آید و در این حالت

$$I_{j,k}^*(\eta) = \frac{\partial^2}{\partial \eta_j \partial \eta_k} A(\eta) \Rightarrow I^*(\eta) = C = [Cov(T_j, T_k)]_{jk}$$

همچنین $\tau_j = E(T_j(X)) = \frac{\partial}{\partial \eta_j} A(\eta)$ پس طبق مطالب فوق

$$J = \left[\frac{\partial \tau_j}{\partial \eta_i} \right]_{i,j} = \left[\frac{\partial^2}{\partial \eta_i \partial \eta_j} A(\eta) \right]_{i,j} = [Cov(T_i, T_j)]_{i,j} = C$$

و بنابراین

$$C = I^*(\eta) = JI(\tau)J' = CI(\tau)C' \Rightarrow I(\tau) = C^{-1} \quad \blacksquare$$

برای ساختن نامساوی اطلاع ابتدا قضیه زیر را می آوریم.

قضیه ۶-۱: برای هر برآوردگر T برای $g(\theta)$ و هر تابع $\psi_i(x, \theta)$ با گشتاور مرتبه دوم متناهی داریم که

$$V_\theta(T) \geq \gamma' C^{-1} \gamma \quad (۱)$$

که در آن $\gamma' = (\gamma_1, \dots, \gamma_k)$ و $C = (C_{ij})$ و $\gamma_i = Cov(T, \psi_i)$ و $C_{ij} = Cov(\psi_i, \psi_j)$.
طرف راست رابطه (۱) تنها از طریق $g(\theta) = E_\theta(T)$ بستگی به T دارد به شرط آنکه هر کدام از ψ_i ها در شرط $Cov(U, \psi_i) = 0$ برای هر $U \in \Delta$ صدق کنند.

اثبات: برای هر ثابت $a_1, a_2, \dots, a_k$ داریم که

$$\left[Cov\left(T, \sum_{i=1}^k a_i \psi_i\right) \right]^2 \leq V(T) V\left(\sum_{i=1}^k a_i \psi_i\right)$$

$$\Rightarrow V(T) \geq \frac{\left[Cov\left(T, \sum_{i=1}^k a_i \psi_i\right) \right]^2}{V\left(\sum_{i=1}^k a_i \psi_i\right)}$$

$$Cov\left(T, \sum_{i=1}^k a_i \psi_i\right) = \sum_{i=1}^k a_i Cov(T, \psi_i) = \sum_{i=1}^k a_i \gamma_i = a' \gamma$$

اما

$$V \left(\sum_{i=1}^k a_i \psi_i \right) = V(\underline{a}' \underline{\psi}) = \underline{a}' V(\underline{\psi}) \underline{a} = \underline{a}' C \underline{a}$$

با استفاده از نامساوی کشی - شوارتز تعمیم یافته ($\underline{a}' C \underline{a} (\underline{\gamma}' C^{-1} \underline{\gamma}) \leq (\underline{a}' \underline{\gamma})^2, \forall \underline{a} \neq \underline{z}$) که در آن C معین مثبت است) داریم که

$$V(T) \geq \frac{[\underline{a}' \underline{\gamma}]^2}{\underline{a}' C \underline{a}} \quad \forall \underline{a} \neq \underline{z} \Rightarrow V(T) \geq \max_{\underline{a} \neq \underline{z}} \frac{[\underline{a}' \underline{\gamma}]^2}{\underline{a}' C \underline{a}} = \underline{\gamma}' C^{-1} \underline{\gamma}$$

$$\blacksquare \quad \therefore V(T) \geq \underline{\gamma}' C^{-1} \underline{\gamma}$$

قضیه ۶-۲: فرض کنید X دارای تابع چگالی $f_{\theta}(x)$ باشد اگر T برآوردگر برای $g(\theta)$ باشد که $E(T^2) < \infty$ و شرایط نظم برقرار باشد آنگاه

$$V_{\theta}(T) \geq \underline{\alpha}' I^{-1}(\theta) \underline{\alpha} \quad (2)$$

که در آن $\underline{\alpha}' = (\alpha_1, \dots, \alpha_k)$ و $\alpha_i = \frac{\partial}{\partial \theta_i} E_{\theta}[T(X)]$ و نامساوی فوق را نامساوی اطلاع گویند.

اثبات: در قضیه ۶-۱ قرار می دهیم $\psi_i(x, \theta) = \frac{\partial}{\partial \theta_i} \ln f_{\theta}(x)$ بنابراین ماتریس C حاصل بنابر شرط نظم (ه) معین مثبت است و نتیجه به دست می آید.

$$\blacksquare \quad \gamma_i = E(T \psi_i) = \int T \frac{\partial}{\partial \theta_i} f_{\theta}(x) d\mu(x) = \frac{\partial}{\partial \theta_i} E_{\theta}(T(X)) = \alpha_i$$

مثال ۶-۲: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\mu, \sigma^2)$ باشند. CRLB برای واریانس برآوردگر نااریب پارامتر $E(X^2) = \mu^2 + \sigma^2$ را بدست آورید.

$$\ln f_{\theta}(\underline{X}) = -\frac{n}{2} \ln(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \mu)^2 \quad \text{حل:}$$

$$\frac{\partial}{\partial \mu} \ln f_{\theta}(\underline{x}) = \frac{1}{\sigma^2} \sum_{i=1}^n (x_i - \mu) \quad \frac{\partial^2}{\partial \mu^2} \ln f_{\theta}(\underline{x}) = -\frac{n}{\sigma^2}$$

$$\frac{\partial}{\partial \sigma^2} \ln f_{\theta}(\underline{x}) = -\frac{n}{2\sigma^2} + \frac{1}{2\sigma^4} \sum_{i=1}^n (x_i - \mu)^2$$

$$\frac{\partial^2}{\partial (\sigma^2)^2} \ln f_{\theta}(\underline{x}) = \frac{n}{2\sigma^4} - \frac{1}{\sigma^6} \sum_{i=1}^n (x_i - \mu)^2 \quad \frac{\partial^2}{\partial \mu \partial \sigma^2} \ln f_{\theta}(\underline{x}) = -\frac{1}{\sigma^4} \sum_{i=1}^n (x_i - \mu)$$

$$\therefore I(\mu, \sigma^2) = \begin{bmatrix} \frac{n}{\sigma^2} & \cdot \\ \cdot & \frac{n}{2\sigma^4} \end{bmatrix}$$

$$\blacksquare \quad \alpha' = (2\mu, 1) \Rightarrow CRLB = \alpha' I^{-1} \alpha = \begin{bmatrix} 2\mu & 1 \end{bmatrix} \begin{bmatrix} \frac{\sigma^2}{n} & \cdot \\ \cdot & \frac{2\sigma^4}{n} \end{bmatrix} \begin{bmatrix} 2\mu \\ 1 \end{bmatrix} = \frac{2\mu^2 \sigma^2}{n} + \frac{2\sigma^4}{n}$$

توجه: اگر $I_{ii}(\theta)$ اطلاع فشر در حالتی باشد که θ_i نامعلوم و مابقی اعضای θ بجز θ_i معلوم باشند T آنگاه $I_{ii}^{-1}(\theta) \leq [I^{-1}(\theta)]_{ii}$ و در نتیجه

$$\frac{\left[\frac{\partial g(\theta)}{\partial \theta_i} \right]^2}{I_{ii}(\theta)} \leq \left[\frac{\partial g(\theta)}{\partial \theta_i} \right]^2 [I^{-1}(\theta)]_{ii}$$

یعنی اینکه تعریف اطلاع فشر چند پارامتری و نامساوی اطلاع چندپارامتری باعث می‌شود که CRLB (در یک پارامتری) تیزتر^۱ شود. برای رسیدن به این نامساوی توجه کنید که اگر در قضیه ۶-۱ قرار دهیم $\alpha' = (\cdot, \dots, \cdot, 1, \cdot, \dots, \cdot)$ آنگاه

$$CRLB = \frac{[Cov(T, \psi_i)]^2}{V(\psi_i)} = \frac{\alpha' \alpha}{\alpha' I(\theta) \alpha} \leq \alpha' I^{-1}(\theta) \alpha, \quad \alpha_k = \begin{cases} \frac{\partial g(\theta)}{\partial \theta_i} & k = i \\ \cdot & k \neq i \end{cases}$$

$$\Rightarrow \frac{\left[\frac{\partial g(\theta)}{\partial \theta_i} \right]^2}{I_{ii}(\theta)} \leq \left[\frac{\partial g(\theta)}{\partial \theta_i} \right]^2 [I^{-1}(\theta)]_{ii} \Rightarrow I_{ii}^{-1}(\theta) \leq [I^{-1}(\theta)]_{ii}$$

همچنین توجه کنید که اگر پارامترها بر هم متعامد باشند آنگاه $\forall i \neq j \quad I_{ij}(\theta) = 0$ و در این صورت نامساوی فوق به تساوی تبدیل می‌شود. مثال ۶-۲ این حالت را نشان می‌دهد.

^۱ Sharp

فصل ۳

پایایی

بخش ۱: مفهوم پایایی^۱ و کلیات

در فصل قبل کلاس برآوردگرها را محدود به کلاس برآوردگرهای نااریب کردیم و در این کلاس بهترین برآوردگرهای نااریب را پیدا کردیم. در این قسمت می‌خواهیم خود را محدود به کلاس دیگری از برآوردگرها، یعنی برآوردگرهای پایا کنیم و در این کلاس بهترین برآوردگر پایا را پیدا کنیم. برای پی بردن به مفهوم پایایی به مثال زیر توجه کنید.

مثال ۱-۱: فرض کنید متغیر تصادفی X طول عمر یک قطعه برحسب ساعت باشد که دارای توزیع

$$f_{\theta}(x) = \frac{1}{\theta} e^{-x/\theta} \quad x > 0, \theta > 0.$$

نمایی با تابع چگالی زیر است:

می‌خواهیم براساس یک مشاهده X پارامتر θ را براساس تابع زیان $L(\theta, \delta) = (1 - \frac{\delta}{\theta})^2$ برآورد کنیم. فرض کنید $\delta(x)$ یک برآورد پارامتر θ براساس X (طول عمر برحسب ساعت) باشد.

حال اگر شخص دیگری طول عمر این قطعه را برحسب دقیقه اندازه‌گیری کند و در این حالت متغیر تصادفی Y را طول عمر قطعه برحسب دقیقه در نظر بگیرد در این صورت بایستی $Y = 60X$ باشد. اگر قرار دهیم $\eta = 60\theta$ در این صورت به راحتی دیده می‌شود که Y دارای تابع چگالی زیر است:

$$f_{\eta}(y) = \frac{1}{60} f_{\theta}\left(\frac{y}{60}\right) = \frac{1}{60\theta} e^{-\frac{y}{60\theta}} \Rightarrow f_{\eta}(y) = \frac{1}{\eta} e^{-y/\eta} \quad y > 0, \eta > 0.$$

حال اگر در این اندازه‌گیری جدید برآورد δ^* را برای پارامتر η بدست آوریم انتظار داریم که

$$\delta^*(y) = 60\delta(X) \quad (1)$$

باشد (زیرا منطقی به نظر می‌رسد که برآوردگرهای یکسانی را صرف نظر از نوع اندازه‌گیری بکار ببریم). در این حالت:

$$L(\eta, \delta^*) = \left(1 - \frac{\delta^*}{\eta}\right)^2 = \left(1 - \frac{60\delta}{60\theta}\right)^2 = \left(1 - \frac{\delta}{\theta}\right)^2 = L(\theta, \delta)$$

بنابراین ساختار کلی مسئله با اندازه‌گیری برحسب دقیقه (یعنی کلاس توابع چگالی، کلاس فضای پارامتری و کلاس توابع زیان) دقیقاً با ساختار کلی مسئله با اندازه‌گیری برحسب ساعت یکی می‌شود.

^۱ Invariance

در نتیجه منطقی به نظر می‌رسد که در این مسئله برآوردهایی را مورد استفاده قرار دهیم که تحت این دو ساختار یکسان باشند یعنی:

$$\delta(y) = \delta^*(y) \quad (۲)$$

بنابراین از روابط (۱) و (۲) داریم که:

$$\delta(\epsilon \cdot x) = \epsilon \cdot \delta(x) \quad (۳)$$

برآوردهای δ که در شرط (۳) صدق می‌کند را برآوردهای هم پایا^۱ می‌نامند و همانطور که دیده می‌شود کلاس برآوردهای هم پایا یک قسمت کوچک از کلاس کلیه برآوردها است.

مسئله فوق برای حالت کلی‌تر یعنی $Y = cX$, $c > 0$ نیز برقرار است و در این حالت برآوردهای هم پایا به صورت زیر است:

$$\delta(cx) = c\delta(x) \quad \forall c > 0 \quad (۴)$$

چون این رابطه برای هر x برقرار است پس برای $x = \frac{1}{c}$ نیز برقرار است و در نتیجه از رابطه (۴) داریم که

$$\delta(1) = \frac{1}{x} \delta(x) \Rightarrow \delta(x) = \delta(1)x \Rightarrow \delta(x) = kx$$

بنابراین در این حالت کلاس برآوردهای هم پایا عبارت است از:

$$D_E = \{\delta \in D \mid \delta(x) = kx, \forall x, \forall k > 0\}$$

در این فصل می‌خواهیم در بین برآوردهای هم پایا برای یک پارامتر مجهول آن برآوردهای را پیدا کنیم که دارای کمترین مخاطره^۲ باشد، که به آن برآوردهای هم پایا با کمترین مخاطره (MRE^3) گویند. در این مثال داریم که:

$$\begin{aligned} R(\theta, \delta) &= E_{\theta} \left[\left(1 - \frac{\delta(X)}{\theta} \right)^2 \right] = \frac{1}{\theta^2} E_{\theta} \left[(kX - \theta)^2 \right] = \frac{1}{\theta^2} \left[V(kX) + [E(kX - \theta)]^2 \right] \\ &= \frac{1}{\theta^2} \left[k^2 \theta^2 + (k - 1)^2 \theta^2 \right] = k^2 + (k - 1)^2 = 2k^2 - 2k + 1 \end{aligned}$$

^۱ Equivariance

^۲ Risk

^۳ Minimum Risk Equivariant Estimator

بنابراین تابع مخاطره ثابت است و به θ بستگی ندارد. در نتیجه برآوردگر MRE با مینیمم کردن $R(\theta, \delta)$ نسبت به k بدست می آید:

$$\frac{\partial R(\theta, \delta)}{\partial k} = 4k - 2 = 0 \Rightarrow k = \frac{1}{2}, \quad \frac{\partial^2 R(\theta, \delta)}{\partial k^2} = 4 > 0.$$

بنابراین برآوردگر $\delta^*(X) = \frac{1}{2}X$ برآوردگری هم پایا با کمترین مخاطره است. پس برآوردگر MRE پارامتر θ است. ■

اصول پایایی:

با توجه به مثال بالا می توان خود را محدود به کلاس برآوردگرهای پایا کنیم. اما پایایی را می توان از دو جنبه در نظر گرفت.

۱- پایایی اندازه گیری^۱ (پایایی طبیعی یا منطقی یا تابعی)

این جنبه از پایایی می گوید که استنباط نبایستی به نوع اندازه گیری داده ها بستگی داشته باشد. (مثلا در مثال ۱-۱ اگر $\delta(X)$ برای برآورد θ بکار برده می شود طبیعی به نظر می رسد که $\delta^*(Y) = c\delta(X)$ برای برآورد $\eta = c\theta$ بکار برده شود.)

۲- پایایی ساختاری^۲ (یا ریاضی)

این جنبه از پایایی می گوید که اگر دو مسئله استنباطی دارای یک ساختار از لحاظ فضای پارامتری، تابع چگالی و تابع زیان باشند آنگاه بایستی روش استنباط یکسانی را در دو مسئله مورد استفاده قرار دهیم. (مثلا در مثال ۱-۱ برآورد $\delta(Y) = \delta(cX)$ را برای برآورد $\eta = c\theta$ بکار ببریم). با تلفیق این دو جنبه می توان برآوردگرهای هم پایا را پیدا کرد (مثلا در مثال ۱-۱ $\delta(cX) = c\delta(X)$)

فرمول بندی و تعاریف پایایی

با توجه به مطالب بالا در حالت کلی ما یک نمونه تصادفی $\underline{X} = (X_1, X_2, \dots, X_n)$ با مقادیر $\underline{x}$ و تابع چگالی یا احتمال توام $f_\theta(\underline{x}) = f_\theta(x_1, \dots, x_n)$ داریم. پس بوسیله یک تبدیل یک به یک و

^۱ Measurement Invariance

^۲ Formal Invariance

پوشای $Y = g(X)$ داده‌های X را به داده‌های جدید Y تبدیل می‌کنیم که این تبدیلات یک گروه از تبدیلات G را بوجود می‌آورند.

تعریف ۱-۱: یک گروه از تبدیلات روی $\mathcal{X}$ که آن را با G نمایش می‌دهیم یک مجموعه تبدیلات اندازه‌پذیر یک به یک و پوشا روی $\mathcal{X}$ است که در شرایط زیر صدق می‌کند:

- (i) if $g_1, g_2 \in G \Rightarrow g_1 g_2 \in G$
- (ii) $\forall g \in G \quad \exists g^{-1} \in G \quad \exists g^{-1}(g(x)) = x$
- (iii) $\exists e \in G \quad \exists e(x) = x$

مثال ۱-۲: در مثال ۱-۱ گروه تبدیلات زیر را داریم $G = \{g_c : g_c(X) = cX, c > 0\}$ که به آن گروه تبدیلات ضربی یا مقیاسی گویند. توابع g_c یک به یک و پوشا هستند.

$$g_{c_2} g_{c_1}(x) = g_{c_2}(g_{c_1}(x)) = g_{c_2}(c_1 x) = c_2 c_1 x = g_{c_2 c_1}(x) \Rightarrow g_{c_2} g_{c_1} = g_{c_2 c_1} \in G$$

$$g_c^{-1} = g_{c^{-1}} \in G, e = g_1 \in G$$

تعریف ۱-۲: فرض کنید $F = \{f_\theta(x) : \theta \in \Theta\}$ یک خانواده از توزیع‌های احتمال برای X باشند و G یک گروه از تبدیلات روی $\mathcal{X}$ باشد. خانواده F را تحت گروه G پایا^۱ گویند اگر برای هر $\theta \in \Theta$ و $g \in G$ یک $\theta' \in \Theta$ یکتا وجود داشته باشد بطوریکه $Y = g(X)$ دارای تابع چگالی $f_{\theta'}(y)$ متعلق به خانواده F باشد در این حالت θ' را با $\bar{g}(\theta)$ نمایش می‌دهند. ■

مثلا در مثال ۱-۱ تحت گروه تبدیلات $G = \{g_c : g_c(X) = cX, c > 0\}$ خانواده توزیع نمایی پایا می‌باشد.

$$f_\theta(x) = \frac{1}{\theta} e^{-x/\theta} \quad \begin{matrix} x > 0 \\ \theta > 0 \end{matrix} \quad \xrightarrow[\eta=c\theta]{Y=cX} f_\eta(y) = \frac{1}{\eta} e^{-y/\eta} \quad \begin{matrix} y > 0 \\ \eta > 0 \end{matrix}$$

اگر قرار دهیم $P_\theta^X(A) = P_\theta(X \in A)$ آنگاه تعریف ۱-۲ می‌گوید که خانواده F تحت تبدیلات G پایا است اگر برای هر $\theta \in \Theta$ و $g \in G$ یک $\bar{g}(\theta) \in \Theta$ موجود باشد بطوری که $P_\theta^{g(X)} = P_{\bar{g}(\theta)}^X$ و یا:

^۱ invariant

$$P_{\theta}(g(X) \in A) = P_{\bar{g}(\theta)}(X \in A) \quad (5)$$

زیرا

$$\begin{aligned} P_{\theta}^{g(X)}(A) &= P_{\theta}(g(X) \in A) = \int_{\{x|g(x) \in A\}} f_{\theta}(x) d\mu(x) \stackrel{(*)}{=} \int_{\{x|g(x) \in A\}} f_{\bar{g}(\theta)}(g(x)) d\mu(x) \\ &\stackrel{Y=g(X)}{=} \int_{\{y|y \in A\}} f_{\bar{g}(\theta)}(y) dv(y) = P_{\bar{g}(\theta)}(X \in A) = P_{\bar{g}(\theta)}^X(A) \end{aligned}$$

که در آن تساوی (*) از پایایی توزیع نتیجه می‌گردد. از رابطه (5) برای هر تابع انتگرال پذیر h نتیجه می‌شود که

$$E_{\theta}[h(g(X))] = E_{\bar{g}(\theta)}[h(X)] \quad (6)$$

قضیه ۱-۱: اگر توزیع‌های $f_{\theta} \in F$ مجزا باشند و خانواده F تحت تبدیلات G پایا باشد آنگاه گروه تبدیلات $\bar{G} = \{\bar{g} : g \in G\}$ روی Θ تشکیل یک گروه از تبدیلات را می‌دهد.

اثبات: مجزا بودن f_{θ} بدین معنی است که $\theta_1 \neq \theta_2 \Rightarrow p_{\theta_1}^X \neq p_{\theta_2}^X$. حال شرایط گروه تبدیل بودن G را بررسی می‌کنیم:

(i) $\bar{g} \in \bar{G}$ یک به یک است

$$\text{Let } C = g(A), \bar{g}(\theta_1) = \bar{g}(\theta_2) \Rightarrow P_{\bar{g}(\theta_1)}(X \in C) = P_{\bar{g}(\theta_2)}(X \in C)$$

$$\stackrel{(\Delta)}{\Rightarrow} P_{\theta_1}(g(X) \in C) = P_{\theta_2}(g(X) \in C)$$

$$\Rightarrow P_{\theta_1}(X \in g^{-1}(C)) = P_{\theta_2}(X \in g^{-1}(C)) \Rightarrow P_{\theta_1}(X \in A) = P_{\theta_2}(X \in A) \stackrel{(*)}{\Rightarrow} \theta_1 = \theta_2$$

(ii) $\bar{g} \in \bar{G}$ is onto, i.e., $\forall \theta \in \Theta \exists \theta^* \in \Theta \ni \bar{g}(\theta^*) = \theta$

$$P_{\theta}(X \in A) = P_{\theta}(g^{-1}(g(X)) \in A) \stackrel{(\Delta)}{=} P_{\bar{g}^{-1}(\theta)}(g(X) \in A)$$

$$= P_{\bar{g}(g^{-1}(\theta))}(X \in A) = P_{\bar{g}(\theta^*)}(X \in A)$$

$$\stackrel{(*)}{\Rightarrow} \bar{g}(\theta^*) = \theta \text{ where } \theta^* = \overline{g^{-1}(\theta)}$$

$$(iii) \quad \bar{g}_1, \bar{g}_r \in \bar{G}^- \Rightarrow \bar{g}_r \bar{g}_1 \in \bar{G}^-$$

ابتدا نشان می‌دهیم که $\overline{g_r g_1} = \bar{g}_r \bar{g}_1$

$$\begin{aligned} P_{\overline{g_r g_1}(\theta)}(X \in A) &\stackrel{(i)}{=} P_{\theta}(g_r g_1(X) \in A) = P_{\theta}(g_1(X) \in g_r^{-1}(A)) \stackrel{(i)}{=} P_{\bar{g}_1(\theta)}(X \in g_r^{-1}(A)) \\ &\stackrel{(i)}{=} P_{\bar{g}_1(\theta)}(g_r(X) \in A) = P_{\bar{g}_r \bar{g}_1(\theta)}(X \in A) \\ &\stackrel{(*)}{\Rightarrow} \overline{g_1 g_r} = \bar{g}_r \bar{g}_1 \in \bar{G}^- \quad (\text{since } g_r g_1 \in G^-) \end{aligned}$$

$$(iv) \quad \forall \bar{g} \in \bar{G}^- \quad \exists \bar{g}^{-1} \in \bar{G}^- \ni \bar{g} \bar{g}^{-1}(\theta) = \theta$$

ابتدا نشان می‌دهیم که $\overline{g^{-1}} = \bar{g}^{-1}$

$$\begin{aligned} P_{\theta}(X \in A) &= P_{\theta}(g^{-1}g(X) \in A) \stackrel{(i)}{=} P_{\overline{g^{-1}g}(\theta)}(g(X) \in A) = P_{\bar{g}(\overline{g^{-1}(\theta)})}(X \in A) \\ &\stackrel{(*)}{\Rightarrow} \overline{g(g^{-1}(\theta))} = \theta \Rightarrow \overline{g^{-1}} = \bar{g}^{-1} \end{aligned}$$

$$g \in G \Rightarrow g^{-1} \in G \Rightarrow \bar{g} \in \bar{G}^- \text{ and } \bar{g}^{-1} \in \bar{G}^- \quad \text{از طرفی}$$

$$\Rightarrow \bar{g}^{-1} = \overline{g^{-1}} \in \bar{G}^- \quad \text{and} \quad \overline{g g^{-1}(\theta)} = \bar{g}(\bar{g}^{-1}(\theta)) = \overline{g g^{-1}(\theta)} = \bar{e}(\theta) = \theta$$

$$(v) \quad \exists \bar{e} \in \bar{G}^- \ni \bar{e}(\theta) = \theta$$

$$P_{\theta}(X \in A) = P_{\theta}(e(X) \in A) \stackrel{(i)}{=} P_{\bar{e}(\theta)}(X \in A) \stackrel{(*)}{\Rightarrow} \bar{e}(\theta) = \theta, \quad \bar{e} \in \bar{G}^-$$

■ که در آن کلیه روابط (*) از مجزا بودن بدست آمده است.

برای مثال، در مثال ۱-۱، $\bar{G}^- = \{\bar{g}_c \mid \bar{g}_c(\theta) = c\theta, c > 0\}$ ، گروه تبدیلات بوجود آمده تحت خانواده توزیع های پایای نمایی، روی فضای پارامتری Θ می‌باشد.

تعریف ۱-۳: فرض کنید خانواده $F = \{f_{\theta}(x) : \theta \in \Theta\}$ تحت گروه تبدیلات G پایا باشد. در این صورت تابع زیان $L(\theta, \delta)$ را تحت گروه G پایا گویند اگر برای هر $g \in G$ و هر $\delta \in D$ (کلاس برآوردگرها) یک $\delta^* \in D$ یکتا وجود داشته باشد بطوری که

$$L(\theta, \delta) = L(\bar{g}(\theta), \delta^*) \quad \forall \theta \in \Theta$$

در این حالت δ^* را بوسیله $\tilde{g}(\delta)$ نشان می‌دهند و مسئله مورد نظر را تحت گروه تبدیلات G پایا گویند. ■

برای مثال، در مثال ۱-۱ تابع زیان $L(\theta, \delta) = (1 - \frac{\delta}{\theta})^2$ تحت تبدیلات $G = \{g_c \mid g_c(X) = cX, c > 0\}$ پایا بود زیرا

$$L(\bar{g}(\theta), \tilde{g}(\delta)) = L(c\theta, c\delta) = (1 - \frac{c\delta}{c\theta})^2 = (1 - \frac{\delta}{\theta})^2 = L(\theta, \delta)$$

قضیه ۱-۲: اگر خانواده F تحت تبدیلات G پایا باشد، انگاه گروه تبدیلات $\tilde{G} = \{\tilde{g} : g \in G\}$ (یعنی گروه تبدیلات بوجود آمده تحت این خانواده روی کلاس برآوردگرها) تشکیل یک گروه از تبدیلات را می‌دهند.

اثبات: تمرین (همانند قضیه ۱-۱). ■

برای مثال، در مثال ۱-۱، $\tilde{G} = \{\tilde{g}_c \mid \tilde{g}_c(\delta) = c\delta, c > 0\}$ گروه تبدیلات بوجود آمده تحت خانواده توزیع‌های پایای نمایی روی کلاس برآوردگرهای D می‌باشد.

برآوردگرهای هم‌پایا^۱

در قسمت قبل اصول پایایی را به صورت زیر معرفی کردیم:

۱- پایایی اندازه‌گیری (طبیعی یا منطقی یا تابعی): اگر $\delta(x)$ را برای برآورد $h(\theta)$ بکار

می‌بریم طبیعی به نظر می‌رسد که $\tilde{g}(\delta(X))$ را برای برآورد $\tilde{g}(h(\theta))$ بکار ببریم.

۲- پایایی ساختاری (یا ریاضی): اگر دو مسئله استنباطی دارای یک ساختار از لحاظ فضای

پارامتری، تابع چگالی و تابع زیان باشند آنگاه بایستی روش یکسانی را در دو مسئله مورد

استفاده قرار دهیم. یعنی $\delta(g(X))$ را برای برآورد $\tilde{g}(h(\theta))$ بکار ببریم.

با توجه به تعاریف بالا و اصول پایایی، برآوردگر هم‌پایا را بصورت زیر تعریف می‌کنیم.

^۱ Equivariant Estimators

تعریف ۱-۴: فرض کنید مسئله مورد نظر ما تحت گروه تبدیلات G پایا باشد (تعاریف ۱-۲ و ۱-۳) برآوردگر δ را تحت گروه تبدیلات G هم پایا گویند هرگاه برای هر $g \in G$ و هر $x \in \mathcal{X}$ داشته باشیم که

$$\delta(g(x)) = \tilde{g}(\delta(x)) \quad \blacksquare$$

برای مثال، در مثال ۱-۱ شرط فوق معادل است با $\delta(cX) = c\delta(X)$. با توجه به تعریف هم پایایی این سوال مطرح می شود که تحت یک تبدیل G کدامیک از توابع $h(\theta)$ دارای برآوردگر هم پایا هستند. اگر در تعریف ۱-۴ G را با G^- و $\mathcal{X}$ را با Θ عوض کنیم آنگاه یک تابع $h(\theta)$ طبق اصول پایایی دارای برآوردگر هم پایا است اگر در شرط زیر صدق می کند:

$$h(\bar{g}(\theta)) = \tilde{g}(h(\theta))$$

برای مثال، در مثال ۱-۱، $h(\theta) = \theta$ در شرط فوق صدق می کند زیرا:

$$h(\bar{g}(\theta)) = h(c\theta) = c\theta$$

$$\tilde{g}(h(\theta)) = c(h(\theta)) = c\theta$$

اما اگر بخواهیم $h(\theta) = \theta^r$ را برآورد کنیم آنگاه $h(\theta)$ در شرط فوق صدق نمی کند زیرا

$$h(\bar{g}(\theta)) = h(c\theta) = c^r \theta^r \quad \tilde{g}(h(\theta)) = c\theta^r$$

بنابراین، در این مثال اگر بخواهیم $h(\theta) = \theta^r$ را برآورد کنیم بایستی از ابتدا خانواده $\tilde{G}$ را به صورتی در نظر بگیریم که $\tilde{g}(h(\theta)) = c^r \theta^r$ باشد یعنی $\tilde{G} = \{\tilde{g}_c \mid \tilde{g}_c(\delta) = c^r \delta\}$.

مثال ۱-۳: فرض کنید $X \sim B(n, p)$ که در آن n معلوم و p پارامتر مجهول است و می خواهیم آن را تحت تابع زیان $L(p, \delta) = (\delta - p)^2$ برآورد کنیم. فرض کنید $T(X)$ برآوردگر پارامتر p براساس X باشد.

حال فرض کنید شخصی بخواهد احتمال شکست یعنی $q = 1 - p$ را تحت تابع زیان درجه دوم فوق برآورد نماید. این شخص از تعداد شکست ها در n آزمایش یعنی $Y = n - X \sim B(n, q)$ استفاده می نماید. بنابراین در این مسئله گروه تبدیلات زیر را داریم:

$$G = \{g_1, g_2\} \quad g_1(x) = n - x, \quad g_2(x) = x$$

چون براساس این تبدیل $Y = g_1(X) \sim B(n, 1-p)$ و $Y = g_2(X) \sim B(n, p)$ پس گروه تبدیلات G^- تولید شده توسط G عبارت است از

$$G^- = \{\bar{g}_1, \bar{g}_2\} \quad \bar{g}_1(p) = 1-p \quad \bar{g}_2(p) = p$$

حال چون می‌خواهیم $h(p) = p$ یا $h(p) = 1-p$ را برآورد کنیم پس بایستی

$$\tilde{g}(1-p) = \tilde{g}(h(p)) = h(\bar{g}_1(p)) = h(1-p) = p$$

$$\tilde{g}_2(p) = \tilde{g}_2(h(p)) = h(\bar{g}_2(p)) = h(p) = p$$

بنابراین گروه تبدیلات زیر روی D بوجود می‌آید

$$\tilde{G} = \{\tilde{g}_1, \tilde{g}_2\} \quad \tilde{g}_1(\delta) = 1-\delta, \quad \tilde{g}_2(\delta) = \delta$$

همچنین توجه کنید که

$$L(\bar{g}_1(p), \tilde{g}_1(\delta)) = L(1-p, 1-\delta) = (\delta - p)^2 = L(p, \delta)$$

یعنی تابع زیان درجه دم تحت این تبدیل پایا است.

حال طبق تعریف ۱-۴ برآوردگر δ در این مسئله هم‌پایا است اگر

$$\delta(g(X)) = \tilde{g}(\delta(X)) \Leftrightarrow \delta(n-X) = 1-\delta(X)$$

از جمله برآوردگرهای پایا در این مسئله $T_1(x) = \frac{x}{n}$ و $T_2(x) = \alpha(\frac{x}{n}) + (1-\alpha)(0.5)$ می‌باشد.

اگر برآوردگرهای پایا بفرم بالا را در نظر بگیریم در این صورت تعیین این برآوردگر در نقاط

$$T(\cdot), T(1), \dots, T\left(\left[\frac{n}{2}\right]\right)$$

$$T(n) = -T(\cdot), \quad T(n-1) = -T(1), \dots$$

در حالی که اگر یک برآوردگر عادی را بکار ببریم بایستی این برآوردگر را در نقاط

$T(\cdot), T(1), \dots, T(n)$ تعیین کنیم و بنابراین برآوردگرهای پایا نیز روشی برای خلاصه کردن داده‌ها

ارایه می‌دهند. ■

بهترین برآوردگر هم پایا

حال این سوال مطرح می‌شود که در بین برآوردگرهای هم پایا کدام برآوردگر بهتر است. جواب این سوال این است که برآوردگر هم پایایی که دارای کمترین مخاطره باشد بهترین برآوردگر هم پایا می‌باشد. برای بدست آوردن بهترین برآوردگر هم پایا از تعاریف و قضایای زیر استفاده می‌کنیم.

تعریف ۵-۱: دو نقطه θ_1, θ_2 در Θ را معادل^۱ گویند اگر $\bar{g} \in G^-$ موجود باشد که $\theta_2 = \bar{g}(\theta_1)$.
 یک مدار^۲ در Θ عبارت از یک کلاس (مجموعه) از نقاط معادل در Θ است. بنابراین مدار θ در Θ عبارت از مجموعه زیر است: $\Theta(\theta) = \{\bar{g}(\theta) \mid \bar{g} \in G^-\}$ ■

معادل بودن در تعریف ۵-۱ یک رابطه هم‌ارزی را بدست می‌دهد. بنابراین تعریف فوق Θ را به کلاس‌های هم‌ارزی افراز می‌کند. هر یک از این کلاس‌ها را یک مدار (تحت G^-) می‌نامیم.

تعریف ۶-۱: یک گروه G^- از تبدیلات روی Θ را انتقالی یا گذرا^۳ گویند، هرگاه Θ دارای تنها یک مدار باشد و یا معادلا اگر برای هر دو نقطه $\theta_1, \theta_2 \in \Theta$ یک $\bar{g} \in G^-$ موجود باشد بطوریکه $\theta_2 = \bar{g}(\theta_1)$. ■

برای مثال، در مثال ۱-۱ G^- تنها دارای یک مدار است $(\theta_2 = \frac{\theta_1}{\theta_1} \theta_1 = c\theta_1)$ پس یک گروه انتقالی است ولی در مثال ۳-۱ چون $\bar{g}(p) = 1-p$ پس مدارها همان زوج‌های مرتب $(p, 1-p)$ هستند و در نتیجه G^- انتقالی نیست. در این گروه نقاط $p_1 = \frac{1}{3}, p_2 = \frac{2}{3}$ را می‌توان به هم انتقال داد ولی نقاط $p_1 = \frac{1}{3}, p_2 = \frac{2}{3}$ را نمی‌توان به هم انتقال داد.

قضیه ۳-۱ (مهم): تابع مخاطره هر برآوردگر هم‌پایا روی مدارهای Θ مقداری ثابت دارد و یا معادلا

$$R(\theta, \delta) = R(\bar{g}(\theta), \delta) \quad \forall \theta \in \Theta, \bar{g} \in G^-$$

^۱ Equivalent

^۲ Orbit

^۳ Transitive

اثبات: $R(\theta, \delta) = E_{\theta}[L(\theta, \delta(X))]$

$$\begin{aligned}
 &= E_{\theta}[L(\bar{g}(\theta), \tilde{g}(\delta(X)))] && \text{(پایایی تابع زیان)} \\
 &= E_{\theta}[L(\bar{g}(\theta), \delta(g(X)))] && \text{(هم پایایی برآوردگر } \delta \text{)} \\
 &= E_{\bar{g}(\theta)}[L(\bar{g}(\theta), \delta(X)))] && \text{(پایایی توزیع، رابطه (۶))} \\
 &= R(\bar{g}(\theta), \delta) && \blacksquare
 \end{aligned}$$

نتیجه ۱-۱: اگر G^- یک گروه انتقالی از تبدیلات روی فضای پارامتری Θ باشد آنگاه تابع مخاطره هر برآوردگر هم پایا ثابت است و به θ بستگی ندارد. $\blacksquare$

توجه: با توجه به نتیجه فوق اگر G^- یک گروه انتقالی باشد آنگاه تابع مخاطره برآوردگرهای هم پایا نسبت به θ ثابت است و در نتیجه بهترین برآوردگر هم پایا (MRE) آن برآوردگری است که این تابع مخاطره ثابت را مینیمم می کند.

برای مثال، در مثال ۱-۱ دیدیم که تابع مخاطره برآوردگر هم پایای $\delta(X) = kX$ برابر $R(\theta, \delta) = 2k^2 - 2k + 1$ بود که به θ بستگی نداشت و مینیمم آن به ازای $k = \frac{1}{2}$ بدست آمد، پس برآوردگر MRE عبارت از $\delta^*(X) = \frac{1}{2}X$ بود.

تعریف ۱-۷: فرض کنید G گروهی از تبدیلات روی فضای $\mathcal{X}$ باشد. تابع $T(X)$ را روی $\mathcal{X}$ نسبت به G پایا^۱ گویند هرگاه برای هر $x \in \mathcal{X}$ و $g \in G$ داشته باشیم که $T(g(x)) = T(x)$. تابع $T(x)$ را پایای ماکزیمال^۲ نسبت به G گویند هرگاه T پایا باشد و در شرط زیر صدق کند

$$\blacksquare \quad T(x_1) = T(x_2) \Rightarrow x_1 = g(x_2) \text{ for some } g \in G$$

مفاهیم فوق برای درک این مطلب که $\mathcal{X}$ می تواند به مدارهایی از نقاط که تحت G معادل هستند، مفید می باشد. یک تابع پایا تابعی است که روی هر مدار مقداری ثابت دارد. $(x' \sim x \Rightarrow T(x') = T(x))$ و یک تابع پایای ماکزیمال تابعی است که روی هر مدار مقداری ثابت

^۱ Invariant

^۲ Maximal Invariant

دارد و مقادیر متفاوتی را به هر یک از مدارهای متفاوت می‌دهد. توجه کنید که اگر G یک گروه انتقالی روی $\mathcal{X}$ باشد آنگاه $\mathcal{X}$ خود یک مدار است و بنابراین تنها تابع پایایی ماکزیمال روی $\mathcal{X}$ تابع ثابت است.

بدیهی است که اگر T یک آماره پایایی ماکزیمال باشد، هر تابع یک به یک از آن نیز یک آماره پایایی ماکزیمال است. پس آماره پایایی ماکزیمال یکتا نیست ولی همه آنها توابعی یک به یک از یکدیگر هستند.

مثال ۴-۱: فرض کنید $\mathcal{X} = \mathbb{R}^n$ و G گروه تبدیلات مکانی باشد.

$$G = \{g_c : g_c(x_1, x_2, \dots, x_n) = (x_1 + c, \dots, x_n + c), c \in \mathbb{R}\}$$

در این صورت $T(\underline{x}) = (x_1 - x_n, \dots, x_{n-1} - x_n)$ پایایی ماکزیمال است. چون $(x_i + c) - (x_n + c) = x_i - x_n$ پس $T(g_c(\underline{x})) = T(\underline{x})$ پس T پایا است. حال اگر $T(\underline{x}) = T(\underline{x}')$ آنگاه $(x_i - x_n) = (x_i' - x_n)$ و در نتیجه با $c = x_n - x_n'$ داریم که $g_c(\underline{x}') = \underline{x}$ پس T پایایی ماکزیمال است. ■

مثال ۵-۱: فرض کنید $\mathcal{X} = \mathbb{R}^n$ و G گروه تبدیلات مقیاسی باشد.

$$G = \{g_c : g_c(x_1, \dots, x_n) = (cx_1, \dots, cx_n), c > 0\}$$

$$T(\underline{X}) = \begin{cases} \cdot & \text{if } Z = \cdot \\ \left(\frac{X_1}{Z}, \dots, \frac{X_n}{Z}\right) & \text{if } Z \neq \cdot \end{cases} \quad \text{آنگاه } Z = \sum_{i=1}^n X_i^2$$

مثال ۶-۱: فرض کنید $\mathcal{X} = \mathbb{R}^n$ و G گروه تبدیلات مکانی-مقیاسی باشد.

$$G = \{g_{b,c} \mid g_{b,c}(x_1, \dots, x_n) = (cx_1 + b, \dots, cx_n + b), c > 0, b \in \mathbb{R}\}$$

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 \quad \text{و} \quad \bar{x} = \frac{1}{n} \sum_{i=1}^n x_i \quad \text{اگر}$$

$$T(\underline{x}) = \begin{cases} \cdot & s = \cdot \\ \left(\frac{x_1 - \bar{x}}{s}, \dots, \frac{x_n - \bar{x}}{s}\right) & s \neq \cdot \end{cases} \quad \text{آنگاه}$$

توجه: پایایی و پایایی ماکزیمال را روی Θ بطور مشابه می توان تعریف کرد.

قضیه ۱-۴: فرض کنید G گروهی از تبدیلات روی فضای $\mathcal{X}$ باشد و $T(X)$ یک پایای ماکزیمال باشد. در این صورت تابع $h(X)$ نسبت به G پایا است اگر و فقط اگر h تابعی از T باشد.

اثبات:

(کفایت): فرض کنید $h(x) = s(T(x))$ در این صورت

$$h(g(x)) = s(T(g(x))) = s(T(x)) = h(x) \text{ پس } h \text{ پایا است.}$$

(لزوم): فرض کنید h پایا باشد و $T(x_1) = T(x_2)$ در این صورت چون T پایای ماکزیمال است پس $x_1 = g(x_2)$ و در نتیجه $h(x_1) = h(g(x_2)) = h(x_2)$ یعنی h تابعی از T است. ■

قضیه ۱-۵: اگر h یک آماره پایا باشد آنگاه توزیع آن روی هر مدار Θ مقداری ثابت است.

اثبات: فرض کنید θ_1, θ_2 در یک مدار Θ باشند در این صورت $\bar{g} \in G^-$ وجود دارد که $\theta_2 = \bar{g}(\theta_1)$ و بنابراین

$$P_{\theta_2}(h(X) \in A) = P_{\bar{g}(\theta_1)}(h(X) \in A) \stackrel{(\Delta)}{=} P_{\theta_1}(h(g(x)) \in A) \stackrel{(*)}{=} P_{\theta_1}(h(X) \in A)$$

یعنی توزیع روی هر مدار Θ ثابت است. توجه داشته باشید که تساوی $(*)$ از پایایی h حاصل می گردد. ■

نتیجه ۱-۲: اگر G^- تقارنی باشد آنگاه Θ تنها یک مدار دارد و در نتیجه توزیع هر آماره پایای h روی آن به θ بستگی ندارد، یعنی h یک آماره فرعی^۱ است.

قضیه زیر حالت کلی تر قضیه ۱-۵ را بیان می کند.

قضیه ۱-۶: فرض کنید یک مسئله تحت تبدیلات G پایا باشد و $v(\theta)$ یک تابع پایای ماکزیمال روی Θ نسبت به G^- باشد. اگر $h(X)$ یک تابع پایا تحت G باشد آنگاه توزیع $h(X)$ تنها به $v(\theta)$ بستگی دارد.

اثبات: برای هر $A \subset \mathcal{X}$ داریم که

^۱ Ancillary

$$P_{\bar{g}(\theta)}(h(X) \in A) \stackrel{(\Delta)}{=} P_{\theta}(h(g(X)) \in A) \stackrel{(*)}{=} P_{\theta}(h(X) \in A)$$

توجه داشته باشید که تساوی $(*)$ از پایایی h حاصل می‌گردد. یعنی $S(\theta) = P_{\theta}(h(X) \in A)$ یک تابع پایا روی Θ است $(S(\bar{g}(\theta)) = S(\theta))$ بنابراین طبق قضیه ۱-۴ $S(\theta)$ تابعی از پایای ماکزیمال $\nu(\theta)$ است. چون این مطلب برای هر A برقرار است پس توزیع $h(X)$ تنها به $\nu(\theta)$ بستگی دارد.

■

توسط قضیه زیر برآوردهای MRE را بدست می‌آوریم.

قضیه ۱-۷: فرض کنید مسئله مورد نظر تحت تبدیلات G پایا باشد (تعاریف ۱-۲ و ۱-۳) و G^- یک گروه انتقالی روی Θ باشد و $Y = T(X)$ یک تابع پایای ماکزیمال روی $\mathcal{X}$ باشد در این صورت هر برآوردهای هم‌پایای $\delta(X)$ با مخاطره متناهی که امید شرطی $E_{\bar{e}}[L(\bar{e}, \delta(X)) | Y = y]$ را مینیمم کند برآوردهای MRE پارامتر θ است.

اثبات: برآوردهای MRE از مینیمم کردن $R(\theta, \delta) = E_{\theta}[L(\theta, \delta(X))]$ بدست می‌آید. اما طبق نتیجه ۱-۱ تابع مخاطره هر برآوردهای هم‌پایا نسبت به θ ثابت است (چون G^- انتقالی است). بنابراین کافی است این مخاطره را در نقطه $\theta = \bar{e}$ مینیمم کنیم. بنابراین کافی است مخاطره زیر را نسبت به δ مینیمم کنیم

$$\begin{aligned} R(\bar{e}, \delta) &= E_{\bar{e}}[L(\bar{e}, \delta(X))] = E_{\bar{e}}\{E_{\bar{e}}[L(\bar{e}, \delta(X)) | Y]\} \\ &= \int E_{\bar{e}}[L(\bar{e}, \delta(X)) | Y = y] f_Y(y) d(\mu_Y) \end{aligned}$$

این انتگرال موقعی مینیمم می‌شود که $E_{\bar{e}}[L(\bar{e}, \delta(X)) | Y = y]$ مینیمم شود و در نتیجه قضیه ثابت است. ■

توجه: اگر در مسئله‌ای آماره بسنده $S(X)$ وجود داشته باشد که پایا باشد یعنی $S(g(X)) = \tilde{g}(S(X))$ آنگاه می‌توان برآوردهای هم‌پایای $\delta(X)$ را همین آماره بسنده $S(X)$ در نظر گرفت که با توجه به قضیه باسو و قضیه ۱-۵ این آماره هم‌پایا از تابع پایای ماکزیمال $Y = T(X)$ مستقل است.

بخش ۲: برآوردگر MRE پارامتر مکان^۱

فرض کنید $\underline{X} = (X_1, X_2, \dots, X_n)$ دارای توزیع احتمال توام

$$f(\underline{x} - \theta) = f(x_1 - \theta, \dots, x_n - \theta)$$

باشد که در آن f معلوم و θ پارامتر مکان نامعلوم است. اگر گروه تبدیلات زیر را در نظر بگیریم.

$$G = \{g_a : g_a(x_1, \dots, x_n) = (x_1 + a, \dots, x_n + a) = \underline{x} + a, a \in R\}$$

آنگاه خانواده توزیع‌های فوق نسبت به این تبدیل پایا است، هرگاه گروه تبدیلات G^- بصورت زیر

$$G^- = \{\bar{g}_a : \bar{g}_a(\theta) = \theta + a, a \in R\} \quad \text{باشد}$$

که یک گروه انتقالی از تبدیلات است. در برآورد پارامتر $h(\theta) = \theta$ تحت تبدیل G ، $h(\theta)$ دارای برآوردگر هم‌پایا است اگر

$$h(\bar{g}(\theta)) = \tilde{g}(h(\theta)) \Leftrightarrow \theta + a = \tilde{g}(\theta)$$

بنابراین گروه تبدیلات G^- بصورت زیر است $\tilde{G} = \{\tilde{g}_a : \tilde{g}_a(\delta) = \delta + a, a \in R\}$

در نتیجه تحت گروه G تابع زیان $L(\theta, \delta)$ پایا است اگر و فقط اگر

$$L(\theta, \delta) = L(\bar{g}(\theta), \tilde{g}(\delta)) \Leftrightarrow L(\theta, \delta) = L(\theta + a, \delta + a) \stackrel{?}{\Leftrightarrow} L(\theta, \delta) = \rho(\delta - \theta)$$

(اثبات رابطه سوم:

$$\Rightarrow) a = -\theta \Rightarrow L(\theta, \delta) = L(\cdot, \delta - \theta) = \rho(\delta - \theta)$$

$$(\Leftarrow) L(\theta + a, \delta + a) = \rho(\delta + a - \theta - a) = \rho(\delta - \theta) = L(\theta, \delta)$$

بنابراین اگر تابع زیان بفرم $L(\theta, \delta) = \rho(\delta - \theta)$ باشد آنگاه مسئله، تحت گروه تبدیلات G پایا

است. برای مثال $L(\theta, \delta) = |\delta - \theta|$ یا $L(\theta, \delta) = (\delta - \theta)^2$ از این نوع تابع زیان می باشند.

همچنین برآوردگرهای هم‌پایا در این حالت بصورت زیر می‌باشند:

$$\delta(g(\underline{X})) = \tilde{g}(\delta(\underline{X})) \Leftrightarrow \delta(\underline{X} + a) = \delta(\underline{X}) + a$$

از جمله برآوردگرهای هم‌پایا در این حالت $\bar{X}$ و $\frac{X_{(1)} + X_{(n)}}{2}$ می‌باشند.

برای بدست آوردن برآوردگرهای MRE در این حالت از قضیه زیر استفاده می‌کنیم.

^۱ MRE Estimator of Location Parameter

قضیه ۱-۲: اگر δ هر برآوردگر هم پایایی θ باشد آنگاه شرط لازم و کافی برای آنکه برآوردگر δ برای θ هم پایا باشد آن است که

$$\delta(\underline{x}) = \delta(\underline{x}) + U(\underline{x}) \quad (۱)$$

که در آن $U(\underline{x})$ هر تابع پایا نسبت به G است یعنی:

$$U(\underline{x} + a) = U(\underline{x}) \quad (۲)$$

اثبات: فرض کنید روابط (۱) و (۲) برقرار باشند در این صورت:

$$\delta(\underline{x} + a) = \delta(\underline{x} + a) + U(\underline{x} + a) = \{\delta(\underline{x}) + a\} + U(\underline{x}) = \delta(\underline{x} + a)$$

پس $\delta(\underline{x})$ یک برآوردگر هم پایا است. حال فرض کنید $\delta(\underline{x})$ هم پایا باشد و قرار می‌دهیم

$$U(\underline{x}) = \delta(\underline{x}) - \delta(\underline{x})$$

در این صورت رابطه (۱) برقرار است و همچنین

$$U(\underline{x} + a) = \delta(\underline{x} + a) - \delta(\underline{x} + a) = \delta(\underline{x}) + a - \delta(\underline{x}) - a = \delta(\underline{x}) - \delta(\underline{x}) = U(\underline{x})$$

پس رابطه (۲) برقرار است. ■

اگر قرار دهیم $\underline{Y} = (Y_1, \dots, Y_{n-1})$ که $Y_i = X_i - X_n$, $i = 1, \dots, n-1$ و $Y = c$ اگر $n = 1$ باشد، آنگاه $\underline{Y}$ یک تابع پایای ماکزیمال نسبت به G است و طبق قضیه ۱-۴ تابع $U(\underline{x})$ پایا است اگر و فقط اگر تابعی از تابع پایای ماکزیمال $\underline{Y}$ باشد بنابراین از قضیه ۱-۲ نتیجه می‌شود که

نتیجه ۱: اگر δ هر برآوردگر هم پایایی θ باشد آنگاه شرط لازم و کافی برای اینکه برآوردگر δ برای θ هم پایا باشد این است که تابع v از $n-1$ متغیر $y_i = x_i - x_n$, $i = 1, \dots, n-1$ وجود داشته باشد بطوری که $\delta(\underline{x}) = \delta(\underline{x}) - v(\underline{y})$. ■

حال چون گروه G^- انتقالی است و $\underline{Y}$ یک تابع پایای ماکزیمال است در این صورت طبق ۱-۷ برآوردگر هم پایای $\delta(\underline{X}) = \delta(\underline{X}) - v(\underline{Y})$ که امید شرطی $E_{\theta=}[\rho(\delta(\underline{X}) - v(\underline{Y})) | \underline{y}]$ را مینیمم کند برآوردگر MRE پارامتر θ است. حال اگر $v^*(\underline{y})$ این امید شرطی را مینیمم کند آنگاه $\delta^*(\underline{x}) = \delta(\underline{x}) - v^*(\underline{y})$ برآوردگر MRE پارامتر θ است.

چگالی $f(x_i - \theta)$ با $\theta = E(X_i)$ باشد و $r_n(F)$ مخاطره برآوردگر MRE پارامتر θ تحت تابع زیان درجه دوم باشد. در این صورت $r_n(F)$ موقعی مقدار ماکزیمم خود را روی F بدست می‌آورد که F توزیع نرمال باشد.

اثبات: در حالت توزیع نرمال برآوردگر MRE پارامتر θ برابر $\bar{X}$ با مخاطره $r_n(F) = \frac{1}{n}$ است. چون مخاطره $\bar{X}$ در هر توزیعی برابر $\frac{1}{n}$ است بنابراین مخاطره برآوردگر MRE در هر توزیع F بایستی کمتر یا مساوی $\frac{1}{n}$ باشد (MRE مخاطره را مینم می‌کند) بنابراین توزیع نرمال دارای ماکزیمم مقدار مخاطره برآوردگر MRE در بین کلیه توزیع‌های F است. ■

مثال ۲-۳: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع نمایی دو پارامتری $E(\theta, b)$ باشند که b مقداری معلوم است. برآوردگر MRE پارامتر θ را تحت توابع زیان درجه دوم و قدر مطلق خطا بدست آورید.

$$f_{\theta}(x) = \frac{1}{b} e^{-\frac{1}{b}(x-\theta)} I_{(\theta, +\infty)}(x) = \frac{1}{b} e^{-\frac{1}{b}(x-\theta)} I_{(+\infty)}(x - \theta) = f(x - \theta) \quad \text{حل:}$$

قرار می‌دهیم $\delta(X_{(1)}) = X_{(1)}$ بنابراین $\delta(X_{(1)})$ آماره بسنده کامل و همچنین برآوردگر هم‌پایای مکان θ است و طبق قضیه باسو از آماره فرعی $Y_{(1)}$ (پایای ماکزیمال) مستقل است. بنابراین الف- تحت تابع زیان درجه دوم

$$v^*(y) = E[X_{(1)} | y] = E.(X_{(1)}) = \int_{-\infty}^{+\infty} x \left(\frac{n}{b} e^{-\frac{n}{b}x} \right) dx = \frac{b}{n}$$

پس برآوردگر MRE پارامتر θ عبارت است از $\delta^*(X_{(1)}) = X_{(1)} - \frac{b}{n}$ که برآوردگر UMVU نیز هست.

ب- تحت تابع زیان قدر مطلق خطا $v^*(y) = \text{med}.[X_{(1)} | y] = \text{med}.(X_{(1)})$

$$\int_{-\infty}^v \frac{n}{b} e^{-\frac{n}{b}x} dx = \frac{1}{2} \Rightarrow -e^{-\frac{n}{b}x} \Big|_{-\infty}^v = \frac{1}{2} \Rightarrow 1 - e^{-\frac{n}{b}v} = \frac{1}{2} \Rightarrow v^* = \frac{b}{n} \ln 2$$

بنابراین برآوردگر MRE پارامتر θ عبارت است از $\delta^*(X) = X_{(1)} - \frac{b}{n} \ln 2$ که این برآوردگر نااریب نیست. ■

برآوردگر پیتمن^۱

اگر تابع زیان درجه دوم باشد، آنگاه برآوردگر MRE مکان را به فرم راحت‌تری می‌توان بدست آورد. قضیه ۲-۳: تحت گروه تبدیلات مکانی G و در خانواده توزیع‌های مکان و تحت تابع زیان درجه دوم، برآوردگر MRE مکان عبارت است از:

$$\delta^*(\tilde{X}) = \delta(\tilde{X}) - E[\delta(\tilde{X}) | y] = \frac{\int_{-\infty}^{+\infty} u f(x_1 - u, \dots, x_n - u) du}{\int_{-\infty}^{+\infty} f(x_1 - u, \dots, x_n - u) du}$$

این فرم را بعنوان برآوردگر پیتمن پارامتر θ می‌شناسند.

اثبات: فرض کنید $\delta(\tilde{X}) = X_n$ برآوردگر هم‌پایای θ باشد. بنابراین

$$v^*(y) = E(X_n | y)$$

حال تبدیل زیر را در نظر می‌گیریم:

$$\begin{cases} Y_i = X_i - X_n, & i = 1, \dots, n-1 \\ W = X_n \end{cases}$$

در این صورت $J = 1$ و $f_{Y,W}(y_1, \dots, y_{n-1}, w) = f(y_1 + w, \dots, y_{n-1} + w, w)$

و بنابراین

$$\begin{aligned} f_{X_n | \tilde{X}}(w | y) &= f_{W | \tilde{X}}(w | y) = \frac{f_{Y,W}(y_1, \dots, y_{n-1}, w)}{\int f_{Y,W}(y_1, \dots, y_{n-1}, w) dw} \\ &= \frac{f(y_1 + w, \dots, y_{n-1} + w, w)}{\int f(y_1 + w, \dots, y_{n-1} + w, w) dw} \\ \Rightarrow v^*(y) &= E(X_n | Y) = E(W | Y) = \frac{\int w f(y_1 + w, \dots, y_{n-1} + w, w) dw}{\int f(y_1 + w, \dots, y_{n-1} + w, w) dw} \end{aligned}$$

^۱ Pitman Estimator

$$\begin{aligned}
&= \frac{\int w f(x_1 - x_n + w, \dots, x_{n-1} - x_n + w, w) dw}{\int f(x_1 - x_n + w, \dots, x_{n-1} - x_n + w, w) dw} \\
&= \frac{\int_{u=x_n-w}^u f(x_1 - u, \dots, x_{n-1} - u, x_n - u) du}{\int f(x_1 - u, \dots, x_{n-1} - u, x_n - u) du} \\
\therefore \delta^*(x_{\sim}) &= x_n - E(X_n | Y_{\sim}) = \frac{\int u f(x_1 - u, \dots, x_n - u) du}{\int f(x_1 - u, \dots, x_n - u) du} \quad \blacksquare
\end{aligned}$$

مثال ۲-۴: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $U(\theta - \frac{1}{r}, \theta + \frac{1}{r})$ باشد برآوردگر

MRE پارامتر θ را بدست آورید.

حل:

$$\begin{aligned}
f_{X_{\sim}}(x_{\sim}) &= \begin{cases} \prod_{i=1}^n I_{(\theta - \frac{1}{r}, \theta + \frac{1}{r})}(x_i) = \prod_{i=1}^n I_{(-\frac{1}{r}, \frac{1}{r})}(x_i - \theta) = f(x_1 - \theta, \dots, x_n - \theta) \\ I_{(x_{(n)} - \frac{1}{r}, x_{(1)} + \frac{1}{r})}(\theta) \end{cases} \\
\therefore \delta^*(x_{\sim}) &= \frac{\int u f(x_{\sim} - u) du}{\int f(x_{\sim} - u) du} = \frac{\int u I_{(x_{(n)} - \frac{1}{r}, x_{(1)} + \frac{1}{r})} du}{\int I_{(x_{(n)} - \frac{1}{r}, x_{(1)} + \frac{1}{r})} du} = \dots = \frac{x_{(1)} + x_{(n)}}{2}
\end{aligned}$$

پس $\delta^*(X_{\sim}) = \frac{X_{(1)} + X_{(n)}}{2}$ برآوردگر MRE برای θ است ولی UMVUE نیست. $\blacksquare$

مثال ۲-۵: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $E(\theta, b)$ باشند که b معلوم است

برآوردگر MRE پارامتر θ را بدست آورید.

حل:

$$f_{X_{\sim}}(x_{\sim}) = \begin{cases} \prod_{i=1}^n \frac{1}{b} e^{-(x_i - \theta)/b} I_{(., +\infty)}(x_i - \theta) = f(x_1 - \theta, \dots, x_n - \theta) \\ \frac{1}{b^n} e^{(-n\bar{x} + n\theta)/b} I_{(-\infty, x_{(1)})}(\theta) \end{cases}$$

$$\begin{aligned} \therefore \delta^*(x) &= \frac{\int u f(x-u) du}{\int f(x-u) du} = \frac{\int_{-\infty}^{x_{(1)}} \frac{u}{b^n} e^{(-n\bar{x}+nu)/b} du}{\int_{-\infty}^{x_{(1)}} \frac{1}{b^n} e^{(-n\bar{x}+nu)/b} du} = \frac{\frac{b}{n} (u - \frac{b}{n}) e^{\frac{nu}{b}} \Big|_{-\infty}^{x_{(1)}}}{\frac{b}{n} e^{\frac{-nu}{b}} \Big|_{-\infty}^{x_{(1)}}} \\ &= x_{(1)} - \frac{b}{n} \end{aligned}$$

■ که همان برآوردگر بدست آمده در مثال ۲-۳-الف می باشد.

قضیه ۲-۴: برآوردگر پیتمن پارامتر مکان تابعی از آماره بسنده است.

اثبات: فرض کنید $\underline{S} = (S_1, \dots, S_k)$ یک آماره بسنده توام برای θ باشد. بنابراین طبق قضیه تجزیه به عوامل داریم که $f(x_1 - \theta, \dots, x_n - \theta) = g(\underline{S}, \theta) h(x_1, \dots, x_n)$ بنابراین برآوردگر پیتمن پارامتر θ عبارت است از

$$\blacksquare \quad \delta^*(x) = \frac{\int u f(x-u) du}{\int f(x-u) du} = \frac{\int u g(\underline{S}, u) h(x) du}{\int g(\underline{S}, u) h(x) du} = \frac{\int u g(\underline{S}, u) du}{\int g(\underline{S}, u) du} = k(\underline{S})$$

با توجه به مثال‌های بالا دیده می‌شود که برای برآورد پارامتر مکان ارتباطی بین برآوردگرهای MRE و ناریب و UMVUE وجود دارد. این ارتباط در قضیه زیر آورده شده است:

قضیه ۲-۵: فرض کنید تابع زیان درجه دوم باشد در این صورت برای برآورد پارامتر مکان

الف- اگر $\delta(X)$ هر برآورد هم‌پایا با اریب ثابت b باشد آنگاه $\delta(X) - b$ یک برآوردگر هم‌پایا و ناریب می‌باشد و دارای مخاطره کمتری نسبت به $\delta(X)$ است.

ب- برآوردگر یکتای MRE یک برآوردگر ناریب است.

ج- اگر برآوردگر UMVU موجود باشد و هم‌پایا باشد آنگاه یک برآوردگر MRE است.

اثبات: (الف) هم‌پایا بودن و ناریب بودن $\delta(X) - b$ واضح است.

$$\begin{aligned} R(\theta, \delta - b) &= E_{\theta}[(\delta(X) - b - \theta)^2] \\ &= E_{\theta}\left[(\delta(X) - \theta)^2\right] + b^2 - 2bE_{\theta}[\delta(X) - \theta] \\ &= R(\theta, \delta) - b^2 < R(\theta, \delta) \end{aligned}$$

ب- فرض کنید $\delta(X)$ یک برآوردگر یکتای MRE باشد ولی نااریب نباشد و دارای اریب b باشد در این صورت طبق قسمت (الف) $\delta(X) - b$ یک برآوردگر هم‌پایا با مخاطره کمتر نسبت به $\delta(X)$ است که این با یکتا بودن برآوردگر MRE، $\delta(X)$ متناقض است. پس $\delta(X)$ نااریب می‌باشد.

ج- اگر $\delta(X)$ یک برآوردگر هم‌پایا و UMVU باشد آنگاه این برآوردگر مقدار مخاطره را مینیمم می‌کند. پس یک برآوردگر MRE است. ■

توجه: در مثال‌های فوق دیدیم که اگر تابع زیان درجه دوم نباشد آنگاه ممکن است برآوردگر MRE یک برآوردگر نااریب نباشد. مفهوم معمول نااریبی این است که

$$E_{\theta}[\delta(X)] = \theta \quad \forall \theta \in \Theta$$

و اگر تابع زیان درجه دوم باشد آنگاه $R(\theta, \delta) = E_{\theta}[(\delta(X) - \theta)^2]$ موقعی به ازای هر $\theta \in \Theta$ مینیمم مقدار خود را بدست می‌آورد که $\theta = E_{\theta}[\delta(X)]$ باشد بنابراین

$$E_{\theta}[(\delta(X) - \theta)^2] \leq E_{\theta}[(\delta(X) - \theta')^2] \quad \forall \theta' \neq \theta$$

یعنی متوسط فاصله $\delta(X)$ با مقدار واقعی پارامتر، یعنی $E_{\theta}[(\delta(X) - \theta)^2]$ کمتر است از متوسط فاصله $\delta(X)$ از هر مقدار دیگر پارامتر. این مطلب راه گشای تعریفی برای نااریبی تحت توابع زیان دیگر است. ■

تعریف ۱-۲: (مخاطره نااریب)^۱: فرض کنید در یک مسئله برآوردیابی تابع زیان $L(\theta, \delta(X))$ باشد. برآوردگر $\delta(X)$ را مخاطره نااریب گویند هرگاه

$$E_{\theta}[L(\theta, \delta(X))] \leq E_{\theta}[L(\theta', \delta(X))] \quad \forall \theta' \neq \theta \quad \blacksquare$$

مثال ۲-۶: (نااریبی میانه): اگر تابع زیان قدرمطلق خطا باشد آنگاه

$$\begin{aligned} g(\theta) &= E_{\theta}[L(\theta, \delta(X))] = E[|\delta(X) - \theta|] = \int_{-\infty}^{+\infty} |\delta - \theta| f(\delta) d\delta \\ &= \int_{-\infty}^{\theta} (\theta - \delta) f(\delta) d\delta + \int_{\theta}^{+\infty} (\delta - \theta) f(\delta) d\delta \\ g'(\theta) &= 0 \quad \Rightarrow \quad \int_{-\infty}^{\theta} f(\delta) d\delta = \int_{\theta}^{+\infty} f(\delta) d\delta \Rightarrow \theta = med_{\theta}(\delta(X)) \quad \blacksquare \end{aligned}$$

^۱ Risk-unbiased

با توجه به قضیه ۲-۵-ب می‌خواهیم نشان دهیم که برآوردگر MRE یک برآوردگر مخاطره نااریب است. برای این منظور از لم زیر استفاده می‌کنیم.

لم ۲-۱: اگر $\tilde{G}$ یک گروه جابجایی (آبلی) باشد و δ یک برآوردگر هم پایا تحت گروه G باشد آنگاه $\tilde{G}(\delta)$ برای هر $\tilde{g} \in \tilde{G}$ یک برآوردگر هم پایا است.

اثبات: برای هر $g_* \in G$ داریم که

$$\tilde{g}(\delta(g_*(\underline{X}))) \stackrel{(1)}{=} \tilde{g}(\tilde{g}_*(\delta(\underline{X}))) \stackrel{(2)}{=} \tilde{g}_*(\tilde{g}(\delta(\underline{X})))$$

که در آن (۱) از هم پایا بودن δ و (۲) از آبلی بودن $\tilde{G}$ بدست می‌آید. یعنی $\tilde{g}(\delta)$ یک برآوردگر هم پایا است. ■

قضیه ۲-۶: اگر G^- یک گروه انتقالی و $\tilde{G}$ یک گروه جابجایی باشد آنگاه برآوردگر MRE یک برآوردگر مخاطره نااریب است.

اثبات: فرض کنید δ برآوردگر MRE باشد و $\theta', \theta \in G^-$. چون G^- انتقالی است پس $\bar{g} \in G^-$ موجود است بطوری که $\theta = \bar{g}(\theta')$ و بنابراین

$$E_{\theta}[L(\theta', \delta(\underline{X}))] = E_{\theta}[L(\bar{g}^{-1}(\theta), \delta(\underline{X}))] \stackrel{(*)}{=} E[L(\theta, \tilde{g}(\delta(\underline{X}))) \stackrel{(**)}{\geq} E_{\theta}[L(\theta, \delta(\underline{X}))]$$

که در آن تساوی (*) از پایایی تابع زیان و نامساوی (**) از لم ۲-۱ و MRE بودن δ حاصل می‌گردند. پس δ برآوردگر مخاطره نااریب است. ■

نتیجه ۲-۳: تحت گروه تبدیلات مکان G و تحت تابع زیان درجه دوم، برآوردگر MRE پارامتر θ مخاطره-نااریب (نااریب - میانگین) است. ■

بخش ۳: برآوردگر MRE پارامتر مقیاس^۱

فرض کنید $\underline{X} = (X_1, \dots, X_n)$ دارای توزیع احتمال توام

$$\frac{1}{\sigma^n} f\left(\frac{\underline{x}}{\sigma}\right) = \frac{1}{\sigma^n} f\left(\frac{x_1}{\sigma}, \dots, \frac{x_n}{\sigma}\right) \quad \sigma > 0.$$

^۱ MRE Estimator of scale-parameter

باشد که در آن f معلوم و σ پارامتر مقیاس نامعلوم است. اگر گروه تبدیلات زیر را در نظر بگیریم

$$G = \{g_c : g_c(x_1, x_2, \dots, x_n) = (cx_1, \dots, cx_n) = cx, c > 0\}$$

آنگاه خانواده توزیع‌های فوق نسبت به این تبدیل پایا است هرگاه گروه تبدیلات G^- بصورت زیر

$$G^- = \{\bar{g}_c : \bar{g}_c(\sigma) = c\sigma, c > 0\} \quad \text{باشد}$$

که یک گروه انتقالی از تبدیلات است. در برآورد پارامتر $h(\sigma) = \sigma^r$ تحت تبدیل G ، $h(\sigma)$ دارای برآوردگر هم‌پایا است اگر

$$h(\bar{g}(\sigma)) = \tilde{g}(h(\sigma)) \Leftrightarrow (c\sigma)^r = \tilde{g}(\sigma^r)$$

بنابراین گروه تبدیلات G^- بصورت زیر است $\tilde{G} = \{\tilde{g}_c : \tilde{g}_c(\delta) = c^r \sigma, c > 0\}$

بنابراین تحت گروه G تابع زیان $L(\sigma, \delta)$ پایا است اگر و فقط اگر

$$L(\bar{g}(\sigma), \tilde{g}(\delta)) = L(\sigma, \delta) \Leftrightarrow L(c\sigma, c^r \delta) = L(\sigma, \delta) \stackrel{(*)}{\Leftrightarrow} L(\sigma, \delta) = \gamma\left(\frac{\delta}{\sigma^r}\right)$$

$$\Rightarrow c = \frac{1}{\sigma} \Rightarrow L(\sigma, \delta) = L\left(1, \frac{\delta}{\sigma^r}\right) = \gamma\left(\frac{\delta}{\sigma^r}\right) \quad \text{اثبات (*):}$$

$$\Leftrightarrow L(c\sigma, c^r \delta) = \gamma\left(\frac{c^r \delta}{(c\sigma)^r}\right) = \gamma\left(\frac{\delta}{\sigma^r}\right) = L(\sigma, \delta)$$

بنابراین اگر تابع زیان بفرم $L(\sigma, \delta) = \gamma\left(\frac{\delta}{\sigma^r}\right)$ باشد آنگاه مسئله ما تحت گروه تبدیلات G پایا

است. برای مثال $L(\sigma, \delta) = \left(1 - \frac{\delta}{\sigma^r}\right)^2$ و $L(\sigma, \delta) = \left|1 - \frac{\delta}{\sigma^r}\right|$ از این نوع تابع زیان هستند.

همچنین برآوردهای هم‌پایا در این حالت بصورت زیر می‌باشند

$$\delta(g(X)) = \tilde{g}(\delta(X)) \Leftrightarrow \delta(cX) = c^r \delta(X)$$

از جمله برآوردهای هم‌پایای پارامتر σ ($r=1$) عبارتند از

$$D = \frac{1}{n} \sum_{i=1}^n |X_i - \bar{X}| \quad S = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}$$

برای بدست آوردن برآوردگر MRE در این حالت از قضیه زیر استفاده می‌کنیم.

قضیه ۱-۳: اگر $\delta(X)$ هر برآوردگر هم پایایی σ^r باشد آنگاه شرط لازم و کافی برای آنکه برآوردگر $\delta(X)$ برای σ^r هم پایا باشد آن است که

$$\delta(x) = \frac{\delta.(x)}{U(x)} \quad (۱)$$

که در آن $U(x)$ هر تابع پایا نسبت به G است یعنی

$$U(cx) = U(x) \quad (۲)$$

اثبات: فرض کنید (۱) و (۲) برقرار باشند در این صورت

$$\delta(cx) = \frac{\delta.(cx)}{U(cx)} = \frac{c^r \delta.(x)}{U(x)} = c^r \delta(x)$$

پس $\delta(x)$ یک برآوردگر هم پایا است. حال فرض کنید $\delta(x)$ برآوردگری هم پایا باشد قرار می دهیم.

$$U(x) = \frac{\delta.(x)}{\delta(x)} \quad \text{در این صورت رابطه (۱) برقرار است و همچنین}$$

$$U(cx) = \frac{\delta.(cx)}{\delta(cx)} = \frac{c^r \delta.(x)}{c^r \delta(x)} = \frac{\delta.(x)}{\delta(x)} = U(x)$$

■

اگر قرار دهیم $Z = (Z_1, \dots, Z_n)$ که $Z_i = \frac{X_i}{X_n} i = 1, \dots, n-1, Z_n = \frac{X_n}{|X_n|}$ و $Z = c$ اگر $n=1$

باشد) آنگاه Z یک تابع پایای ماکزیمال نسبت به G است و طبق قضیه ۱-۴ تابع $U(x)$ پایا است اگر و فقط اگر تابعی از تابع پایای ماکزیمال Z باشد. بنابراین از قضیه ۱-۳ نتیجه می شود که:

نتیجه ۱-۳: اگر δ هر برآوردگر هم پایایی σ^r باشد آنگاه شرط لازم و کافی برای این که برآوردگر δ برای σ^r هم پایا باشد این است که تابع w از n متغیر

$$Z_i = \frac{X_i}{X_n} i = 1, \dots, n-1, Z_n = \frac{X_n}{|X_n|}$$

$$\delta(x) = \frac{\delta.(x)}{w(z)} \quad \text{وجود داشته باشد بطوری که}$$

■

حال چون گروه G^- انتقالی است و Z یک تابع پایای ماکزیمال است در این صورت طبق قضیه ۷-۱ برآوردگر هم پایای $\delta(X) = \frac{\delta(X)}{w(Z)}$ که امید شرطی $E_{\sigma=1}[\gamma(\frac{\delta(X)}{w(Z)})|Z]$ را مینیمم کند برآوردگر MRE پارامتر σ^r است. حال اگر $w^*(Z)$ این امید شرطی را مینیمم کند آنگاه $\delta(X) = \frac{\delta(X)}{w^*(Z)}$ برآوردگر MRE پارامتر σ^r است.

لم ۳-۱: فرض کنید X یک متغیر تصادفی مثبت باشد

الف- اگر $E(X^2) < \infty$ آنگاه $E\left(\frac{X}{c} - 1\right)^2$ مقدار مینیمم خود را به ازای $c = \frac{E(X^2)}{E(X)}$ بدست می‌آورد.

ب- اگر $E(X) < \infty$ و X دارای تابع چگالی f باشد آنگاه $E\left(\frac{X}{c} - 1\right)$ مقدار مینیمم خود را به ازای مقدار c بدست می‌آورد که در رابطه زیر صدق کند

$$\int_c^{\infty} xf(x)dx = \int_c^{\infty} xf(x)dx$$

(هر مقداری c که در رابطه فوق صدق می‌کند را میانه - مقیاس^۱ توزیع X می‌نامند)

اثبات: تمرین. ■

نتیجه ۳-۲: در خانواده توزیع‌های مقیاس و تحت تبدیلات مقیاس G داریم که

الف- اگر $\gamma\left(\frac{\delta}{\sigma^r}\right) = \left(\frac{\delta}{\sigma^r} - 1\right)^2$ (فرم A) آنگاه برآوردگر MRE پارامتر σ^r عبارت است از

$$\delta^*(X) = \frac{\delta(X)E[\delta(X)|Z]}{E[\delta^2(X)|Z]}$$

ب- اگر $\gamma\left(\frac{\delta}{\sigma^r}\right) = \left|\frac{\delta}{\sigma^r} - 1\right|$ (فرم B) آنگاه $\delta^*(X) = \frac{\delta(X)}{w^*(Z)}$ که در آن $w^*(Z)$ هر میانه مقیاس

توزیع $\delta(X)$ به شرط Z با $\sigma = 1$ می‌باشد.

اثبات: تمرین. ■

^۱ Scale-median

مثال ۳-۱: (حالت $n=1$) اگر X داری توزیع احتمال $\frac{1}{\sigma} f\left(\frac{x}{\sigma}\right)$ باشد در این صورت $w^*(z)$ یک

مقدار ثابت است و طبق مطالب فوق $\delta^*(X) = \frac{X^r}{w^*}$ برآوردگر MRE پارامتر σ^r است که w^*

مقداری است که $E_1\left[\gamma\left(\frac{X^r}{w}\right)\right]$ را مینیمم می‌کند. بنابراین تحت تابع زیان فرم A ،

$$\delta^*(X) = \frac{X^r E_1(X^r)}{E_1(X^{2r})} \text{ و تحت تابع زیان بفرم } B, \delta^*(X) = \frac{X^r}{w^*} \text{ که } w^* \text{ هر میانه مقیاس}$$

توزیع X^r می‌باشد، برآوردگرهای MRE پارامتر σ^r هستند. ■

مثال ۳-۲: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(0, \sigma^2)$ باشند برآوردگر MRE پارامتر σ^2 را تحت تابع زیان بفرم A بدست آورید.

$$f_X(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x}{\sigma}\right)^2} = \frac{1}{\sigma} f\left(\frac{x}{\sigma}\right) \quad \text{حل:}$$

قرار می‌دهیم $\delta(X) = \sum_{i=1}^n X_i^2$ در این صورت $\delta(X)$ آماره بسنده کامل برای σ^2 و همچنین

برآوردگر هم‌پایای σ^2 است (زیرا $(\delta(cX)) = c^2 \delta(X)$) و طبق قضیه باسو از آماره فرعی Z (پایای ماکزیمال) مستقل است. بنابراین

$$\delta^*(X) = \frac{\delta(X) E_1[\delta(X)]}{E_1[\delta^2(X)]} = \frac{\sum_{i=1}^n X_i^2 E(\chi_{(n)}^2)}{E_1[(\chi_{(n)}^2)^2]} = \frac{n \sum_{i=1}^n X_i^2}{2n + n^2} = \frac{1}{n+2} \sum_{i=1}^n X_i^2$$

■ نشان دهید که این برآوردگر مخاطره نااریب است.

مثال ۳-۳: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $U(0, \theta)$ باشند برآوردگر MRE پارامتر θ را تحت توابع زیان A و B بدست آورید.

$$f_X(x) = \frac{1}{\theta} I_{(0, \theta)}(x) = \frac{1}{\theta} I_{(0, 1)}\left(\frac{x}{\theta}\right) = \frac{1}{\theta} f\left(\frac{x}{\theta}\right) \quad \text{حل:}$$

قرار می دهیم $\delta(X) = X_{(n)}$ بنابراین $\delta(X)$ آماره بسنده کامل برای θ است و همچنین برآوردگر هم پایای θ است. بنابراین از آماره فرعی (پایای ماکزیمال) Z مستقل است. در نتیجه

الف- تحت تابع زیان بفرم A با $r=1$,

$$\delta^*(X) = \frac{X_{(n)} E_1(X_{(n)})}{E_1(X_{(n)}^r)} = \frac{X_{(n)} \left(\frac{n}{n+1} \right)}{\frac{n}{n+1}} = \frac{n+1}{n+1} X_{(n)} \quad (X_{(n)} \sim Be(n, 1))$$

ب- تحت تابع زیان بفرم B با $r=1$ ، $\delta^*(X) = \frac{X_{(n)}}{w^*}$ که در آن w^* هر میانه مقیاس توزیع $X_{(n)}$ با $\theta=1$ است. یعنی

$$\int_{\cdot}^w x (nx^{n-1}) dx = \int_{\cdot}^1 x (nx^{n-1}) dx \Rightarrow \dots \Rightarrow w^* = n+1 \sqrt{\frac{1}{r}} \Rightarrow \delta^*(X) = n+1 \sqrt{\frac{1}{r}} X_{(n)}$$

■

برآوردگر بیتمن:

اگر تابع زیان بفرم A باشد آنگاه برآوردگر MRE پارامتر مقیاس را بفرم راحت تری می توان بدست آورد.

قضیه ۳-۲: تحت گروه تبدیلات مقیاسی G و در خانواده توزیع های مقیاسی و تحت تابع زیان بفرم A برآوردگر MRE پارامتر مقیاس σ^r عبارت است از

$$\delta^*(x) = \frac{\delta(x) E_1[\delta(x) | z]}{E_1[\delta^r(x) | z]} = \frac{\int_{\cdot}^{+\infty} v^{n+r-1} f(vx_1, \dots, vx_n) dv}{\int_{\cdot}^{+\infty} v^{n+r-1} f(vx_1, \dots, vx_n) dv}$$

این فرم را بعنوان برآورد بیتمن پارامتر σ^r می شناسند.

اثبات: تمرین (همانند اثبات قضیه ۳-۲ می باشد) $\delta(X) = X_n^r$ ■

مثال ۳-۴: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع نمایی با پارامتر λ باشند. برآوردگر بیتمن پارامتر λ را بدست آورید.

حل: $r = 1$

$$f_{\underline{X}}(\underline{x}) = \prod_{i=1}^n \frac{1}{\lambda} e^{-\frac{x_i}{\lambda}} = \frac{1}{\lambda^n} e^{-\sum_{i=1}^n \frac{x_i}{\lambda}} = \frac{1}{\lambda^n} f\left(\frac{x_1}{\lambda}, \dots, \frac{x_n}{\lambda}\right) = v^n f(vx_1, \dots, vx_n)$$

$$\delta^*(\underline{x}) = \frac{\int_0^{+\infty} v^n f(v\underline{x}) dv}{\int_0^{+\infty} v^{n+1} f(v\underline{x}) dv} = \frac{\int_0^{+\infty} v^n e^{-v \sum x_i} dv}{\int_0^{+\infty} v^{n+1} e^{-v \sum x_i} dv} = \frac{\frac{\Gamma(n+1)}{(\sum x_i)^{n+1}}}{\frac{\Gamma(n+2)}{(\sum x_i)^{n+2}}} = \frac{\sum x_i}{n+1}$$

پس برآوردگر MRE پارامتر λ برابر $\delta^*(\underline{X}) = \frac{\sum X_i}{n+1}$ است (که مخاطره نااریب نیز است). ■

مثال ۳-۵: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $U(\cdot, \theta)$ باشند. برآوردگر پیتمن پارامتر θ را بدست آورید.

حل:

$$f_{\underline{X}}(\underline{x}) = \begin{cases} \prod_{i=1}^n \frac{1}{\theta} I_{(\cdot, \cdot)}\left(\frac{x_i}{\theta}\right) = \frac{1}{\theta^n} \prod_{i=1}^n I_{(\cdot, \cdot)}\left(\frac{x_i}{\theta}\right) = \frac{1}{\theta^n} f\left(\frac{x_1}{\theta}, \dots, \frac{x_n}{\theta}\right) = v^n f(v\underline{x}) \\ \frac{1}{\theta^n} I_{(x_{(n)}, +\infty)}(\theta) = v^n I_{(\cdot, \frac{1}{x_{(n)}})}(v) \quad v = \frac{1}{\theta} \end{cases}$$

$$\delta^*(\underline{x}) = \frac{\int_0^{+\infty} v^n f(v\underline{x}) dv}{\int_0^{+\infty} v^{n+1} f(v\underline{x}) dv} = \frac{\int_0^{1/x_{(n)}} v^n dv}{\int_0^{1/x_{(n)}} v^{n+1} dv} = \frac{\frac{v^{n+1}}{n+1} \Big|_0^{1/x_{(n)}}}{\frac{v^{n+2}}{n+2} \Big|_0^{1/x_{(n)}}} = \frac{n+2}{n+1} x_{(n)}$$

■ که همان برآوردگر بدست آمده در مثال ۳-۳ است.

قضیه ۳-۳: برآوردگر پیتمن پارامتر مقیاس تابعی از آماره بسنده است.

■ اثبات: تمرین.

بخش ۴: برآوردگر MRE پارامترهای مکان-مقیاس

فرض کنید $X = (X_1, \dots, X_n)$ دارای توزیع احتمال توام $\frac{1}{\sigma^n} f\left(\frac{x_1 - \theta}{\sigma}, \dots, \frac{x_n - \theta}{\sigma}\right)$ باشند که در آن f معلوم و θ و σ هر دو نامعلوم هستند. در زیر برآوردگرهای MER برای θ و σ را به طور جداگانه بدست می‌آوریم.

الف- برآوردگرهای MRE پارامتر مقیاس σ :

فرض کنید بخواهیم پارامتر σ^r را برآورد کنیم. اگر گروه تبدیلات زیر را در نظر بگیریم.

$$G = \{g_{a,b} \mid g_{a,b}(x) = a + bx, a \in R, b > 0\}$$

آنگاه خانواده توزیع‌های فوق نسبت به این تبدیل پایا است هرگاه گروه تبدیلات G^- به صورت زیر باشد.

$$G^- = \{\bar{g}_{a,b} \mid \bar{g}_{a,b}(\theta, \sigma) = (a + b\theta, b\sigma), a \in R, b > 0\}$$

که یک گروه انتقالی از تبدیلات است. در برآورد پارامتر $h(\theta, \sigma) = \sigma^r$ تحت تبدیل G ، $h(\sigma)$ دارای برآوردگر هم‌پایا است اگر

$$h(\bar{g}(\theta, \delta)) = \tilde{g}(h(\theta, \sigma)) \Leftrightarrow h(a + b\theta, b\sigma) = \tilde{g}(\sigma^r) \Leftrightarrow b^r \sigma^r = \tilde{g}(\sigma^r)$$

بنابراین گروه تبدیلات G^- به صورت زیر است.

$$\tilde{G} = \{\tilde{g}_{a,b} \mid \tilde{g}_{a,b}(\delta) = b^r \delta, b > 0, a \in R\}$$

بنابراین تحت گروه G تابع زیان $L(\theta, \sigma, \delta)$ پایا است اگر و فقط اگر

$$L(\theta, \sigma, \delta) = L(a + b\theta, b\sigma, b^r \delta) \quad \Leftrightarrow \quad \begin{matrix} a = -\theta \\ \sigma = \frac{1}{b} \end{matrix} L(\theta, \sigma, \delta) = \gamma\left(\frac{\delta}{\sigma^r}\right)$$

(برای مثال تابع زیان‌های بفرم A و B).

در این حالت برآوردگرهای هم‌پایا بفرم زیر می‌باشند.

$$\delta(g(X)) = \tilde{g}(\delta(X)) \Leftrightarrow \delta(a + bX) = b^r (\delta(X)) \quad (1)$$

برای بدست آوردن برآوردگر MRE در این حالت به صورت زیر عمل می‌کنیم. اگر تنها یک تغییر در مکان داشته باشیم یعنی:

$$g(\underline{x}) = \underline{x} + a \quad \tilde{g}(\theta, \sigma) = (\theta + a, \sigma) \quad (b = 1)$$

در این صورت رابطه (۱) با $b = 1$ نتیجه می‌دهد که بایستی

$$\delta(a + b\underline{X}) = \delta(\underline{X})$$

یعنی δ یک تابع پایا نسبت به G است پس طبق نتیجه ۱-۲، δ تابعی از $n-1$ تفاضل $Y_i = X_i - X_n$ ، $i = 1, \dots, n-1$ می‌باشد. اما تابع چگالی توام Y_i ها بصورت زیر بدست می‌آید:

$$W_i = X_i - \theta \quad i = 1, \dots, n \Rightarrow f_W(w) = f_X(w + \theta) = \frac{1}{\sigma^n} f\left(\frac{w}{\sigma}\right)$$

$$\begin{cases} Y_i = X_i - X_n = W_i - W_n & i = 1, \dots, n-1 \\ T = W_n \end{cases} \Rightarrow J = 1$$

$$f_{Y,T}(y, t) = f_W(y_1 + t, \dots, y_{n-1} + t, t) = \frac{1}{\sigma^n} f\left(\frac{y_1 + t}{\sigma}, \dots, \frac{y_{n-1} + t}{\sigma}, \frac{t}{\sigma}\right)$$

$$\Rightarrow f_Y(y) = \frac{1}{\sigma^n} \int_{-\infty}^{+\infty} f\left(\frac{y_1 + t}{\sigma}, \dots, \frac{y_{n-1} + t}{\sigma}, \frac{t}{\sigma}\right) dt$$

$$= \frac{1}{\sigma^{n-1}} \int_{-\infty}^{+\infty} f\left(\frac{y_1}{\sigma} + u, \dots, \frac{y_{n-1}}{\sigma} + u, u\right) du = \frac{1}{\sigma^{n-1}} f^*\left(\frac{y}{\sigma}\right)$$

بنابراین توزیع $Y = (Y_1, \dots, Y_{n-1})$ به یک خانواده مقیاس متعلق است در نتیجه در این خانواده برای

برآورد پارامتر σ^r از برآوردگرهای MRE بدست آمده در بخش قبل استفاده می‌کنیم. بنابراین

برآوردگر MRE پارامتر σ^r بفرم $\delta(\underline{X}) = \frac{\delta(\underline{Y})}{w^*(\underline{Z})}$ است که $\delta(\underline{Y})$ هر برآوردگر پایای مقیاس

پارامتر σ^r بر اساس $\underline{Y}$ است $(\delta(b\underline{Y}) = b^r \delta(\underline{Y}))$ که در آن $Z_{n-1} = \frac{Y_{n-1}}{|Y_{n-1}|}$ و

$\underline{Z} = (Z_1, \dots, Z_{n-1})$ ، $Z_i = \frac{Y_i}{Y_{n-1}}$ ، $i = 1, 2, \dots, n-2$ و $w^*(\underline{z})$ هر عددی است که

$$E_{\sigma=1} \left\{ \gamma \left(\frac{\delta(\underline{Y})}{w(\underline{Z})} \right) \middle| \underline{z} \right\} \text{ را مینیمم می‌کند.}$$

نتیجه ۴-۱: اگر تابع زیان بفرم $A \left(L(\theta, \sigma, \delta) = \left(\frac{\delta}{\sigma^r} - 1 \right)^2 \right)$ باشد آنگاه برآوردگر MRE

$$\delta^*(X) = \frac{\delta(Y) E_{\sigma=1}[\delta(Y) | Z]}{E_{\sigma=1}[\delta^2(Y) | Z]}$$

پارامتر σ^r عبارت است از

و اگر تابع زیان بفرم $B \left(L(\theta, \sigma, \delta) = \left(\frac{\delta}{\sigma^r} - 1 \right) \right)$ باشد آنگاه برآوردگر MRE پارامتر σ^r عبارت

است از $\delta^*(X) = \frac{\delta(X)}{w^*(Z)}$ که در آن $w^*(Z)$ هر میانه مقیاس توزیع $\delta(Y)$ به شرط Z با $\sigma = 1$ می باشد. ■

مثال ۴-۱: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\mu, \sigma^2)$ باشند که در آن μ و σ^2 هر دو نامعلوم هستند. برآورد MRE پارامتر σ^2 را تحت تابع زیانهای بفرم A و B بدست آورید.

حل: $T = (\bar{X}, \sum_{i=1}^n (X_i - \bar{X})^2)$ یک آماره بسنده کامل برای μ و σ^2 است. قرار می دهیم

$$\delta = \sum_{i=1}^n (X_i - \bar{X})^2 \quad \text{در این صورت } \delta \text{ یک برآوردگر هم پایا برای } \sigma^2 \text{ است}$$

$(\delta(Y) = b^2 \delta(Y), \delta(a + bX) = b^2 \delta(X))$ و از آماره پایای ماکزیمال Z (فرعی) مستقل

$$\delta \sim \chi^2_{(n-1)} \quad (\sigma = 1) \quad \text{است و}$$

الف- اگر تابع زیان بفرم A باشد آنگاه

$$\delta^*(X) = \frac{\delta E_1(\delta | Z)}{E_1(\delta^2 | Z)} = \frac{\delta E_1(\delta)}{E_1(\delta^2)} = \frac{\delta (n-1)}{2(n-1) + (n-1)^2} = \frac{\delta}{n+1} = \frac{1}{n+1} \sum_{i=1}^n (X_i - \bar{X})^2$$

ب- اگر تابع زیان بفرم B باشد آنگاه w^* میانه-مقیاس توزیع $\delta | Z = D$ است.

$$\int_{\cdot}^{w*} \frac{x}{\frac{n-1}{2} \Gamma\left(\frac{n-1}{2}\right)} x^{\frac{n-1}{2}-1} e^{-\frac{1}{2}x} dx = \int_{w*}^{+\infty} \frac{x}{\frac{n-1}{2} \Gamma\left(\frac{n-1}{2}\right)} x^{\frac{n-1}{2}-1} e^{-\frac{1}{2}x} dx$$

$$\Leftrightarrow \dots \Leftrightarrow \int_{\cdot}^{w*} \frac{1}{\frac{n+1}{2} \Gamma\left(\frac{n+1}{2}\right)} x^{\frac{n+1}{2}-1} e^{-\frac{1}{2}x} dx = \int_{w*}^{+\infty} \frac{x}{\frac{n+1}{2} \Gamma\left(\frac{n+1}{2}\right)} x^{\frac{n+1}{2}-1} e^{-\frac{1}{2}x} dx$$

یعنی w^* میانه توزیع $\chi^2_{(n+1)}$ است و بنابراین $w^* = \chi^2_{(n+1)}/2$ و در نتیجه

$$\delta^*(\tilde{X}) = \frac{1}{\chi^2_{(n+1)}/2} \sum_{i=1}^n (X_i - \bar{X})^2$$

مثال ۴-۲: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع $U(\theta - \frac{1}{4}\sigma, \theta + \frac{1}{4}\sigma)$ باشد.

برآوردگر MRE پارامتر σ را تحت توابع زیان بفرم A و B بدست آورید.

$$f_{\theta, \sigma}(x) = \frac{1}{\sigma} I_{(\theta - \frac{1}{4}\sigma, \theta + \frac{1}{4}\sigma)}(x) = \frac{1}{\sigma} I_{(-\frac{1}{4}, \frac{1}{4})}\left(\frac{x - \theta}{\sigma}\right) \quad \text{حل:}$$

$T = (X_{(1)}, X_{(n)})$ یک آماره بسنده مینیمال برای (θ, σ) است و $\delta = X_{(n)} - X_{(1)}$ برآوردگر هم پایای σ است که از آماره پایای ماکزیمال (فرعی) Z مستقل است و اگر $\sigma = 1$ باشد.

$$\delta = X_{(n)} - X_{(1)} = (X_{(n)} - \theta + \frac{1}{4}) - (X_{(1)} - \theta + \frac{1}{4}) = V_{(n)} - V_{(1)}$$

$$\stackrel{D}{=} V_{(n-1)} \sim Be(n-1, 2)$$

الف- اگر تابع زیان بفرم A باشد آنگاه

$$\delta^*(\tilde{X}) = \frac{\delta E_1[\delta]}{E_1[\delta^2]} = \frac{\delta \cdot \left(\frac{n-1}{n+1}\right)}{\frac{n(n-1)}{(n+1)(n+2)}} = \frac{n+2}{n} \delta = \frac{n+2}{n} (X_{(n)} - X_{(1)})$$

ب- اگر تابع زیان به فرم B باشد برآوردگر MRE، σ را بدست آورید.

توجه: توجه کنید که اگر گروه تبدیلات $G = \{g_b(x) = bx, b > 0\}$ را به کار می‌بریم (همانند $G^- = \{\bar{g}_b(\sigma) = b\sigma, b > 0\}$)

بخش قبل) آنگاه G^- دیگر یک گروه انتقالی نبود و از روش بکار برده شده در این فصل نمی‌توان

برآوردگر MRE، σ^r را بدست آورد. (اشتاین^۱ (۱۹۶۴) نشان داد که $\delta(\underline{x}) = \psi\left(\frac{\bar{x}}{s}\right)s^2$ که

$$\psi\left(\frac{\bar{x}}{s}\right) = \min \left\{ \frac{1}{n+1}, \frac{1 + \frac{n\bar{x}^2}{s}}{n+2} \right\}$$

تحت تبدیل فوق برای σ^2 در توزیع نرمال دارای مخاطره

$$\delta = \frac{1}{n+1} \sum_{i=1}^n (X_i - \bar{X})^2 \text{ به } \delta \text{ در مثال ۴-۱ است}$$

ب- برآوردگر MRE پارامتر مکان θ

فرض کنید بخواهیم پارامتر مکان θ را در خانواده توزیع‌های $\frac{1}{\sigma^n} f\left(\frac{x_1 - \theta}{\sigma}, \dots, \frac{x_n - \theta}{\sigma}\right)$ برآورد

کنیم. همان‌گونه که دیدیم این خانواده تحت گروه‌های تبدیلات

$$\begin{aligned} G &= \{g_{a,b} \mid g_{a,b}(\underline{x}) = a + b\underline{x}\} \\ G^- &= \{\bar{g}_{a,b} \mid \bar{g}_{a,b}(\theta, \sigma) = (a + b\theta, b\sigma)\} \end{aligned}$$

پایا است و در برآورد پارامتر θ گروه تبدیلات G^- به صورت زیر می‌باشد

$$\tilde{G} = \{\tilde{g}_{a,b}(\delta) = a + b\delta\}$$

در این صورت مسئله مورد نظر تحت گروه G پایا است اگر و فقط اگر تابع زیان بفرم زیر باشد

$$L(\theta, \sigma, \delta) = L(a + b\theta, b\sigma, a + b\delta) \quad \Leftrightarrow \quad \overset{a = \frac{-\theta}{\sigma}, b = \frac{1}{\sigma}}{L(\theta, \sigma, \delta) = \rho\left(\frac{\delta - \theta}{\sigma}\right)}$$

در این حالت برآوردهای هم‌پایا به فرم زیر می‌باشند

$$\delta(g(\underline{X})) = \tilde{g}(\delta(\underline{X})) \Leftrightarrow \delta(a + b\underline{x}) = a + b\delta(\underline{x}) \quad (2)$$

چون G^- انتقالی است پس مخاطره هر برآوردگر هم‌پایا ثابت است. برای بدست آوردن برآوردگر

MRE پارامتر θ دو حالت زیر پیش می‌آید.

حالت اول: برای یک مقدار ثابت σ قرار می‌دهیم

$$g_{\sigma}(x_1, \dots, x_n) = \frac{1}{\sigma^n} f\left(\frac{x_1}{\sigma}, \dots, \frac{x_n}{\sigma}\right)$$

^۱ Stein

بنابراین خانواده توزیع‌های مکان مقیاس به صورت زیر تبدیل می‌شود.

$$g_{\sigma}(x_1 - \theta, \dots, x_n - \theta) \quad (۳)$$

که به یک خانواده مکان متعلق است.

لم ۴-۱: فرض کنید برای خانواده مکان (۳) و تحت تابع زیان بفرم $\rho\left(\frac{\delta - \theta}{\sigma}\right)$ یک برآوردگر MRE

δ^* برای θ تحت تبدیلات مکان $G^* = \{g_a(x) = x + a\}$ موجود باشد و

(i) δ^* به σ بستگی نداشته باشد

(ii) δ^* در رابطه (۲) صدق کند.

آنگاه δ^* مخاطره را در بین تمام برآوردگرهایی که در (۲) صدق می‌کنند، مینیمم می‌کند.

اثبات: فرض کنید δ هر برآوردگر هم‌پایای دیگر که در (۲) صدق می‌کند باشد، بنابراین δ تحت G^*

نیز هم‌پایا است. همچنین فرض کنید که σ مقداری معلوم باشد. بنابراین طبق فرض قضیه در مورد δ^* ,

برای این مقدار σ مخاطره δ^* کمتر یا مساوی δ است و چون این موضوع برای هر σ برقرار است

پس δ^* برآوردگر MRE در بین تمام برآوردگرهای هم‌پایای (۲) است. ■

مثال ۴-۳: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\mu, \sigma^2)$ باشند که هر دوی μ و σ

نامعلوم هستند. طبق مثال ۲-۲، $\delta^* = \bar{X}$ برآوردگر MRE پارامتر μ تحت تابع زیان

$$\rho\left(\frac{\delta - \theta}{\sigma}\right) = \left(\frac{\delta - \theta}{\sigma}\right)^2$$

است. چون شرایط (i) و (ii) لم ۴-۱ برقرار است بنابراین $\delta^* = \bar{X}$

برآوردگر MRE پارامتر μ تحت تبدیلات G می‌باشد. ■

مثال ۴-۴: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $U\left(\theta - \frac{1}{2}\sigma, \theta + \frac{1}{2}\sigma\right)$ باشند. در

این صورت طبق مثال ۲-۴، $\delta^* = \frac{X_{(1)} + X_{(n)}}{2}$ برآوردگر MRE پارامتر θ تحت تابع زیان

$$\rho\left(\frac{\delta - \theta}{\sigma}\right) = \left(\frac{\delta - \theta}{\sigma}\right)^2$$

است. چون شرایط (i) و (ii) لم ۴-۱ برقرار است، بنابراین δ^* برآوردگر

MRE پارامتر θ تحت تبدیلات G می‌باشد. ■

توجه: در مثال ۳-۲ دیدیم که در خانواده توزیع‌های $E(\theta, b)$ با b معلوم $\delta^* = X_{(1)} - \frac{b}{n}$ برآوردگر MRE پارامتر θ است. چون شرط (i) لم ۴-۱ برای این برآوردگر صدق نمی‌کند پس از این روش نمی‌توان برآوردگر MRE پارامتر θ را پیدا کرد.

حالت دوم: اگر شرایط حالت اول برقرار نباشد، برای یافتن برآوردگر MRE پارامتر θ همانند بخش‌های قبل به صورت زیر عمل می‌کنیم.

قضیه ۴-۱: فرض کنید δ هر برآوردگر θ باشد که در رابطه زیر صدق کند

$$\delta(a + bx) = a + b\delta(x)$$

و δ_1 هر برآوردگر σ باشد که مقادیر مثبت را بگیرد و در رابطه زیر صدق کند

$$\delta_1(a + bx) = b\delta_1(x)$$

در این صورت δ یک برآوردگر هم‌پایای مکان $(\delta(a + bx) = a + b\delta(x))$ می‌باشد اگر و فقط اگر

$$\delta(x) = \delta_1(x) - w(z)\delta_1(x)$$

$$Y_i = X_i - X_n \quad i = 1, \dots, n-1 \quad \text{که در آن}$$

$$Z = (Z_1, \dots, Z_{n-1}), Z_n = \frac{Y_{n-1}}{|Y_{n-1}|}, Z_i = \frac{Y_i}{Y_{n-1}} \quad i = 1, \dots, n-1$$

اثبات: همانند اثبات قضیه ۲-۱ می‌باشد و بعنوان تمرین واگذار می‌گردد. ■

نتیجه ۴-۲: اگر δ و δ_1 برآوردهای هم‌پایای θ و σ باشند که در شرایط قضیه ۴-۱ صدق می‌کنند آنگاه برآوردگر MRE پارامتر مکان θ عبارت است از

$$\delta^*(x) = \delta_1(x) - w^*(z)\delta_1(x)$$

که در آن برای هر z ، $w(z)$ هر عددی است که $E_{\theta, \sigma}\{\rho[\delta_1(x) - w^*(z)\delta_1(x)] | z\}$ را مینیمم کند. ■

نتیجه ۳-۴: اگر تابع زیان بفرم $\rho\left(\frac{\delta-\theta}{\sigma}\right)=\left(\frac{\delta-\theta}{\sigma}\right)^2$ (فرم C) باشد آنگاه $w^*(z)$ از رابطه زیر

$$w^*(z) = \frac{E_{\gamma_1}[\delta_1(X) \delta_1(X) | Z]}{E_{\gamma_1}[\delta_1^2(X) | Z]}$$
 بدست می‌آید.

اثبات: تمرین. ■

مثال ۵-۴: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع نمایی $E(\theta, \sigma)$ باشند که در آن θ و σ هر دو نامعلوم هستند. برآوردگرهای MRE پارامترهای θ تحت تابع زیان به فرم C و σ تحت تابع زیان به فرم A را بدست آورید.

حل: به راحتی دیده می‌شود که $T = (T_1, T_2) = \left(X_{(1)}, \sum_{i=1}^n (X_i - X_{(1)}) \right)$ یک آماره بسنده کامل

برای (θ, σ) است و $T_1 \sim E(\theta, \frac{\sigma}{n})$ و $T_2 \sim \frac{\sigma}{\gamma} \chi_{(2n-2)}^2$ و T_1 و T_2 از یکدیگر مستقل هستند.

بنابراین $\frac{1}{n-1} \sum_{i=1}^n (X_i - X_{(1)})$ و $X_{(1)} - \frac{1}{n(n-1)} \sum_{i=1}^n (X_i - X_{(1)})$ به ترتیب برآوردگرهای UMVU پارامترهای θ و σ هستند.

الف- در برآورد σ تحت تابع زیان $L(\theta, \sigma, \delta) = \left(\frac{\delta}{\sigma} - 1\right)^2$ (فرم A با $r=1$)

$\delta_1 = \sum_{i=1}^n (X_i - X_{(1)})$ یک برآوردگر هم‌پایا برای σ است $(\delta_1(a + bX) = b\delta_1(X))$ و طبق قضیه

باسو از آماره پایای ماکزیمال (فرعی) Z مستقل است. بنابراین برآوردگر MRE پارامتر σ عبارت است از

$$\delta^*(X) = \frac{\delta_1 E_1[\delta_1 | Z]}{E_1[\delta_1^2 | Z]} = \frac{\delta_1 E_1[\delta_1]}{E_1[\delta_1^2]} = \frac{\delta_1 (n-1)}{(n-1) + (n-1)^2} = \frac{\delta_1}{n} = \frac{1}{n} \sum_{i=1}^n (X_i - X_{(1)})$$

(توجه کنید در بخش قبل دیدیم که اگر $\theta = 0$ باشد، آنگاه برآوردگر MRE پارامتر σ عبارت

$$\text{از } \frac{1}{n+1} \sum_{i=1}^n X_i \text{ بود.)}$$

ب- در برآورد θ تحت تابع زیان به فرم (C)،

$$\delta_1(\underline{X}) = \sum_{i=1}^n (X_i - X_{(1)}) \quad \text{و} \quad \delta_1(\underline{X}) = X_{(1)}$$

به ترتیب برآوردهای هم‌پایا برای θ و σ هستند $\left(\begin{array}{l} \delta_1(a+b\underline{x}) = a+b\delta_1(\underline{x}) \\ \delta_1(a+b\underline{x}) = b\delta_1(\underline{x}) \end{array} \right)$ و بنابر قضیه

باسو (δ_1, δ_1) و $\underline{Z}$ از یکدیگر مستقل هستند. بنابراین برآوردهای MRE پارامتر σ عبارت است از

$$w^*(\underline{z}) = \frac{E_{\cdot,1}[\delta_1 \delta_1 | \underline{z}]}{E_{\cdot,1}[\delta_1^2 | \underline{z}]} = \frac{E_{\cdot,1}[\delta_1] E_{\cdot,1}[\delta_1]}{E_{\cdot,1}[\delta_1^2]} = \frac{\frac{1}{n}(n-1)}{(n-1) + (n-1)^2} = \frac{\frac{1}{n}}{\frac{1}{n} + \frac{1}{n^2}} = \frac{1}{1 + \frac{1}{n}}$$

$$\Rightarrow \delta^*(\underline{X}) = \delta_1(\underline{X}) - w^*(\underline{z})\delta_1(\underline{X}) = X_{(1)} - \frac{1}{n^2} \sum_{i=1}^n (X_i - X_{(1)})$$

(توجه کنید که در بخش‌های قبل دیدیم که برآوردهای MRE پارامتر مکان θ با فرض معلوم بودن σ

برابر $X_{(1)} - \frac{\sigma}{n}$ بود). ■

توجه: تحت تابع زیان $\left(\frac{\delta - \theta}{\sigma} \right)^2$ هیچ برآوردهای مخاطره نااریب وجود ندارد. توجه کنید که گروه $\tilde{G}$

در این حالت غیرجابجایی است. ■

فصل ۴

تصمیم ییزی

بخش ۱: مقدمه

تمامی روش‌های استنباط آماری که تاکنون بررسی کرده‌ایم یک نوع تصمیم‌گیری می‌باشند و در آنها به کمک نمونه‌های آماری و داده‌ها درباره پارامتر یا توزیع مجهول به تصمیم‌گیری می‌پرداختیم و به آنها آمار کلاسیک گویند. در نظریه تصمیم به کمک نمونه‌های آماری و دیگر جوانب، با بررسی مقدار سود و زیان می‌کوشیم تا به بهترین تصمیم‌گیری دست یابیم.

بنابراین نظریه تصمیم شامل روش‌های دیگری از مسائل استنباط آماری است که روش‌های قبلی را شامل می‌باشد و جنبه‌های جدیدی را نیز در نظر می‌گیرد. در فصل‌های قبل با عناصر یک مسئله تصمیم آشنا شده‌ایم ولی آنها را بطور کلی معرفی نکرده‌ایم. در زیر این عناصر را معرفی می‌کنیم.

داده‌ها: داده‌ها بوسیله یک بردار تصادفی X با فضای نمونه (مجموعه مقادیر) $\mathcal{X}$ نمایش داده می‌شود.

وضع طبیعت^۱: پارامتر واقعی و مجهول θ را که می‌خواهیم روی آن استنباط انجام دهیم را وضع طبیعت گویند.

فضای پارامتری^۲: مجموعه تمام مقادیر ممکن θ را فضای پارامتری گویند و با نماد Θ نمایش می‌دهند.

مدل: مدل عبارت است از مجموعه کلیه توزیع‌های احتمال برای بردار X که دارای پارامتر مجهول θ است. یعنی مجموعه $\{f(x|\theta): \theta \in \Theta\}$ که در آن $f(x|\theta)$ تابع چگالی یا تابع احتمال روی $\mathcal{X}$ است.

فضای کارها^۳: بعد از اینکه $X = \tilde{x}$ مشاهده گردید، تصمیمی برای θ اتخاذ می‌گردد. مجموعه کلیه تصمیم‌های ممکن را فضای کارها گویند و با نماد A نمایش می‌دهند و اعضای آن را با $a \in A$ نشان می‌دهند. فضای کارها معین می‌کند که مسئله تصمیم مورد مطالعه، یک مسئله برآورد نقطه‌ای یا فاصله اطمینان و یا آزمون فرض می‌باشد. در مسئله برآورد نقطه‌ای می‌خواهیم نقاطی از فضای کارها را

^۱ State of Nature^۲ Parameter Space^۳ Action Space

برای θ حدس بزنیم پس $\theta \in A$ خواهد بود (مثلاً) در برآورد پارامتر توزیع دو جمله‌ای $0 < p < 1$ و می‌تواند $A = [0, 1]$ باشد) در مسئله آزمون فرض، ما دو کار "قبول H " و "رد H " را داریم که آن را با a_0 و a_1 نمایش می‌دهیم و بنابراین $A = \{a_0, a_1\}$. در یک مسئله فاصله اطمینان، کارها فواصل یا زیرمجموعه‌هایی از فضای پارامتر هستند و بنابراین $A \subset \Theta$.

تابع زیان: اگر $\theta \in \Theta$ مقدار واقعی وضع طبیعت باشد در این صورت کار $a \in A$ ممکن است درست یا تا اندازه‌ای نادرست و یا کلاً نادرست باشد. میزان این نادرستی را با کمیت تابع زیان $L(\theta, a)$ اندازه‌گیری می‌کنیم که مقدار زیان بکار بردن کار a موقعی که θ مقدار وضع طبیعت واقعی باشد را اندازه می‌گیرد. بنابراین $L: \Theta \times A \rightarrow [0, +\infty)$ و $L(\theta, a) = 0$ به معنای این است که اگر θ وضع طبیعت واقعی است آنگاه a تصمیم درستی می‌باشد.

تابع تصمیم^۱: تابع تصمیم (یا قاعده تصمیم) قاعده‌ای است که با مشاهده $\tilde{X} = X$ معین می‌کند که چه کار $a \in A$ را بایستی بکار ببریم و معمولاً آنرا با $\delta(\tilde{X})$ نمایش می‌دهند. بنابراین

$$\delta(\tilde{X}): \mathcal{X} \rightarrow A$$

(در مسئله برآورد نقطه‌ای، توابع تصمیم همان برآوردها می‌باشند)

فضای توابع تصمیم: مجموعه کلیه توابع تصمیم ممکن را با D نمایش می‌دهیم. در نظریه تصمیم هدف این است که معین کنیم که کدامیک از توابع تصمیم $\delta(\tilde{X}) \in D$ یک قاعده تصمیم خوب می‌باشد.

تابع مخاطره^۲ (ریسک): در مسئله تصمیم روی پارامتر θ براساس تابع تصمیم $\delta(\tilde{X})$ میزان دقت تابع تصمیم بوسیله تابع مخاطره یعنی $R(\theta, \delta) = E_{\theta}[L(\theta, \delta(\tilde{X}))]$ اندازه گرفته می‌شود. چون مقدار θ نامعلوم است علاقه‌مند هستیم توابع تصمیمی را بکار ببریم که $R(\theta, \delta)$ را به ازای هر $\theta \in \Theta$ مینیمم کند. اما در حالت کلی چنین کاری ممکن نیست و بنابراین بایستی محدودیت‌هایی روی فضای توابع تصمیم (کلاس برآوردها) برای دستیابی به تصمیم بهینه اعمال کنیم. یک روش

^۱ Decision Rule

^۲ Risk Function

برای پیدا کردن تصمیم بهینه، برقراری یک رابطه ترتیب بین تصمیم‌ها می‌باشد. دو روش اساسی در برقراری رابطه‌ای ترتیبی بین تصمیمات، اصل کمین بیشینه^۱ (مینیماکس) و اصل بیز^۲ می‌باشد. در این فصل و فصل بعد این دو روش را برای بدست آوردن برآورد نقطه‌ای برای پارامتر مجهول θ بررسی می‌کنیم.

بخش ۲- تصمیم بیزی^۳

در روش‌هایی که قبلاً مطالعه کردیم (روش‌های کلاسیک) پارامتر θ را یک مقدار ثابت نامعلوم در نظر گرفتیم و یک نمونه تصادفی $X_1, X_2, \dots, X_n$ از جمعیت که دارای توزیع $f_\theta(x)$ بود جمع آوری کرده و براساس آن در مورد θ تصمیم‌گیری می‌کردیم. در روش بیزی θ را کمیتی در نظر می‌گیریم که خود یک متغیر تصادفی است و تغییرات آن توسط یک توزیع احتمال (که آنرا توزیع پیشین می‌نامیم) بیان می‌گردد. این توزیع پیشین براساس اعتقادات و تجربیات قبلی آزمایشگر و قبل از مشاهده داده‌ها تعیین می‌گردد. سپس از جمعیت یک نمونه جمع آوری می‌شود و بر اساس آن توزیع پیشین تصحیح می‌گردد. توزیع پیشین تصحیح شده را توزیع پسین می‌نامند و براساس این توزیع پسین تصمیم‌گیری انجام می‌گیرد. برای مثال فرض کنید قطعات تولیدی یک کارخانه دارای توزیع نمایی با میانگین θ ساعت می‌باشند و به علت فرسودگی دستگاه‌ها این میانگین θ نیز در سال تغییر می‌کند و تغییرات آن طبق یک توزیع احتمال $\pi(\theta)$ باشد.

مخاطره و تصمیم بیزی:

فرض کنید $\pi(\theta)$ توزیع احتمال روی فضای پارامتری Θ باشد که به آن توزیع پیشین^۴ θ گویند چون تابع مخاطره $R(\theta, \delta) = E_\theta[L(\theta, \delta(X))]$ تابعی از θ است و θ متغیر تصادفی است پس $R(\theta, \delta)$ نیز یک متغیر تصادفی است و تابعی از θ است. امید ریاضی (متوسط) $R(\theta, \delta)$ را نسبت به چگالی پیشین θ مخاطره بیز می‌نامند یعنی

^۱ Minimax

^۲ Bayes

^۳ Bayesian Decision

^۴ Prior Distribution

$$r(\pi, \delta) = E[R(\theta, \delta)] = \int_{\Theta} R(\theta, \delta) d\Pi(\theta) = \begin{cases} \int R(\theta, \delta) \pi(\theta) d\theta & \text{اگر } \pi(\theta) \text{ پیوسته باشد} \\ \sum_{\Theta} R(\theta, \delta) \pi(\theta) & \text{اگر } \pi(\theta) \text{ گسسته باشد} \end{cases}$$

تصمیم بیزی نسبت به توزیع پیشین π آن تصمیم $\delta^\pi(x)$ است که مخاطره بیزی را در بین تمام توابع تصمیم مینیمم کند، یعنی تصمیم بیزی آن تابع تصمیمی است که در رابطه زیر صدق کند.

$$r(\pi, \delta^\pi) = \inf_{\delta \in D} r(\pi, \delta)$$

روش بدست آوردن تصمیم بیزی:

فرض کنید θ دارای توزیع پیشین $\pi(\theta)$ باشد (در اینجا θ هم بعنوان پارامتر و هم به عنوان متغیر تصادفی در نظر گرفته می شود) و $X = (X_1, \dots, X_n)$ یک نمونه تصادفی جمع آوری شده از جمعیتی باشد که توزیع آن به θ بستگی دارد. براساس این نمونه تصادفی در توزیع پیشین $\pi(\theta)$ تجدید نظر می کنیم و چگالی پسین θ را به صورت زیر بدست می آوریم

$$\pi(\theta | x) = \frac{f(x, \theta)}{m(x)}$$

که در آن $f(x, \theta)$ چگالی توام X و θ و $m(x)$ چگالی کناری X است بنابراین

$$f(x, \theta) = \pi(\theta) f(x | \theta), \quad m(x) = \int f(x, \theta) d\Pi(\theta)$$

$$\pi(\theta | x) = \frac{\pi(\theta) f(x, \theta)}{m(x)} = c(x) \pi(\theta) f(x | \theta) \propto \pi(\theta) f(x | \theta) = \pi(\theta) L(\theta) \quad \text{پس}$$

یعنی برای محاسبه توزیع احتمال پسین θ کافیت حاصلضرب $\pi(\theta)$ در $L(\theta)$ (تابع درست نمایی) را بدست آورده و آنرا به یک توزیع احتمال تبدیل کنیم.

مثال ۱-۲: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از توزیع $B(1, \theta)$ باشند و θ دارای توزیع پیشین $Be(\alpha, \beta)$ باشد که در آن α و β معلوم هستند. چگالی پسین θ را بدست آورید.

حل:

$$\pi(\theta | \underline{x}) \propto \pi(\theta) L(\theta) = \left[\frac{1}{Be(\alpha, \beta)} \theta^{\alpha-1} (1-\theta)^{\beta-1} \right] \left[\theta^{\sum x_i} (1-\theta)^{n-\sum x_i} \right]$$

$$\propto \theta^{\alpha+\sum x_i-1} (1-\theta)^{n+\beta-\sum x_i-1}$$

$$\Rightarrow \theta | \underline{x} \sim Be(\alpha + \sum x_i, n + \beta - \sum x_i)$$

برای بدست آوردن تصمیم بیزی از قضیه زیر استفاده می‌کنیم.

قضیه ۱-۲: فرض کنید θ دارای توزیع پیشین $\pi(\theta)$ باشد و به شرط θ ، $\underline{X}$ دارای توزیعی از خانواده $\{f(\underline{x} | \theta) | \theta \in \Theta\}$ باشد. فرض کنید در مسئله تصمیم با تابع زیان غیرمنفی $L(\theta, \delta)$ شرایط زیر برقرار باشد.

الف- یک تابع تصمیم δ با مخاطره متناهی موجود باشد.

ب- برای هر $\underline{x} \in \mathcal{X}$ یک تابع تصمیم $\delta^\pi(\underline{x})$ موجود باشد که

$$r(\underline{x}, \delta(\underline{x})) = E[L(\theta, \delta(\underline{x})) | \underline{X} = \underline{x}] = \int_{\Theta} L(\theta, \delta(\underline{x})) d\pi(\theta | \underline{x})$$

را مینیمم کند (یعنی $r(\underline{x}, \delta^\pi) = \inf_{\delta \in D} r(\underline{x}, \delta)$).

در این صورت $\delta^\pi(\underline{x})$ یک تصمیم بیز است $(r(\underline{x}, \delta))$ را مخاطره پسین گویند).

اثبات: فرض کنید δ هر تصمیمی با مخاطره متناهی باشد. در این صورت چون تابع زیان نامنفی است.

$$E\{L(\theta, \delta(\underline{x})) | \underline{X} = \underline{x}\} < \infty \quad a.e.$$

پس

و در نتیجه طبق شرط (ب) داریم که

$$E\{L(\theta, \delta(\underline{x})) | \underline{X} = \underline{x}\} \geq E\{L(\theta, \delta^\pi(\underline{x})) | \underline{X} = \underline{x}\} \quad a.e.$$

$$\Rightarrow E[E\{L(\theta, \delta(\underline{x})) | \underline{X} = \underline{x}\}] \geq E[E\{L(\theta, \delta^\pi(\underline{x})) | \underline{X} = \underline{x}\}]$$

$$\Rightarrow E[L(\theta, \delta(\underline{x}))] \geq E[L(\theta, \delta^\pi(\underline{x}))]$$

که امید ریاضی نسبت به توزیع توام θ و $\underline{X}$ است پس $r(\pi, \delta) \geq r(\pi, \delta^\pi)$ یعنی $\delta^\pi(\underline{X})$ یک

تصمیم بیز است زیرا مخاطره بیز را مینیمم می‌کند. ■

با توجه به قضیه بالا برای بدست آوردن تصمیم بیزی بایستی مخاطره پسین را مینیمم کنیم. در برآورد نقطه‌ای، تصمیم بیزی بدست آمده را برآوردگر بیز^۱ می‌نامیم و در قضیه زیر تحت توابع زیان مختلف این برآوردگر را بدست می‌آوریم.

قضیه ۲-۲: تحت فرضیات قضیه ۱-۲ در برآورد نقطه‌ای پارامتر $\gamma(\theta)$ داریم که

الف- تحت تابع زیان درجه دوم $L(\gamma(\theta), \delta(\underline{x})) = [\delta(\underline{x}) - \gamma(\theta)]^2$ برآورد بیز عبارت است از:

$$\delta^\pi(\underline{x}) = E[\gamma(\theta) | \underline{x}]$$

ب- تحت تابع زیان قدر مطلق $L(\gamma(\theta), \delta(\underline{x})) = |\delta(\underline{x}) - \gamma(\theta)|$ برآورد بیز، میانه توزیع $\gamma(\theta) | \underline{x}$ است.

ج- تحت تابع زیان درجه دوم وزنی $L(\gamma(\theta), \delta(\underline{x})) = w(\theta)[\delta(\underline{x}) - \gamma(\theta)]^2$ برآورد بیز عبارت

$$\delta^\pi(\underline{x}) = \frac{E[w(\theta)\gamma(\theta) | \underline{x}]}{E[w(\theta) | \underline{x}]}$$

است از:

اثبات (ج):

$$\begin{aligned} r(\underline{x}, \delta(\underline{x})) &= \int_{\Theta} w(\theta)[\delta(\underline{x}) - \gamma(\theta)]^2 \pi(\theta | \underline{x}) d\theta \\ &= \delta^2(\underline{x}) E[w(\theta) | \underline{x}] - 2\delta(\underline{x}) E[w(\theta)\gamma(\theta) | \underline{x}] + E[w(\theta)\gamma^2(\theta) | \underline{x}] \\ \blacksquare \quad \frac{\partial r(\underline{x}, \delta(\underline{x}))}{\partial \delta(\underline{x})} &= 0 \Rightarrow \delta_\pi(\underline{x}) = \dots\dots \end{aligned}$$

مثال ۲-۲: در مثال ۱-۲ برآوردگر بیز پارامتر θ را تحت تابع زیان درجه دوم بدست آورید.

$$X_1, \dots, X_n \sim B(1, \theta) \quad \theta \sim Be(\alpha, \beta) \Rightarrow \theta | \underline{x} \sim Be(\alpha + \sum x_i, \beta + n - \sum x_i)$$

$$\delta^\pi(\underline{x}) = E[\theta | \underline{x}] = \frac{\alpha + \sum x_i}{\alpha + \beta + n} = \left(\frac{\alpha + \beta}{\alpha + \beta + n} \right) \frac{\alpha}{\alpha + \beta} + \left(\frac{n}{\alpha + \beta + n} \right) \bar{x}$$

چون $E(\theta) = \frac{\alpha}{\alpha + \beta}$ بنابراین برآوردگر بیز بصورت یک میانگین وزنی از میانگین توزیع پیشین

$\frac{\alpha}{\alpha + \beta}$ (برآورد θ قبل از مشاهده داده‌ها) و میانگین نمونه $\bar{x}$ (برآورد θ بوسیله مشاهدات) می‌باشد.

^۱ Bayes Estimator

توجه کنید که این برآوردگر ناریب نیست و در صورتی که $n \rightarrow \infty$ و یا $\alpha = \beta = 0$ انگاه برآوردگر $\delta^\pi(x)$ با برآوردگر UMVU یکی می‌شود.

البته حالت $\alpha = \beta = 0$ یک توزیع پیشین نمی‌باشد زیرا بایستی $\alpha > 0, \beta > 0$ باشد (حالت $\alpha \rightarrow 0$ و $\beta \rightarrow 0$ را بعداً بحث خواهیم کرد). ■

مثال ۲-۳: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی n تایی از توزیع $N(\theta, \sigma^2)$ باشند که در آن σ^2 یک مقدار معلوم است. اگر $\theta \sim N(\mu, \tau^2)$ و تابع زیان بفرم درجه ۲ باشد برآوردگر بیز پارامتر θ را بدست آورید.

حل:

$$\begin{aligned} \pi(\theta | x) &\propto \pi(\theta) L(\theta) = \frac{1}{\sqrt{2\pi\tau^2}} e^{-\frac{1}{2\tau^2}(\theta-\mu)^2} \left(\frac{1}{\sqrt{2\pi\sigma^2}} \right)^n e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n (X_i - \theta)^2} \\ &\propto e^{-\frac{\theta^2}{2\tau^2} + \frac{\mu}{\tau^2}\theta} e^{\frac{\theta}{\sigma^2} \sum x_i - \frac{n\theta^2}{2\sigma^2}} = e^{-\frac{\theta^2}{2} \left(\frac{1}{\tau^2} + \frac{n}{\sigma^2} \right) + \theta \left(\frac{\mu}{\tau^2} + \frac{n\bar{x}}{\sigma^2} \right)} \\ &= e^{-\frac{1}{2} \left(\frac{n}{\sigma^2} + \frac{1}{\tau^2} \right) \left[\theta^2 - 2\theta \frac{\frac{\mu}{\tau^2} + \frac{n\bar{x}}{\sigma^2}}{\frac{1}{\tau^2} + \frac{n}{\sigma^2}} \right]} = \exp \left\{ -\frac{1}{2} \left(\frac{n}{\sigma^2} + \frac{1}{\tau^2} \right) \left(\theta^2 - 2\theta \frac{\mu \left(\frac{\sigma^2}{n} \right) + \tau^2 \bar{x}}{\tau^2 + \frac{\sigma^2}{n}} \right) \right\} \\ &\propto \exp \left\{ -\frac{1}{2} \left(\frac{n}{\sigma^2} + \frac{1}{\tau^2} \right) \left(\theta - \frac{\mu \left(\frac{\sigma^2}{n} \right) + \tau^2 \bar{x}}{\tau^2 + \frac{\sigma^2}{n}} \right)^2 \right\} \\ &\therefore \theta | x \sim N \left(\frac{\mu \left(\frac{\sigma^2}{n} \right) + \tau^2 \bar{x}}{\frac{\sigma^2}{n} + \tau^2}, \frac{1}{\frac{n}{\sigma^2} + \frac{1}{\tau^2}} \right) \end{aligned}$$

$$\Rightarrow \delta^\pi(x) = E(\theta | x) = \frac{\frac{\sigma^2}{n}}{\frac{\sigma^2}{n} + \tau^2} \mu + \frac{\frac{\tau^2}{n}}{\frac{\sigma^2}{n} + \tau^2} \bar{x}$$

یعنی برآوردگر بیز بصورت یک میانگین وزنی از میانگین توزیع پیشین (μ) و میانگین نمونه $\bar{x}$ (برآورد عادی θ) می‌باشد و به سادگی دیده می‌شود که ناریب نیست. ■

مثال ۲-۴: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی n تایی از توزیع $U(\cdot, \theta)$ باشند اگر

$\theta \sim U(\cdot, 1)$ باشد و تابع زیان بصورت $L(\theta, \delta) = \frac{(\delta - \theta)^2}{\theta^2}$ باشد برآوردگر بیز پارامتر θ را بدست

آورید.

حل:

$$\begin{aligned} \pi(\theta | x) &\propto \pi(\theta) L(\theta) = I_{(\cdot, 1)}(\theta) \prod_{i=1}^n \frac{1}{\theta} I_{(\cdot, \theta)}(x_i) = I_{(\cdot, 1)}(\theta) \frac{1}{\theta^n} I_{(x_{(n)}, +\infty)}(\theta) \\ &= \frac{1}{\theta^n} I_{(x_{(n)}, 1)}(\theta) \Rightarrow \pi(\theta | x) = c(x) \frac{1}{\theta^n} I_{(x_{(n)}, 1)}(\theta) \\ \Rightarrow \delta^\pi(x) &= \frac{E[\theta w(\theta) | x]}{E[w(\theta) | x]} = \frac{E\left(\frac{1}{\theta} | x\right)}{E\left(\frac{1}{\theta^2} | x\right)} = \frac{\int_{x_{(n)}}^1 \frac{1}{\theta^{n+1}} d\theta}{\int_{x_{(n)}}^1 \frac{1}{\theta^{n+2}} d\theta} = \frac{-\frac{1}{n\theta^n} \Big|_{x_{(n)}}^1}{\frac{-1}{(n+1)\theta^{n+1}} \Big|_{x_{(n)}}^1} \\ &= \frac{\frac{1}{n} \left[-1 + \frac{1}{x_{(n)}^n}\right]}{\frac{1}{n+1} \left[-1 + \frac{1}{x_{(n)}^{n+1}}\right]} = \frac{n+1}{n} x_{(n)} \left[\frac{1 - x_{(n)}^n}{1 - x_{(n)}^{n+1}} \right] \quad \blacksquare \end{aligned}$$

نکاتی در مورد برآوردگر بیز:

نکته ۱: در رابطه با برآوردگرهای بیز این سوال مطرح می‌شود که چه خانواده از توزیع‌های پیشین را برای یک مسئله انتخاب کنیم. یکی از روش‌های انتخاب خانواده توزیع‌های پیشین، استفاده از خانواده توزیع‌های مزدوج است.

تعریف ۱-۲: فرض کنید F خانواده توزیع‌های $f(x|\theta)$ باشد یعنی $F = \{f(x|\theta) | \theta \in \Theta\}$. خانواده Π از توزیع‌های پیشین برای θ را یک خانواده مزدوج برای F گوئیم هرگاه برای هر $f \in F$ و هر $\pi \in \Pi$ و هر $x \in \mathcal{X}$ توزیع‌های پسین نیز در خانواده Π قرار گیرند. ■

برای مثال در مثال ۲-۲ خانواده $Be(\alpha, \beta)$ یک خانواده مزدوج برای $B(1, \theta)$ بود و در مثال ۲-۳ خانواده $N(\mu, \sigma^2)$ یک خانواده مزدوج برای $N(\theta, \sigma^2)$ بود.

مثال ۲-۵: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع نمایی با تابع چگالی زیر باشند $f(x|\theta) = \theta e^{-\theta x}$, $x > 0$. اگر $\theta \sim \Gamma(\alpha, \beta)$ باشد تحت تابع زیان درجه دوم برآوردگر بیز میانگین توزیع یعنی $\gamma(\theta) = \frac{1}{\theta}$ را بدست آورید.

حل:

$$\pi(\theta|x) \propto \pi(\theta)L(\theta) = \frac{1}{\beta^\alpha \Gamma(\alpha)} \theta^{\alpha-1} e^{-\theta/\beta} \theta^n e^{-n\theta\bar{x}} \propto \theta^{n+\alpha-1} e^{-(n\bar{x} + \frac{1}{\beta})\theta}$$

$$\Rightarrow \theta|x \sim \Gamma(n+\alpha, \frac{1}{n\bar{x} + \beta^{-1}})$$

$$\begin{aligned} \delta^\pi(x) &= E[\gamma(\theta)|x] = E\left[\frac{1}{\theta}|x\right] = \frac{(n\bar{x} + \beta^{-1})^{n+\alpha}}{\Gamma(n+\alpha)} \times \frac{\Gamma(n+\alpha-1)}{(n\bar{x} + \beta^{-1})^{n+\alpha-1}} \\ &= \frac{n\bar{x} + \beta^{-1}}{n+\alpha-1} = \left(\frac{\alpha-1}{n+\alpha-1}\right) \left(\frac{1}{\beta(\alpha-1)}\right) + \left(\frac{n}{n+\alpha-1}\right) \bar{x} \end{aligned}$$

بنابراین برآورد بیز، یک میانگین وزنی از میانگین نمونه $\bar{x}$ (برآوردگر معقول $\frac{1}{\theta}$) و میانگین $\frac{1}{\theta}$ در

توزیع پیشین $\left(\frac{1}{\beta(\alpha-1)}\right)$ می‌باشد. ■

مثال ۲-۶: فرض کنید $X_1, \dots, X_n$ یک نمونه تصادفی از توزیع $P(\theta)$ باشد اگر $\theta \sim \Gamma(\alpha, \beta)$ باشد برآورد بیز پارامتر θ تحت تابع زیان درجه دوم را بدست آورید.

$$\pi(\theta | \underline{x}) \propto \pi(\theta) L(\theta) = \frac{1}{\beta^\alpha \Gamma(\alpha)} \theta^{\alpha-1} e^{-\theta/\beta} \frac{e^{-n\theta} \theta^{n\bar{x}}}{\prod_{i=1}^n x_i!} \propto \theta^{n\bar{x}+\alpha-1} e^{-\theta(n+\frac{1}{\beta})}$$

$$\Rightarrow \theta | \underline{x} \sim \Gamma(n\bar{x} + \alpha, \frac{1}{n + \beta^{-1}}) \Rightarrow E(\theta | \underline{x}) = \frac{n\bar{x} + \alpha}{n + \beta^{-1}}$$

$$\therefore \delta^\pi(\underline{X}) = \frac{\beta^{-1}}{n + \beta^{-1}} (\alpha\beta) + \frac{n}{n + \beta^{-1}} \bar{X}$$

برآوردگر بیز یک میانگین وزنی از میانگین نمونه $\bar{X}$ (برآوردگر معقول θ) و میانگین توزیع پیشین، $\alpha\beta$ است. ■

نکته ۲: در مثال‌های قبل دیدیم که برآوردگرهای بیز نارایب نمی‌باشند. این موضوع در قضیه زیر اثبات می‌شود.

قضیه ۲-۳: فرض کنید θ دارای توزیع پیشین $\pi(\theta)$ باشد و $f(\underline{x} | \theta)$ توزیع شرطی $\underline{X}$ به شرط θ باشد. در برآورد پارامتر $\gamma(\theta)$ تحت تابع زیان درجه دوم، هیچ برآوردگر نارایب $\delta(\underline{X})$ نمی‌تواند

$$r(\pi, \delta) = E_{\underline{X}, \theta} [(\delta(\underline{X}) - \theta)^2] = 0.$$

یک برآوردگر بیز باشد، بجز اینکه

این امید ریاضی نسبت به چگالی توام $\underline{X}$ و θ است.

اثبات: فرض کنید δ یک برآوردگر بیز و نارایب برای $\gamma(\theta)$ باشد. در اینصورت چون δ برآوردگر

$$E[\delta(\underline{X}) | \theta] = \gamma(\theta)$$

نارایب $\gamma(\theta)$ است پس

$$E[\gamma(\theta) | \underline{X}] = \delta(\underline{X})$$

و چون δ برآوردگر بیز $\gamma(\theta)$ است پس

همچنین محاطره بیز عبارت است از

$$r(\Pi, \delta) = E[(\delta(\underline{X}) - \gamma(\theta))^2] = E[\delta^2(\underline{X})] - 2E[(\gamma(\theta)\delta(\underline{X}))] + E[\gamma^2(\theta)] \quad (1)$$

از طرفی

$$E[\gamma(\theta)\delta(\underline{X})] = E[E(\gamma(\theta)\delta(\underline{X}) | \underline{X})] = E[\delta(\underline{X})E(\gamma(\theta) | \underline{X})] = E[\delta^2(\underline{X})] \quad (2)$$

$$E[\gamma(\theta)\delta(\underline{X})] = E[E(\gamma(\theta)\delta(\underline{X}) | \theta)] = E[\gamma(\theta)E(\delta(\underline{X}) | \theta)] = E[\gamma^2(\theta)] \quad (3)$$

و در نتیجه از روابط (۱) و (۲) و (۳) نتیجه می شود که $r(\pi, \delta) = 0$ است. ■

مثال ۷-۲: در مثال ۳-۲ دیدیم که $\delta^\pi(X) = \frac{\sigma^2/n}{\sigma^2/n + \tau^2} \mu + \frac{\tau^2}{\sigma^2/n + \tau^2} \bar{X}$ برآوردگر بیز θ در توزیع $N(\theta, \sigma^2)$ بود. نشان دهید که $\bar{X}$ برآوردگر UMVU پارامتر θ است نمی تواند برآوردگر بیز باشد.

حل: اگر $\bar{X}$ بخواهد برآوردگر ناریب بیز باشد بایستی $E_{X, \theta}((\bar{X} - \theta)^2) = 0$ باشد. اما

$$E_{X, \theta}((\bar{X} - \theta)^2) = E[E(\bar{X} - \theta)^2 | \theta] = E\left[\frac{\sigma^2}{n} | \theta\right] = \frac{\sigma^2}{n} \neq 0.$$

پس $\bar{X}$ نمی تواند برآوردگر بیز باشد.

مثال ۸-۲: در مثال ۲-۲ دیدیم که در خانواده توزیع های $B(1, \theta)$ تحت توزیع پیشین $Be(\alpha, \beta)$ و

تابع زیان درجه دوم برآوردگر بیز برابر $\delta^\pi(X) = \frac{\alpha + n\bar{X}}{\alpha + \beta + n}$ بود. تحت چه توزیع پیشینی $\bar{X}$ یک برآوردگر بیز برای θ است.

حل: $\bar{X}$ در صورتی برآوردگر بیز θ است که

$$\begin{aligned} 0 &= E_{X, \theta}[(\bar{X} - \theta)^2] = E[E(\bar{X} - \theta)^2 | \theta] = E[Var(\bar{X}) | \theta] = E\left[\frac{\theta(1-\theta)}{n}\right] \\ \Rightarrow E\left[\frac{\theta(1-\theta)}{n}\right] &= 0 \stackrel{0 \leq \theta \leq 1}{\Rightarrow} P(\theta(1-\theta) = 0) = 1 \Rightarrow P(\theta = 0 \text{ or } \theta = 1) = 1 \end{aligned}$$

$$\Rightarrow P(\theta = 0) + P(\theta = 1) = 1$$

بنابراین در صورتی $\bar{X}$ برآوردگر بیز است که θ دارای توزیع پیشین برنولی باشد و برای چنین توزیع پیشینی داریم که

$$\delta^\pi(0) = 0 \quad \delta^\pi(1) = 1 \quad (\text{چون } 1 \text{ یا } 0 \text{ است } \theta). \quad (4)$$

و هر برآوردگری که در شرط بالا صدق کند برآوردگر بیز است. مثلاً $\delta^\pi(X) = \bar{X}$ در این شرایط صدق می کند. البته این چگالی پیشین غیرمعقول است.

نکته ۳: در مثال‌های بالا اکثر برآوردهای بیز یکتا بودند و تنها در مثال ۲-۸ با تابع چگالی پیشین غیرمعقول $P(\theta=0)+P(\theta=1)=1$ دیدیم که برآوردهای بیز یکتا نبود. شرط کلی برای یکتا بودن برآوردهای بیز در قضیه زیر آورده شده است.

قضیه ۲-۴: اگر تابع زیان $L(\theta, \delta)$ درجه دوم یا بطور کلی اکیدا محدب باشد و P کلاس کلیه توزیع‌های P_θ باشد، آنگاه برآورد بیز δ^π یکتا ($a.e. P$) است هرگاه

الف- مخاطره بیز آن نسبت به π متناهی باشد $r(\pi, \delta^\pi) < \infty$

ب- اگر Q توزیع کناری X باشد یعنی $Q(A) = \int P_\theta(X \in A) d\Pi(\theta)$

آنگاه $a.e. Q \Rightarrow a.e. P$

(به عبارتی $Q(N)=0 \Rightarrow P(N)=0$ و یا $P(X \in A)=0 \Rightarrow P(X \in A | \theta)=0$)

اثبات: برای تابع زیان درجه دوم از قضیه ۲-۲ نتیجه می‌شود که برآوردهای بیز با مخاطره متناهی بایستی در رابطه $\delta^\pi(X) = E[\gamma(\theta) | X]$ صدق کند بجز روی مجموعه‌ای از نقاط X مانند N که $Q(N)=0$ و این بدان معنی است که برآوردهای بیز یکتا است. در حالت تابع زیان اکیدا محدب نیز نتیجه به صورت فوق است زیرا این نوع توابع دارای نقطه مینیمم یکتا می‌باشند. ■

مثال ۲-۹: در مثال ۲-۸ دیدیم که تحت چگالی پیشین $P(\theta=0)+P(\theta=1)=1$ برآوردهای بیز پارامتر θ در خانواده $B(1, \theta)$ در شرط $\delta^\pi(0)=0$ و $\delta^\pi(1)=1$ صدق می‌کرد و یکتا نبود. علت یکتا نبودن آن برقرار نبودن شرط (ب) قضیه ۲-۴ است زیرا

$$P(X_1=1, \dots, X_n=1 | \theta) = \prod_{i=1}^n \theta = \theta^n \neq 0.$$

$$P(X_1=1, \dots, X_n=1) = \sum_{t=0}^1 P(X_1=1, \dots, X_n=1 | \theta=t) P(\theta=t)$$

$$= 0 \times P(\theta=0) + 1 \times P(\theta=1) = 1.$$

$$\left(f(x | \theta) = \theta^{\sum x_i} (1-\theta)^{n-\sum x_i} = 0 \text{ if } \theta=1 \right)$$

بنابراین $P(X_1=1, \dots, X_n=1) = 1 \neq 0 = P(X_1=1, \dots, X_n=1 | \theta=0)$ ■

نکته ۴: در بعضی از مسایل ممکن است برای بدست آوردن برآوردگر بیز، تابع چگالی پسین را بطور کامل محاسبه کنیم. به مثال زیر توجه کنید.

مثال ۲-۱۰: فرض کنید X دارای توزیع $U(\cdot, \theta)$ باشد و θ دارای چگالی پیشین بصورت

$$\pi(\theta) = \begin{cases} \theta e^{-\theta} & \theta > 0 \\ 0 & \text{otherwise} \end{cases}$$

برآوردگرهای بیز را بدست آورید.

حل:

$$\pi(\theta | x) = \frac{\pi(\theta) f(x | \theta)}{m(x)} = c(x) \pi(\theta) f(x | \theta) = c(x) \theta e^{-\theta} \frac{1}{\theta} I_{(x, +\infty)}(\theta)$$

$$1 = \int \pi(\theta | x) d\theta \Rightarrow 1 = \int_x^{+\infty} c(x) e^{-\theta} d\theta = c(x) [-e^{-\theta}]_x^{+\infty} = c(x) e^{-x}$$

$$\Rightarrow c(x) = e^x \quad \therefore \pi(\theta | x) = e^{-(\theta-x)} I_{(x, +\infty)}(\theta)$$

$$L(\theta, \delta) = (\delta - \theta)^2 \Rightarrow \delta^\pi(x) = E(\theta | x) = x + 1 \Rightarrow \delta^\pi(X) = X + 1 \quad (\text{الف})$$

(ب)

$$L(\theta, \delta) = |\delta - \theta| \Rightarrow \int_x^m \pi(\theta | x) d\theta = \frac{1}{2} \Rightarrow \int_x^m e^{-(\theta-x)} d\theta = \frac{1}{2}$$

$$\Rightarrow -e^{-(\theta-x)} \Big|_x^m = \frac{1}{2} \Rightarrow 1 - e^{-(m-x)} = \frac{1}{2} \Rightarrow m - x = \log 2 \Rightarrow m = x + \log 2$$

$$\Rightarrow \delta^\pi(X) = X + \log 2$$

■

نکته ۵: اگر در مسئله‌ای یک آماره بسنده T موجود باشد آنگاه توزیع پسین به آماره T وابسته است و در حقیقت $\pi(\theta | x) = \pi(\theta | t)$. برای اثبات این مطلب توجه کنید که از قضیه تجزیه به عوامل

$$L(\theta) = f(x | \theta) = g(t | \theta) h(x) \quad \text{داریم که}$$

$$\pi(\theta | x) = \frac{f(x | \theta) \pi(\theta)}{\int f(x | \theta') d\pi(\theta')} = \frac{g(t | \theta) \pi(\theta)}{\int g(t | \theta') d\pi(\theta')} = \pi(\theta | t) \quad \text{پس}$$

یعنی توزیع پسین $\pi(\theta | t)$ بستگی به x از طریق t دارد و بنابراین محاسبه توزیع پسین براساس x

یا t یکسان است. برای مثال، در مثال ۲-۳ داریم که

$$X_1, \dots, X_n \text{ i.i.d } N(\mu, \sigma^2), T = \bar{X} \sim N\left(\theta, \frac{\sigma^2}{n}\right), \theta \sim N(\mu, \tau^2)$$

$$\pi(\theta|t) \propto f(t|\theta)\pi(\theta) \propto e^{-\frac{n}{2\sigma^2}(\bar{x}-\theta)^2} e^{-\frac{1}{2\tau^2}(\theta-\mu)^2} \propto e^{-\frac{\theta^2}{2}\left(\frac{1}{\tau^2}+\frac{n}{\sigma^2}\right)+\theta\left(\frac{\mu}{\tau^2}+\frac{n\bar{x}}{\sigma^2}\right)}$$

که همان تابع به دست آمده در مثال ۲-۳ است بنابراین

$$\theta|T=t \sim N\left(\frac{\mu\left(\frac{\sigma^2}{n}\right)+\tau^2\bar{x}}{\frac{\sigma^2}{n}+\tau^2}, \frac{1}{\frac{n}{\sigma^2}+\frac{1}{\tau^2}}\right) \quad \blacksquare$$

توزیع‌های پیشین ناآگاهی بخش^۱:

همانگونه که قبلاً گفته شد، توزیع پیشین $\pi(\theta)$ براساس اعتقادات و تجربیات قبلی پژوهشگر تشکیل می‌گردد. حال اگر هیچ توزیع پیشینی در دسترس نباشد، می‌توان از توزیع‌های پیشین ناآگاهی بخش استفاده کرد. برای مثال اگر در یک مسئله آزمون فرض ساده در مقابل احتمالات یکسان $\frac{1}{p}$ را به فرض صفر و مقابل نسبت دهیم آنگاه این توزیع پیشین، ناآگاهی بخش است.

یکی از روش‌های بدست آوردن توزیع پیشین ناآگاهی بخش توسط جفریز^۲ (۱۹۶۱) ارایه شده است، که در این روش چگالی پیشین را به صورت $\pi(\theta) \propto \sqrt{|I(\theta)|}$ در نظر می‌گیرد که در آن $|I(\theta)|$ دترمینان ماتریس اطلاع فشر می‌باشد.

مثال ۲-۱۱: اگر $X \sim N(\theta, 1)$ آنگاه توزیع پیشین جفریز را برای θ بدست آورید.

$$\blacksquare \quad f(x|\theta) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-\theta)^2} \Rightarrow I(\theta) = 1 \Rightarrow \pi(\theta) = 1 \quad -\infty < \theta < +\infty$$

توجه کنید که توزیع پیشین به دست آمده در مثال ۲-۱۱ یک تابع چگالی نیست. اگر در مسئله‌ای اندازه پیشین در شرط $\int d\Pi(\theta) = \infty$ صدق کند آنگاه پیشین حاصل را ناسره^۳ گویند و توزیع پیشین

^۱ Noninformative Priors

^۲ Jeffreys

^۳ Improper Prior

که در شرط $\int d\Pi(\theta) = 1$ صدق کند را توزیع پیشین سره^۱ گویند. بنابراین توزیع پیشین بدست آمده در مثال ۱۱-۲ ناسره است.

بایستی توجه کرد که اگر توزیع پیشینی ناسره باشد آنگاه ممکن است مخاطره براساس این توزیع پیشین متناهی باشد و برآوردگر بیز آن قابل محاسبه باشد. برای مثال، در مثال ۱۱-۲ داریم که

$$\pi(\theta|x) \propto f(x|\theta)\pi(\theta) = f(x|\theta) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-\theta)^2}$$

پس توزیع پسین سره است و $\theta|x \sim N(x, 1)$ در نتیجه $\delta^\pi(X) = E(\theta|X) = X$. برآوردگر بیز حاصل را برآوردگر بیز تعمیم یافته گویند.

تعریف ۲-۲: یک برآوردگر $\delta_G^\pi(x)$ را برآوردگر بیز تعمیم یافته^۲ نسبت به اندازه $\pi(\theta)$ گویند اگر مخاطره پسین $E[L(\theta, \delta(x)) | X = x]$ در $\delta = \delta_G^\pi$ برای هر x مینیمم شود. ■

مثال ۲-۲: در مثال ۲-۲ توزیع پیشین جفریز را بدست آورده و سپس براساس توزیع پیشین

$$\pi(\theta) = \frac{1}{\theta(1-\theta)}, \quad \text{برآوردگر بیز (تعمیم یافته) } \theta \text{ تحت تابع زیان درجه دوم را بدست آورید.}$$

حل:

$$X_i \sim B(1, \theta) \quad i = 1, 2, \dots, n \quad f_P(x) = \theta^x (1-\theta)^{1-x} \Rightarrow I(\theta) = \frac{n}{\theta(1-\theta)}$$

$$\pi(\theta) \propto \sqrt{|I(\theta)|} \Rightarrow \Pi(\theta) \propto \frac{1}{\theta^{1/2} (1-\theta)^{1/2}} \quad \text{توزیع سره است}$$

$$\text{براساس توزیع پیشین ناسره داریم } \pi(\theta) = \frac{1}{\theta(1-\theta)}$$

$$\pi(\theta|x) \propto \theta^{n\bar{x}} (1-\theta)^{n-n\bar{x}} \frac{1}{\theta(1-\theta)} = \theta^{n\bar{x}-1} (1-\theta)^{n-n\bar{x}-1} \quad 0 < \theta < 1$$

بنابراین اگر $\bar{x} \neq 0$ یا $\bar{x} \neq 1$ آنگاه توزیع پسین سره است و $\theta|x \sim Be(n\bar{x}, n-n\bar{x})$

$$\delta_G^\pi(x) = E(\theta|x) = \frac{n\bar{x}}{n} = \bar{x} \quad \text{برآوردگر بیز تعمیم یافته } \theta \text{ عبارت است از}$$

^۱ Proper

^۲ Generalized Bayes

حال اگر $\bar{x} = 0$ باشد آنگاه طبق قضیه ۲-۱ برآوردگر بیز آن مقداری از a است که انتگرال زیر را

$$E[(\theta - a)^2 | x] = c(x) \int_0^1 (\theta - a)^2 \theta^{-1} (1 - \theta)^{n-1} d\theta \quad \text{مینیمم کند}$$

انتگرال فوق در صورتی همگرا است که $a = 0$ باشد (بخاطر وجود θ^{-1}) پس در این حالت برآورد بیز تعمیم یافته عبارت است از $\delta_G^\pi(x) = a = 0 = \bar{x}$ به همین ترتیب برای $\bar{x} = 1$ می توان نشان داد که برآورد بیز تعمیم یافته عبارت است از $\delta_G^\pi(x) = a = 1 = \bar{x}$ بنابراین در هر حالت برآورد بیز تعمیم یافته برابر $\delta_G^\pi(x) = \bar{x}$ است.

همچنین توجه کنید که در مثال ۲-۲ تحت توزیع پیشین سره $Be(\alpha, \beta)$ برآورد بیز θ عبارت بود از

$$\delta^\pi(x) = \frac{\alpha + n\bar{x}}{\alpha + \beta + n}$$

حال اگر $\alpha \rightarrow 0$ و $\beta \rightarrow 0$ آنگاه $\delta^\pi(x) \rightarrow \bar{x}$ یعنی $\delta^\pi(x)$ به سمت برآورد بیز تعمیم یافته θ تحت توزیع پیشین جفریز میل می کند. ■

تعریف ۲-۳: یک برآوردگر $\delta_L(x)$ را برآوردگر بیز حدی^۱ گویند اگر دنباله ای از توزیع های پیشین π_m و برآوردگرهای بیز $\delta^{\pi_m}(x)$ موجود باشد بطوری که

$$\delta^{\pi_m}(x) \rightarrow \delta(x) \quad (a.e.) \quad \text{as } m \rightarrow \infty \quad \blacksquare$$

بنابراین برآوردگر $\delta_L(x) = \bar{x}$ در توزیع $B(1, \theta)$ یک برآوردگر بیز حدی برای θ است.

مثال ۲-۱۳: اگر $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\theta, \sigma^2)$ باشند که σ^2 معلوم است، آنگاه تحت توزیع پیشین ناسره $\pi(\theta) = 1$ (جفریز) برآوردگر بیز تعمیم یافته θ برابر $\delta_G^\pi(x) = \bar{x}$ است و

$$\delta^\pi(x) = \frac{\mu \left(\frac{\sigma^2}{n} \right) + \tau^2 \bar{x}}{\frac{\sigma^2}{n} + \tau^2} \quad \text{تحت توزیع پیشین سره } N(\mu, \tau^2) \text{ برآوردگر بیز}$$

برآوردگر $\delta_L(x) = \bar{x}$ برآوردگر بیز حدی θ موقعی که $\tau \rightarrow \infty$ می باشد. ■

^۱Limit of Bayes Estimator

توجه: از دیدگاه بیز، برآوردگرهای بیز حدی مطلوبتر از برآوردگرهای بیز تعمیم یافته هستند. زیرا حد دنباله‌های بیز بایستی به یک برآوردگر بیز (سره) نزدیک باشد، اما برآوردگر بیز تعمیم یافته ممکن است به هیچ برآوردگر بیزی (سره) نزدیک نباشد (مسئله ۲-۱۵ را ببینید).

مثال ۲-۱۴: فرض کنید $X \sim N(\mu, \sigma^2)$ که μ و σ^2 نامعلوم هستند. توزیع پیشین جفریز را برای $\theta = (\mu, \sigma)$ به دست آورید.

$$I_{ij}(\theta) = -E_{\theta} \left[\frac{\partial^2}{\partial \theta_i \partial \theta_j} \log f(x | \theta) \right]$$

$$\log f(x | \theta) = -\frac{1}{2} \log(2\pi) - \log \sigma - \frac{(X - \mu)^2}{2\sigma^2}$$

$$I(\theta) = -E_{\theta} \begin{bmatrix} -\frac{1}{\sigma^2} & \frac{2(\mu - X)}{\sigma^3} \\ \frac{2(\mu - X)}{\sigma^3} & -\frac{1}{\sigma^2} - \frac{2(X - \mu)^2}{\sigma^4} \end{bmatrix} = \begin{bmatrix} \frac{1}{\sigma^2} & 0 \\ 0 & \frac{2}{\sigma^2} \end{bmatrix}$$

$$\pi(\theta) \propto |I(\theta)|^{-\frac{1}{2}} = \left| \frac{2}{\sigma^4} \right|^{-\frac{1}{2}} = \frac{\sqrt{2}}{\sigma^2} \Rightarrow \pi(\theta) \propto \frac{1}{\sigma^2}$$

■ که توزیعی ناسره است.

مثال ۲-۱۵: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(0, \sigma^2)$ باشند. برآوردگر بیز σ^2

را براساس توابع زیان $L(\sigma^2, \delta) = \frac{(\delta - \sigma^2)^2}{\sigma^4}$ (B) و $L(\sigma^2, \delta) = (\delta - \sigma^2)$ (A) بدست آورید. اگر

توزیع پیشین جفریز را بکار ببریم، برآوردگر بیز σ^2 را تحت توابع زیان فوق به دست آورید.
حل:

$$f(x | \sigma^2) = (2\pi\sigma^2)^{-\frac{n}{2}} e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n x_i^2} = c\tau^{\frac{n}{2}} e^{-\tau y} \quad \text{الف-}$$

که در آن $\tau = \frac{1}{2\sigma^2}$ و $y = \sum_{i=1}^n x_i^2$ قرار می‌دهیم

$$\frac{1}{2\sigma^2} = \tau \sim \Gamma(g, \frac{1}{\alpha})$$

$$E\left(\frac{1}{\tau^r}\right) = \frac{\alpha^r}{(g-1)(g-2)}, \quad E\left(\frac{1}{\tau}\right) = \frac{\alpha}{g-1}, \quad E(\tau^r) = \frac{g(g+1)}{\alpha^r}, \quad E(\tau) = \frac{g}{\alpha} \quad \text{پس}$$

$$\pi(\tau | x) \propto \tau^{\frac{n}{r}} e^{-\tau y} \tau^{g-1} e^{-\alpha \tau} = \tau^{g+\frac{n}{r}-1} e^{-(\alpha+y)\tau}$$

بنابراین

$$\Rightarrow \tau | x \sim \Gamma\left(g + \frac{n}{r}, \frac{1}{\alpha+y}\right)$$

$$A) \Rightarrow \delta^\pi(x) = E(\sigma^r | x) = E\left(\frac{1}{\tau^r} | x\right) = \frac{1}{r} \frac{\alpha+y}{g + \frac{n}{r} - 1} = \frac{\alpha+y}{n+rg-2}$$

$$B) \Rightarrow \delta^\pi(x) = \frac{E(\sigma^r \frac{1}{\sigma^r} | x)}{E(\frac{1}{\sigma^r} | x)} = \frac{E(r\tau | x)}{E(r\tau^r | x)} = \frac{1}{r} \frac{(g + \frac{n}{r}) / (\alpha+y)}{(g + \frac{n}{r})(g + \frac{n}{r} + 1) / (\alpha+y)^r}$$

$$= \frac{\alpha+y}{n+rg-2}$$

$$I(\tau) = -E\left[\frac{\partial^r}{\partial \tau^r} \left(\frac{n}{r} \ln \tau - \tau y\right)\right] = -E\left[-\frac{n}{r\tau^r}\right] = \frac{n}{r\tau^r} \quad \text{ب-}$$

$$\pi(\tau) \propto \sqrt{|I(\tau)|} \propto \frac{1}{\tau} \Rightarrow \quad \pi(\tau) = \frac{1}{\tau} \quad \tau > 0.$$

که یک پیشین ناسره است.

$$\pi(\tau | x) \propto \tau^{\frac{n}{r}-1} e^{-\tau y} \Rightarrow \tau | x \sim \Gamma\left(\frac{n}{r}, \frac{1}{y}\right)$$

$$A) \Rightarrow \delta_G^\pi(x) = E(\sigma^r | x) = E\left(\frac{1}{\tau^r} | x\right) = \frac{1}{r} \frac{y}{\frac{n}{r} - 1} = \frac{y}{n-2}$$

$$B) \Rightarrow \delta_G^\pi(x) = \frac{E(\sigma^r \frac{1}{\sigma^r} | x)}{E(\frac{1}{\sigma^r} | x)} = \frac{E(r\tau | x)}{E(r\tau^r | x)} = \frac{1}{r} \frac{(n/r) / y}{(\frac{n}{r})(\frac{n}{r} + 1) / y^r} = \frac{y}{n+2}$$

توجه کنید که برآوردهای بیز تعمیم یافته در این حالت همان برآوردهای بیز حدی با $\alpha \rightarrow 0$ و

■ $g \rightarrow 0$ می باشند.

مثال ۲-۱۶: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\theta, \sigma^2)$ باشند که در آن θ و

σ^2 هر دو نامعلوم هستند. اگر $\tau = \frac{1}{\tau \sigma^2} \sim \Gamma(\alpha, g)$ و $\pi(\theta) = 1; \theta \in R$ ، برآوردگر بیز σ^2 را تحت

توابع زیان A و B مثال قبل و برآوردگر بیز θ را تحت تابع زیان درجه دوم به دست آورید.

$$f(x_i | \theta, \sigma^2) = (\tau \pi \sigma^2)^{-\frac{n}{2}} e^{-\frac{1}{2} \sum_{i=1}^n (x_i - \theta)^2} \propto \tau^{\frac{n}{2}} e^{-\tau [\sum (x_i - \bar{x})^2 + n(\bar{x} - \theta)^2]} \quad \text{حل:}$$

$$\pi(\theta, \tau) \propto \tau^{g-1} e^{-\alpha \tau} \quad \text{اگر } Z = \sum_{i=1}^n (x_i - \bar{x})^2 \text{ آنگاه}$$

$$\therefore \pi(\theta, \tau | x_i) \propto \tau^{\frac{n}{2}} e^{-\tau(Z + n(\bar{x} - \theta)^2)} \tau^{g-1} e^{-\alpha \tau} = \tau^{\frac{n}{2} + g - 1} e^{-\tau[\alpha + Z + n(\bar{x} - \theta)^2]}$$

$$\Rightarrow \pi(\theta, \tau | x_i) = C \tau^{\frac{n}{2} + g - 1} e^{-\tau[\alpha + Z + n(\bar{x} - \theta)^2]}$$

$$\pi(\tau | x_i) = \int_{-\infty}^{+\infty} \pi(\theta, \tau | x_i) d\theta = C^* \tau^{\frac{n-1}{2} + g - 1} e^{-\tau(\alpha + Z)}$$

$$\Rightarrow \tau | x_i \sim \Gamma\left(\frac{n-1}{2} + g, \frac{1}{\alpha + Z}\right)$$

$$(A) \Rightarrow \delta_G^\pi(x_i) = E[\sigma^2 | x_i] = E\left(\frac{1}{\tau} | x_i\right) = \frac{1}{\frac{n-1}{2} + g - 1} \frac{\alpha + Z}{1} = \frac{\alpha + Z}{n + 2g - 2}$$

$$(B) \Rightarrow \delta_G^\pi(x_i) = \frac{E(\tau | x_i)}{E(\tau^2 | x_i)} = \frac{1}{\frac{n-1}{2} + g} \frac{\left(\frac{n-1}{2} + g\right) / (\alpha + Z)}{\left(\frac{n-1}{2} + g\right) \left(\frac{n-1}{2} + g + 1\right) / (\alpha + Z)^2} = \frac{\alpha + Z}{n + 2g + 1}$$

$$\pi(\theta | x_i) = \int_{-\infty}^{+\infty} \pi(\theta, \tau | x_i) d\tau = C \frac{\Gamma\left(\frac{n}{2} + g\right)}{\left[\alpha + Z + n(\bar{x} - \theta)^2\right]^{\frac{n}{2} + g}}$$

این توزیع حول $\bar{x}$ متقارن است پس $\delta_G^\pi(x_i) = E(\theta | x_i) = \bar{x}$ ■

فصل ۵

مینیماکس بودن

و مجاز بودن

بخش ۱: تصمیم مینیماکس^۱

در فصل قبل تصمیم بیز را از طریق مینیم کردن مخاطره بیز (که همان میانگین تابع مخاطره نسبت به یک توزیع پیشین θ یعنی $\int_{\Theta} R(\theta, \delta) d\Pi(\theta)$ بود) بدست آوردیم.

مزیتی که روش بیز داشت این بود که اطلاعات تابع مخاطره در یک عدد (مخاطره بیز) خلاصه می شد و دو تصمیم بوسیله مقایسه این مخاطره بیزشان با یکدیگر مقایسه می شدند. در این فصل روش دیگری برای بدست آوردن یک تصمیم بهینه در نظر می گیریم که در این روش آن تصمیمی را انتخاب می کنیم که ماکسیمم (سوپریمم) تابع مخاطره نسبت به کلیه مقادیر فضای پارامتری Θ را در بین کلیه توابع تصمیم مینیمم کند.

تعریف ۱-۱: فرض کنید که D مجموعه کلیه توابع تصمیم و $R(\theta, \delta)$ تابع مخاطره باشد. تصمیم

$\delta_m \in D$ را تصمیم مینیماکس گویند هرگاه برای هر $\delta \in D$ داشته باشیم که

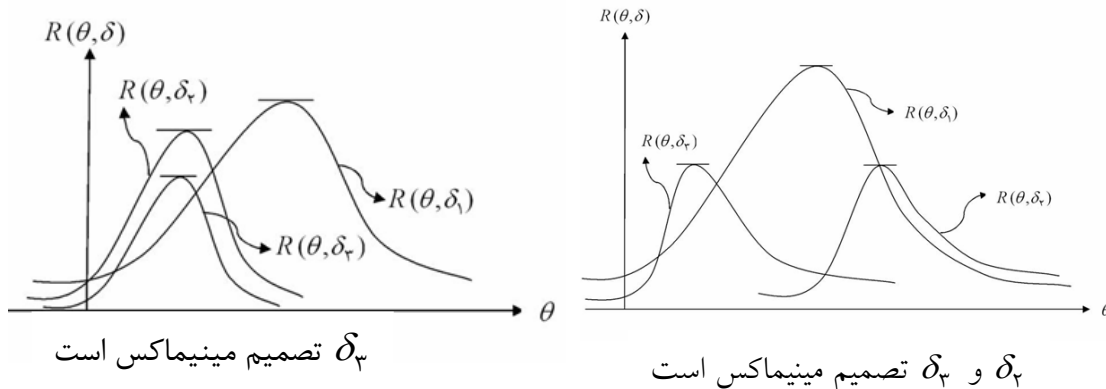
$$\sup_{\theta \in \Theta} R(\theta, \delta_m) \leq \sup_{\theta \in \Theta} R(\theta, \delta)$$

و یا به عبارتی $\delta_m \in D$ را تصمیم مینیماکس گویند هرگاه

$$\sup_{\theta \in \Theta} R(\theta, \delta_m) = \inf_{\delta \in D} \sup_{\theta \in \Theta} R(\theta, \delta)$$

$$\left(\max_{\theta \in \Theta} R(\theta, \delta_m) = \min_{\delta \in D} \max_{\theta \in \Theta} R(\theta, \delta) \right)$$

به شکل های زیر توجه کنید.



^۱ Minimax Decision

با توجه به شکل‌های بالا تصمیم مینیماکس ممکن است یکتا نباشد. تصمیم مینیماکس آن تصمیمی است که در بدترین حالت تابع مخاطره $R(\theta, \delta)$ نسبت به θ ، بهترین حالت را انتخاب می‌کند و در حقیقت یک حالت محافظه کارانه را برای انتخاب بهترین تصمیم در پیش می‌گیرد. با این وجود در بیشتر موارد تصمیم‌های معقولانه‌ای را بدست می‌دهد. مثلاً در بعضی از مسائل برآورد نقطه‌ای برآورد مینیماکس همان برآورد پایا و UMVU نیز می‌باشد.

تصمیم‌های مینیماکس وابستگی زیادی به تصمیم‌های بیز دارند. یک تصمیم مینیماکس بوسیله مینیم کردن، ماکسیمم مخاطره می‌خواهد بهترین تصمیم را در بدترین حالت انتخاب کند. بنابراین انتظار می‌رود که تصمیم مینیماکس همان تصمیم بیز با بدترین حالت توزیع پیشین باشد. چنین توزیعی را نامساعدترین توزیع^۱ گویند. به عبارت دقیقتر اینکه اگر π یک توزیع پیشین و δ^π تصمیم بیز نسبت به این توزیع پیشین و $r(\pi, \delta^\pi)$ مخاطره بیز آن باشد، در این صورت توزیع π را نامساعدترین گویند اگر برای هر توزیع پیشین دیگر π' با تصمیم بیز $\delta^{\pi'}$ داشته باشیم که $r(\pi, \delta^\pi) \geq r(\pi', \delta^{\pi'})$. در قضیه زیر شرط مینیماکس بودن یک تصمیم بیز را می‌آوریم.

قضیه ۱-۱: فرض کنید π یک توزیع پیشین روی Θ باشد و δ^π تصمیم بیز نسبت به این توزیع پیشین باشد. اگر تابع مخاطره تصمیم δ^π در شرط

$$r(\pi, \delta^\pi) = \int R(\theta, \delta^\pi) d\Pi(\theta) = \sup_{\theta \in \Theta} R(\theta, \delta^\pi) \quad (۱)$$

صدق کند آنگاه

الف- δ^π تصمیم مینیماکس است

ب- اگر δ^π تصمیم بیز یکتا نسبت به توزیع پیشین π باشد، آنگاه δ^π تصمیم مینیماکس یکتا است.

ج- توزیع π نامساعدترین توزیع است

اثبات: الف-

$$\sup_{\theta \in \Theta} R(\theta, \delta^\pi) \stackrel{(۱)}{=} \int R(\theta, \delta^\pi) d\Pi(\theta) = r(\pi, \delta^\pi) \leq r(\pi, \delta) = \int R(\theta, \delta) d\Pi(\theta)$$

^۱ Least Favorable Distribution

$$\leq \int \left[\sup_{\theta \in \Theta} R(\theta, \delta) \right] d\Pi(\theta) = \sup_{\theta \in \Theta} R(\theta, \delta)$$

بنابراین δ^π تصمیم مینیماکس است

ب- چون δ^π تصمیم بیز یکتا است پس $\forall \delta \in D$ ، $r(\pi, \delta^\pi) < r(\pi, \delta)$ و با در نظر گرفتن $<$ بجای $\leq$ در اثبات (الف) نتیجه می شود که تصمیم δ^π یک تصمیم مینیماکس یکتا است.

ج- فرض کنید π' هر توزیع پیشین دیگر با تصمیم بیز $\delta^{\pi'}$ باشد در این صورت

$$\begin{aligned} r(\pi', \delta^{\pi'}) &= \int R(\theta, \delta^{\pi'}) d\Pi'(\theta) \stackrel{*}{\leq} \int R(\theta, \delta^\pi) d\Pi'(\theta) \leq \sup R(\theta, \delta^\pi) \\ &\stackrel{(1)}{=} \int R(\theta, \delta^\pi) d\Pi(\theta) = r(\pi, \delta^\pi) \end{aligned}$$

که در آن $*$ از این واقعیت ناشی می گردد که $\delta^{\pi'}$ تصمیم بیز نسبت به π' است. پس توزیع δ^π نامساعدترین توزیع است. ■

شرط (۱) بیان می کند که میانگین تابع مخاطره $R(\theta, \delta^\pi)$ با ماکسیمم این تابع برابر باشد یکی از حالت هایی که این شرط برقرار می شود، این است که تابع مخاطره ثابت باشد، بنابراین:

نتیجه ۱-۱: اگر تصمیم بیز δ^π نسبت به توزیع پیشین π دارای مخاطره ثابت باشد، آنگاه δ^π تصمیم مینیماکس است.

اثبات:

$$R(\theta, \delta^\pi) = C \quad \forall \theta \in \Theta \Rightarrow \left\{ \begin{array}{l} \sup_{\theta \in \Theta} R(\theta, \delta^\pi) = C \\ \int R(\theta, \delta^\pi) d\Pi(\theta) = C \int d\Pi(\theta) = C \end{array} \right\}$$

بنابراین شرط (۱) برقرار است. ■

نتیجه ۲-۱: فرض کنید w_π مجموعه نقاطی از فضای پارامتری باشد که تابع مخاطره δ^π در روی آن ماکسیمم خود را به دست می آورد، یعنی

$$w_\pi = \left\{ \theta : R(\theta, \delta^\pi) = \sup_{\theta'} R(\theta', \delta^\pi) \right\}$$

در این صورت تصمیم بیز δ^π یک تصمیم مینیماکس است اگر $\Pi(w_\pi) = 1$.

(به عبارت دیگر یک شرط کافی برای آنکه برآوردگر δ^π مینیماکس باشد آن است که یک مجموعه w وجود داشته باشد که $\Pi(w) = 1$ و $R(\theta, \delta^\pi)$ ماکسیمم خود را در تمام نقاط w بدست آورد.)

اثبات: فرض کنید δ هر تصمیم دیگر به غیر از δ^π باشد در این صورت

$$\sup_{\Theta} R(\theta, \delta^\pi) = \int_{w^\pi} R(\theta, \delta^\pi) d\Pi(\theta) \stackrel{(**)}{=} \int R(\theta, \delta^\pi) d\Pi(\theta)$$

$$\leq \int R(\theta, \delta) d\Pi(\theta) \leq \sup_{\Theta} R(\theta, \delta)$$

که در آن تساوی (*) به این علت است که روی w^π مخاطره با سوپریمم برابر می شود و تساوی (***) نیز از این مطلب ناشی می گردد که $\Pi(w^\pi) = 0$. پس δ^π یک تصمیم مینیماکس است. ■

در حالت برآورد نقطه ای تصمیم مینیماکس را برآورد نقطه ای مینیماکس گویند.

مثال ۱-۱: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع $B(1, \theta)$ باشند، تحت تابع زیان درجه دوم برآوردگر مینیماکس θ را بدست آورید.

حل: در مثال ۲-۲ فصل قبل تحت چگالی پیشین $Be(\alpha, \beta)$ برآوردگر یکتای بیز عبارت بود از

$$\delta^\pi(x) = \frac{\alpha + \sum x_i}{\alpha + \beta + n}$$

که مخاطره آن برابر است با

$$\begin{aligned} R(\theta, \delta^\pi) &= E_\theta \left[\delta^\pi(x) - \theta \right]^2 = V_\theta(\delta^\pi(X)) + \left\{ E[\delta^\pi(X) - \theta] \right\}^2 \\ &= V_\theta \left(\frac{\alpha + \sum X_i}{\alpha + \beta + n} \right) + \left\{ E \left[\frac{\alpha + \sum X_i}{\alpha + \beta + n} - \theta \right] \right\}^2 \\ &= \frac{n\theta(1-\theta)}{(\alpha + \beta + n)^2} + \left(\frac{\alpha + n\theta - \alpha\theta - \beta\theta - n\theta}{\alpha + \beta + n} \right)^2 \\ &= \frac{1}{(\alpha + \beta + n)^2} \left[n\theta(1-\theta) + (\alpha - (\alpha + \beta)\theta)^2 \right] \\ &= \frac{1}{(\alpha + \beta + n)^2} \left[n\theta - n\theta^2 + \alpha^2 + (\alpha + \beta)^2 \theta^2 - 2\alpha(\alpha + \beta)\theta \right] \\ &= \frac{1}{(\alpha + \beta + n)^2} \left[\{(\alpha + \beta)^2 - n\}\theta^2 + \{n - 2\alpha(\alpha + \beta)\}\theta + \alpha^2 \right] \end{aligned}$$

در صورتی $R(\theta, \delta^*)$ به پارامتر θ بستگی ندارد که

$$\begin{cases} 2\alpha \left\{ (\alpha + \beta)^2 - n = 0 \right. \\ (\alpha + \beta) \left\{ n - 2\alpha(\alpha + \beta) = 0 \right. \end{cases} \Rightarrow (\alpha + \beta)n - 2\alpha n = 0 \Rightarrow \beta - \alpha = 0 \Rightarrow \alpha = \beta$$

$$\Rightarrow (2\alpha)^2 - n = 0 \Rightarrow \alpha = \frac{\sqrt{n}}{2} = \beta$$

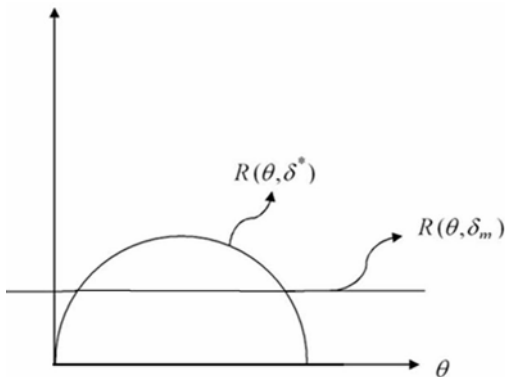
که با جایگذاری در معادله اول بدست می‌دهد:

بنابراین برآوردگر $\delta_m(X) = \frac{\frac{\sqrt{n}}{2} + \sum X_i}{n + \sqrt{n}}$ برآوردگر یکتای بیز با مخاطره ثابت است و در نتیجه برآوردگر مینیماکس یکتا است.

توجه کنید که برآوردگر UMVU پارامتر θ برابر $\delta^*(X) = \bar{X}$ است و به راحتی می‌توان نشان داد که

$$R(\theta, \delta_m) = \frac{1}{4(1 + \sqrt{n})^2}, \quad R(\theta, \delta^*) = \frac{\theta(1 - \theta)}{n}$$

این دو تابع مخاطره در زیر رسم شده‌اند.



دیده می‌شود که در بیشتر فاصله $\theta \in (0, 1)$ برآوردگر مینیماکس از برآوردگر UMVU بهتر عمل می‌کند ولی برآوردگر مینیماکس همواره از برآوردگرهای دیگر بهتر نیست. چون

$$\sup_{\theta \in (0, 1)} R(\theta, \bar{X}) = \sup_{\theta \in (0, 1)} \frac{\theta(1 - \theta)}{n} = \frac{1}{4n} > \frac{1}{4(1 + \sqrt{n})^2} = \sup_{\theta \in (0, 1)} R(\theta, \delta_m)$$

■ بنابراین تحت تابع زیان درجه دوم $\bar{X}$ نمی‌تواند برآوردگر مینیماکس باشد.

مثال ۱-۲: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $B(1, \theta)$ باشند، تحت تابع زیان

$$L(\theta, \delta) = \frac{(\delta - \theta)^2}{\theta(1 - \theta)}$$

و چگالی پیشین $Be(1, 1) = U(0, 1)$ برآوردگر بیز پارامتر θ را بدست آورده و

نشان دهید که یک برآوردگر مینیماکس است.

حل: ابتدا برآوردگر بیز را بدست می آوریم

$$\pi(\theta | \underline{x}) \propto \pi(\theta)L(\theta) = I_{(\cdot, \cdot)}(\theta)\theta^{\sum x_i} (1-\theta)^{n-\sum x_i}$$

$$\Rightarrow \theta | \underline{x} \sim Be(n\bar{x} + 1, n - n\bar{x} + 1)$$

$$\delta^\pi(\underline{x}) = \frac{E\left[\frac{\theta}{\theta(1-\theta)} | \underline{x}\right]}{E\left[\frac{1}{\theta(1-\theta)} | \underline{x}\right]} = \frac{\int_0^1 \theta^{n\bar{x}} (1-\theta)^{n-n\bar{x}-1} d\theta}{\int_0^1 \theta^{n\bar{x}-1} (1-\theta)^{n-n\bar{x}-1} d\theta}$$

قرار می دهیم $y = n\bar{x} = \sum_{i=1}^n x_i$ در این صورت اگر $y \neq 0$ و $y \neq n$ آنگاه

$$\delta^\pi(\underline{x}) = \frac{Be(n\bar{x} + 1, n - n\bar{x})}{Be(n\bar{x}, n - n\bar{x})} = \frac{\Gamma(n\bar{x} + 1)\Gamma(n - n\bar{x})}{\Gamma(n + 1)} \times \frac{\Gamma(n)}{\Gamma(n\bar{x})\Gamma(n - n\bar{x})} = \frac{n\bar{x}}{n} = \bar{x}$$

اگر $y = 0$ باشد آنگاه طبق قضیه ۱-۲ فصل قبل برآوردگر بیز آن مقداری از $\mathbf{a}$ است که انتگرال زیر را مینیمم کند

$$E\left[\frac{(\theta-a)^r}{\theta(1-\theta)} | \underline{x}\right] = \int_0^1 \frac{(\theta-a)^r}{\theta(1-\theta)} (1-\theta)^n d\theta = \int_0^1 (\theta-a)^r \theta^{-1} (1-\theta)^{n-1} d\theta$$

و این انتگرال بخاطر وجود θ^{-1} نامتناهی می شود بجز اینکه $a = 0$ باشد. بنابراین در این حالت برآورد

بیز عبارت است از $\delta^\pi(\underline{x}) = a = 0 = \frac{0}{n} = \frac{y}{n} = \bar{x}$ به همین ترتیب برای $y = n$ می توان نشان داد

که برآورد بیز عبارت است از $\delta^\pi(\underline{x}) = a = 1 = \frac{n}{n} = \frac{y}{n} = \bar{x}$ بنابراین در هر حالت $\delta^\pi(\underline{x}) = \bar{x}$

برآورد بیز است. همچنین

$$R(\theta, \delta^\pi) = E\left[\frac{(\bar{X} - \theta)^r}{\theta(1-\theta)}\right] = \frac{1}{\theta(1-\theta)} V_\theta(\bar{X}) = \frac{1}{\theta(1-\theta)} \frac{\theta(1-\theta)}{n} = \frac{1}{n}$$

بنابراین $\delta^\pi(\underline{X}) = \bar{X}$ یک برآوردگر بیز با مخاطره ثابت است و در نتیجه یک برآوردگر مینیمکس است. ■

لم ۱-۱: فرض کنید δ یک برآوردگر بیز (یا UMVU یا مینیماکس یا مجاز) برای $g(\theta)$ تحت تابع زیان درجه دوم باشد. در این صورت $a\delta + b$ یک برآوردگر بیز (یا UMVU یا مینیماکس یا مجاز) برای $ag(\theta) + b$ است

$$R(ag(\theta) + b, a\delta + b) = a^2 R(g(\theta), \delta) \quad \text{اثبات: چون}$$

بنابراین نتیجه به راحتی حاصل می‌گردد. ■

مثال ۱-۳: فرض کنید $X \sim B(n, p_1)$ ، $Y \sim B(n, p_2)$ و X و Y از یکدیگر مستقل باشند. با استفاده از نتیجه ۱-۲ یک برآوردگر مینیماکس برای $p_2 - p_1$ به دست آورید.

حل: با توجه به مثال ۱-۱ به دنبال برآوردگرهایی به فرم $C(Y - X)$ (این برآوردگر تحت یک تبدیل خاص پایا است) هستیم که مینیماکس باشند. تابع مخاطره‌ی این برآوردگرها عبارتند از

$$\begin{aligned} R(C(Y - X), p_1, p_2) &= E[C(Y - X) - (p_2 - p_1)]^2 \\ &= V(C(Y - X)) + [E(C(Y - X)) - (p_2 - p_1)]^2 \\ &= C^2 n(p_1(1 - p_1) + p_2(1 - p_2)) + (Cn - 1)^2(p_2 - p_1)^2 \end{aligned}$$

می‌خواهیم مجموعه‌ای از نقاط w را پیدا کنیم که تابع مخاطره فوق در این نقاط ماکسیمم باشد. بنابراین از تابع فوق نسبت به p_1 و p_2 مشتق جزئی می‌گیریم.

$$\begin{aligned} \frac{\partial R}{\partial p_1} = 0 &\Rightarrow [2(Cn - 1)^2 - 2C^2 n] p_1 - 2(Cn - 1)^2 p_2 = -C^2 n \\ \frac{\partial R}{\partial p_2} = 0 &\Rightarrow -2(Cn - 1)^2 p_1 + [2(Cn - 1)^2 - 2C^2 n] p_2 = -C^2 n \end{aligned} \quad (1)$$

اگر این دستگاه ۲ معادله و ۲ مجهول ترکیب خطی از یکدیگر نباشند آنگاه این دستگاه تنها دارای یک جواب (p_1^0, p_2^0) است و بنابراین π احتمال ۱ را به نقطه (p_1^0, p_2^0) بایستی نسبت دهد و برآوردگر بیز نسبت به این پیشین برآوردگر $\delta(X, Y) = p_2^0 - p_1^0$ است که دارای ماکسیمم مخاطره در نقطه (p_1^0, p_2^0) نیست. بنابراین در این حالت نمی‌توان از نتیجه ۱-۲ استفاده کرد.

حال اگر دو معادله‌ی دستگاه فوق یک ترکیب خطی از یکدیگر باشند آنگاه

$$2(Cn - 1)^2 - 2C^2 n = 2(Cn - 1)^2 \Rightarrow C^2 n = 2(Cn - 1)^2 \Rightarrow C = \frac{\sqrt{2n}}{n[\sqrt{2n} \pm 1]}$$

چون $-1 < p_2 - p_1 < 1$ می باشد پس حالت منفی در منجر که برآوردگر خارج از فضای پارامتری را می دهد قابل قبول نیست. با قرار دادن این مقدار C در معادله (۱) به معادله $p_1 + p_2 = 1$ می رسیم و

$$\delta(X, Y) = \frac{\sqrt{2n}}{n(\sqrt{2n} + 1)}(Y - X) \quad \text{بنابراین برآوردگر}$$

روی مجموعه $w = \{(p_1, p_2) \mid p_1 + p_2 = 1\}$ دارای مخاطره ماکسیمم است. حال اگر نشان دهیم که $\delta(x, y)$ برآوردگر بیز $p_2 - p_1$ نسبت به پیشینی است که احتمال ۱ به مجموعه w نسبت می دهد آنگاه طبق نتیجه ۱-۲، $\delta(x, y)$ برآوردگر مینیمکس است. اما $p_1 + p_2 = 1$ نتیجه می دهد که $p_2 - p_1 = 2p_2 - 1$ پس طبق لم ۱-۱ کافی است برآوردگر بیز p_2 را نسبت به این پیشین به دست آوریم. حال چون $p_2 = 1 - p_1$ بنابراین برای برآورد p_2 بر اساس n آزمایش با احتمال موفقیت p_2 (بر اساس Y) و n آزمایش با احتمال موفقیت $p_1 = 1 - p_2$ (بر اساس $n - X$) برآوردگر یکتای بیز طبق مثال ۱-۱ عبارت است از:

$$\frac{[Y + (n - X)] + \frac{\sqrt{2n}}{2}}{2n + \sqrt{2n}}$$

پس برآوردگر یکتای بیز $p_2 - p_1 = 2p_2 - 1$ طبق لم ۱-۱ عبارت است از

$$\begin{aligned} 2 \left\{ \frac{Y + (n - X) + \frac{\sqrt{2n}}{2}}{2n + \sqrt{2n}} \right\} - 1 &= \frac{[2(Y - X) + 2n + \sqrt{2n}] - [2n + \sqrt{2n}]}{\sqrt{2n}(\sqrt{2n} + 1)} \\ &= \frac{\sqrt{2n}}{n(\sqrt{2n} + 1)}(Y - X) = \delta(X, Y) \end{aligned}$$

پس $\delta(X, Y)$ برآوردگر مینیمکس یکتای $p_2 - p_1$ است. ■

نکات: (۱) توجه کنید که قضیه ۱-۱ می گوید که اگر برای یک برآوردگر بیز شرط $\int R(\theta, \delta^\pi) d\Pi(\theta) = \sup_{\theta \in \Theta} R(\theta, \delta^\pi)$ برقرار باشد آنگاه آن برآوردگر بیز، مینیمکس است. یکی از راه های برقراری شرط فوق ثابت بودن مخاطره δ^π است که در نتیجه ۱-۱ و مثال های قبل مورد

بررسی قرار گرفت. اما مثال‌هایی نیز وجود دارد که در آنها مخاطره برآوردگر بیز ثابت نیست اما شرط بالا برقرار می‌باشد و در نتیجه برآوردگر بیز حاصله مینیماکس خواهد بود. به تمرین زیر توجه کنید.

تمرین: فرض کنید $X \sim N(\theta, 1)$ که $\theta \in [-m, m]$ و m عددی ثابت و $0 < m < 1$ اگر

$$\delta^m(x) = \text{mtgh}(mx) = m \frac{e^{mx} - e^{-mx}}{e^{mx} + e^{-mx}}$$

الف- نشان دهید که $\delta^m(x)$ برآورد بیز پارامتر θ تحت توزیع پیشین زیر است.

$$\pi(\theta) = \begin{cases} \frac{1}{2} & \theta = m \\ \frac{1}{2} & \theta = -m \end{cases}$$

ب- مخاطره ی بیز δ^m را محاسبه کنید و نشان دهید که $r(\pi, \delta^m) = R(m, \delta^m)$

ج- نشان دهید که $\sup_{-m \leq \theta \leq m} R(\theta, \delta^m) = R(m, \delta^m)$ و بنابراین طبق شرط قضیه ۱-۱ δ^m

برآوردگر مینیماکس θ است.

$$\pi(\theta | x) = \frac{\pi(\theta)L(\theta)}{m(x)} = \frac{\frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-\theta)^2} \pi(\theta)}{m(x)}$$

حل تمرین:

(الف)

$$m(x) = \sum_{\theta=-m, m} \pi(\theta)L(\theta) = \frac{1}{\sqrt{2\pi}} \frac{1}{2} \left\{ e^{-\frac{1}{2}(x-m)^2} + e^{-\frac{1}{2}(x+m)^2} \right\}$$

$$\therefore \pi(\theta | x) = \begin{cases} \frac{e^{-\frac{1}{2}(x-m)^2}}{e^{-\frac{1}{2}(x-m)^2} + e^{-\frac{1}{2}(x+m)^2}} & \theta = m \\ \frac{e^{-\frac{1}{2}(x+m)^2}}{e^{-\frac{1}{2}(x-m)^2} + e^{-\frac{1}{2}(x+m)^2}} & \theta = -m \end{cases} = \begin{cases} \frac{e^{mx}}{e^{mx} + e^{-mx}} & \theta = m \\ \frac{e^{-mx}}{e^{mx} + e^{-mx}} & \theta = -m \end{cases}$$

$$\delta^m(x) = E[\theta | x] = m \frac{e^{mx}}{e^{mx} + e^{-mx}} - m \frac{e^{-mx}}{e^{mx} + e^{-mx}} = \text{mtgh}(mx)$$

$$\begin{aligned}
\text{ب) } r(\pi, \delta^m) &= \frac{1}{\gamma} \sum_{\theta=-m, m} R(\theta, \delta^m) = \frac{1}{\gamma} R(m, \delta^m) + \frac{1}{\gamma} R(-m, \delta^m) \\
&= \frac{1}{\gamma} E_{\theta=m} \left[(mtgh(mX) - m)^\gamma \right] + \frac{1}{\gamma} E_{\theta=-m} \left[(mtgh(mX) - m)^\gamma \right] \\
&= \frac{1}{\gamma} E_{\theta=m} \left[\frac{\gamma m^\gamma e^{-\gamma mX}}{(e^{mX} + e^{-mX})^\gamma} \right] + \frac{1}{\gamma} E_{\theta=-m} \left[\frac{\gamma m^\gamma e^{+\gamma mX}}{(e^{mX} + e^{-mX})^\gamma} \right] \\
g_n^{(x)} &= \frac{e^{-\gamma nx}}{(e^{nx} + e^{-nx})^\gamma}, X \sim N(n, 1), \gamma = E_{\theta=n} [g_n(x)] \\
\therefore r(\pi, \delta^m) &= \gamma m^\gamma \{E_{\theta=m} [g_m(X)] + E_{\theta=-m} [g_{-m}(X)]\} \\
&= \gamma m^\gamma E_{\theta=m} [g_m(X)] = \gamma m^\gamma E \left[\frac{e^{-\gamma mX}}{(e^{mX} + e^{-mX})^\gamma} \right] = R(m, \delta^m) \\
\text{ج) } R(\theta, \delta^m) &= E_\theta \left[(mtgh(mX) - \theta)^\gamma \right] \leq R(m, \delta^m)
\end{aligned}$$

که نامساوی آخر از این واقعیت ناشی می‌گردد که این نامعادله‌ی درجه دوم بر حسب θ در فاصله‌ی $[-m, m]$ ، مینیمم خود را در $\theta = E[mtgh(mX)]$ اختیار می‌کند.

(۲) از قضیه ۱-۱ موقعی می‌توان استفاده کرد که توزیع پیشین سره باشد. حال اگر توزیع پیشین ناسره باشد به دو صورت زیر می‌توان قضیه ۱-۱ را تعمیم داد.

الف- اگر توزیع پیشین ناسره باشد آنگاه ممکن است توزیع پسین سره باشد و بتوان برآوردگر بیز تعمیم یافته را به دست آورد و امیدوار بود که این برآوردگر مینیماکس باشد (یعنی مینیماکس بودن آن مورد بررسی قرار گیرد).

ب- دنباله‌ای از توزیع‌های پیشین ناسره را یافت که به سمت توزیع پیشین سره مورد نظر میل کنند و از روش بیز حدی که در زیرآورد می‌شود برآوردگر مینیماکس را به دست آورد.

تعریف ۱-۲: دنباله‌ای از توزیع‌های پیشین $\{\pi_n\}$ را نامساعدترین گویند اگر برای هر توزیع پیشین π داشته باشیم که

$$r_\pi \leq r = \lim_{n \rightarrow \infty} r_{\pi_n} \quad (۱)$$

که در آن $r_n = r_{\pi_n} = r(\pi_n, \delta^{\pi_n}) = \int R(\theta, \delta^{\pi_n}) d\Pi_n(\theta)$ مخاطره‌ی بیز نسبت به π_n است.

قضیه ۱-۲ (روش بیز حدی): فرض کنید $\pi_1, \pi_2, \pi_3, \dots$ دنباله‌ای از توزیع‌های پیشین باشند و δ^{π_n}

تصمیم بیز نسبت به توزیع پیشین π_n با مخاطره بیز $r_n = r(\pi_n, \delta^{\pi_n}) = \int R(\theta, \delta^{\pi_n}) d\Pi_n(\theta)$

باشد. اگر $\lim_{n \rightarrow \infty} r_n = r$ و δ_m تصمیمی باشد که $\sup_{\theta \in \Theta} R(\theta, \delta_m) = r$ آنگاه

الف) δ_m یک تصمیم مینیماکس است.

ب) دنباله‌ی $\{\pi_n\}$ نامساعدترین است.

اثبات: الف) برای هر تصمیم δ داریم که

$$r_n = r(\pi_n, \delta^{\pi_n}) = \int R(\theta, \delta^{\pi_n}) d\Pi_n(\theta) \stackrel{(*)}{\leq} \int R(\theta, \delta) d\Pi_n(\theta) \leq \sup_{\theta \in \Theta} R(\theta, \delta)$$

$$\Rightarrow \lim_{n \rightarrow \infty} r_n \leq \sup_{\theta \in \Theta} R(\theta, \delta)$$

که نامساوی (*) از این مطلب بدست می‌آید که δ^{π_n} تصمیم بیز نسبت به π_n است.

$$\sup_{\theta \in \Theta} R(\theta, \delta_m) = r = \lim_{n \rightarrow \infty} r_n \leq \sup_{\theta \in \Theta} R(\theta, \delta) \quad \forall \delta$$

بنابراین:

یعنی δ_m یک تصمیم مینیماکس است.

ب) فرض کنید π هر توزیع پیشین باشد.

$$r_n = \int_{\Theta} R(\theta, \delta^{\pi}) d\Pi(\theta) \leq \int_{\Theta} R(\theta, \delta) d\Pi(\theta) \leq \sup_{\theta \in \Theta} R(\theta, \delta) = r = \lim_{n \rightarrow \infty} r_n$$

پس دنباله‌ی $\{\pi_n\}$ نامساعدترین است. ■

نکات: ۱) در قسمت الف قضیه فوق کافی است داشته باشیم که

$$\sup_{\Theta} R(\theta, \delta) = r \leq \lim_{n \rightarrow \infty} r_n \quad (۱)$$

۲) توجه کنید که اگر $\delta^{\pi_n} \rightarrow \delta$ آنگاه دلیلی ندارد که $r_{\pi_n} \rightarrow r$ موقعی که $n \rightarrow \infty$. برای مشاهده

یک مثال مشخص به تمرین زیر توجه کنید.

تمرین: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\theta, 1)$ باشند و $\theta \sim N(b\sigma^2, \sigma^2)$

که $b \neq 0$. تحت تابع زیان درجه دوم، برآوردگر بیز δ^{π_σ} و مخاطره‌ی بیز آن یعنی $r_{\pi_\sigma} = r_\sigma$ را به

دست آورید و نشان دهید که اگر $\sigma \rightarrow \infty$ آنگاه $\delta_\sigma \rightarrow \delta$ که δ برآوردگری با مخاطره ثابت ۱ است،

اما $r_\sigma \not\rightarrow r$.

$$\delta_{\sigma} = \frac{n\bar{x} + b}{n + \frac{1}{\sigma^2}} \rightarrow \bar{x} + \frac{b}{n}, \quad R(\bar{X} + \frac{b}{n}, \theta) = \frac{1}{n} + \frac{b^2}{n} = r$$

$$\blacksquare \quad r_{\sigma} = \frac{1}{n + \frac{1}{\sigma^2}} \rightarrow \frac{1}{n} \neq \frac{1}{n} + \frac{b^2}{n}$$

(۳) قضیه ۱-۲ دارای مطلوبیت کمتری نسبت به قضیه ۱-۱ است زیرا

الف- در قضیه ۱-۲ اگر δ^{π_n} برآوردگر بیز یکتا باشد آنگاه دلیلی ندارد که δ_m برآوردگر مینیماکس یکتا باشد.

ب- برای استفاده از قضیه ۱-۲ ناگزیر هستیم که r و مخاطره ی بیز r_{π_n} را محاسبه کنیم که در بعضی موارد مشکل است. در بعضی موارد می توان برای محاسبه r_{π_n} از لم زیر استفاده کرد. با وجود مشکلات فوق در بسیاری موارد ناگزیر هستیم که از این قضیه یعنی روش بیز حدی استفاده کنیم.

لم ۱-۲: اگر δ^{π} برآوردگر بیز پارامتر $\gamma(\theta)$ نسبت به توزیع پیشین π باشد و اگر $r_{\pi} = E[(\delta^{\pi}(X) - \gamma(\theta))^2]$ مخاطره بیز آن باشد، آنگاه $r_{\pi} = \int \text{Var}[\gamma(\theta) | x] dP(x)$ و اگر واریانس پسین به x بستگی نداشته باشد آنگاه $r_{\pi} = \text{Var}[\gamma(\theta) | x]$.

$$\text{اثبات:} \quad r_{\pi} = E[(\delta_{\pi}(x) - \gamma(\theta))^2] = \int \left\{ E \left[(\gamma(\theta) - \delta^{\pi}(x))^2 | x \right] \right\} dP(x)$$

$$\blacksquare \quad = \int \text{Var}(\gamma(\theta) | x) dP(x)$$

مثال ۱-۴: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\theta, \sigma^2)$ باشد که در آن σ^2 معلوم است. تحت تابع زیان درجه دوم نشان دهید که $\bar{X}$ برآوردگر مینیماکس پارامتر θ است.

حل: دنباله توزیع های پیشین $N(0, m^2)$, $m = 1, 2, 3, \dots$ را در نظر بگیرید، در این صورت

$$\theta | x \sim N\left(\frac{m^2 \bar{x}}{m^2 + \frac{\sigma^2}{n}}, \frac{1}{\frac{n}{\sigma^2} + \frac{1}{m^2}}\right) \quad \delta^{\pi_m}(x) = \frac{m^2 \bar{x}}{m^2 + \frac{\sigma^2}{n}} \xrightarrow{m \rightarrow \infty} \bar{x}$$

مخاطره بیز δ^{π_m} عبارت است از

$$r_m = r(\pi_m, \delta^{\pi_m}) = \int V(\theta | x) dP(x) = \int \frac{1}{\frac{n}{\sigma^2} + \frac{1}{m^2}} dP(x) = \frac{1}{\frac{n}{\sigma^2} + \frac{1}{m^2}}$$

$$\lim_{m \rightarrow \infty} r_m = \lim_{m \rightarrow \infty} \frac{1}{\frac{n}{\sigma^2} + \frac{1}{m^2}} = \frac{\sigma^2}{n} = R(\theta, \bar{x}) = \sup_{\theta} R(\theta, \bar{x})$$

بنابراین طبق قضیه ۱-۲، $\bar{X}$ برآوردگر مینیماکس پارامتر θ است. ■

با توجه به مثال قبل برای هر σ^2 معلوم داریم که

$$\sup_{\theta} R(\theta, \sigma^2, \bar{X}) \leq \sup_{\theta} R(\theta, \sigma^2, \delta) \quad \forall \delta$$

$$\Rightarrow \sup_{\sigma^2} \sup_{\theta} R(\theta, \sigma^2, \bar{X}) \leq \sup_{\sigma^2} \sup_{\theta} R(\theta, \sigma^2, \delta) \quad \forall \delta$$

حال اگر σ^2 نامعلوم باشد از رابطه فوق نتیجه می شود که

$$\infty = \sup_{\sigma^2} \frac{\sigma^2}{n} = \sup_{\sigma^2} [\sup_{\theta} R(\theta, \sigma^2, \bar{X})] \leq \sup_{\sigma^2} \sup_{\theta} R(\theta, \sigma^2, \delta) \quad \forall \delta$$

یعنی مخاطره تمام برآوردگرها نامتناهی است بجز اینکه σ^2 متناهی باشد. بنابراین فرض می کنیم σ^2

نامعلوم و $\sigma^2 \leq M$ باشد در این صورت $\sup_{\sigma^2} \sup_{\theta} R(\theta, \sigma^2, \bar{X}) = \frac{M}{n}$. در این حالت مینیماکس

بودن برآوردگر $\bar{X}$ از قضیه زیر نتیجه می شود.

قضیه ۱-۳: فرض کنید X یک متغیر تصادفی با توزیع F از خانواده F_1 و $g(F)$ پارامتر مورد نظر

باشد و $F_0 \subset F_1$. اگر δ_0 برآوردگر مینیماکس $g(F)$ در F_0 باشد و

$$\sup_{F \in F_0} R(F, \delta_0) = \sup_{F \in F_1} R(F, \delta_0)$$

آنگاه δ_0 برآوردگر مینیماکس $g(F)$ در F_1 است.

اثبات: برای هر برآوردگر δ داریم که

$$\sup_{F \in F_1} R(F, \delta) \geq \sup_{F \in F_0} R(F, \delta) \geq \sup_{F \in F_0} R(F, \delta_0) \stackrel{(*)}{=} \sup_{F \in F_0} R(F, \delta_0) \stackrel{(**)}{=} \sup_{F \in F_1} R(F, \delta_0)$$

که در آن نامساوی (*) از مینیماکس بودن δ_0 در F_0 و تساوی (**) از فرض حاصل می‌گردد. پس δ_0 در F_1 مینیماکس است. ■

اگر قرار دهیم

$$F_0 = \{N(\theta, M) : \theta \in R, M < \infty\} \text{ و } F_1 = \{N(\theta, \sigma^2) : \theta \in R, \sigma^2 \leq M < \infty\}$$

آنگاه $F_0 \subset F_1$ و $\sup_{F_0} R(\theta, \sigma^2, \bar{X}) = \frac{M}{n} = \sup_{F_1} R(\theta, \sigma^2, \bar{X})$. بنابراین طبق قضیه ۱-۳ کافی است نشان دهیم که $\bar{X}$ در F_0 مینیماکس است و در نتیجه $\bar{X}$ در F_1 مینیماکس خواهد بود. اما تحت F_0

$$r_m = \frac{1}{\frac{n}{M} + \frac{1}{m^2}} \xrightarrow{m \rightarrow \infty} \frac{M}{n} = R(\theta, \bar{X}) = \sup_{\theta} R(\theta, \bar{X})$$

پس طبق قضیه ۱-۲ $\bar{X}$ در F_0 مینیماکس است. ■

مثال ۱-۵ (ناپارامتری): فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی با توزیع یکسان F که دارای میانگین متناهی θ است، باشد. برآوردگر مینیماکس پارامتر θ تحت تابع زیان درجه دوم را در هر یک از حالات زیر به دست آورید.

الف- واریانس کراندار $V_F(X_i) \leq M < \infty$

ب- حدود کراندار $-\infty < a < X_i < b < +\infty$

(توجه کنید اگر ماکسیمم مخاطره هر برآوردگر بی‌نهایت شود آنگاه مسئله مینیماکس منتفی است پس شرایط فوق را برای گریز از این موضوع در مسئله اعمال کرده ایم)

حل : الف- اگر قرار دهیم $F_0 = \{N(\theta, \sigma^2) | \sigma^2 \leq M < \infty\}$

$$F_1 = \{F | E_F(X_i) = \theta, V_F(X_i) \leq M < \infty\}$$

در این صورت $F_0 \subset F_1$ و طبق بحث قبل از این مثال $\bar{X}$ در F_0 برآوردگری مینیماکس است و

$$\begin{cases} \sup_{F_0} R(\theta, \sigma^r, \bar{X}) = \frac{M}{n} \\ \sup_{F_1} R(\theta, \sigma^r, \bar{X}) = \sup_{F_1} V(\bar{X}) = \sup_{F_1} \frac{\sigma^r}{n} = \frac{M}{n} \end{cases}$$

پس شرایط قضیه ۱-۳ برقرار است، یعنی $\bar{X}$ در F_1 مینیماکس است.

ب- بدون از دست دادن کلیت مسئله فرض می کنیم $b = 1$ و $a = 0$ ، قرار می دهیم

$$F_1 = \{F \mid F(1) - F(0) = 1\}$$

$$F_0 = \{B(1, \theta) \mid 0 < \theta < 1\}$$

پس $F_0 \subset F_1$ و در مثال ۱-۱ دیدیم که

$$\delta_0(X) = \frac{\sqrt{\frac{n}{2}} + \sum_{i=1}^n X_i}{n + \sqrt{n}} = \frac{\sqrt{n}}{1 + \sqrt{n}} \bar{X} + \frac{1}{2(1 + \sqrt{n})}$$

برآوردگر مینیماکس در F_0 می باشد. حال طبق قضیه ۱-۳ کافی است نشان دهیم که ماکسیمم مخاطره δ_0 در F_0 و F_1 یکی است.

$$F_0 \text{ تحت } R(\theta, \delta_0) = \frac{1}{4(1 + \sqrt{n})^2} \Rightarrow \sup_{F_0} R(\theta, \delta_0) = \frac{1}{4(1 + \sqrt{n})^2}$$

$$\begin{aligned} F_1 \text{ تحت } R(\theta, \delta_0) &= E \left[\frac{\sqrt{n}}{1 + \sqrt{n}} \bar{X} + \frac{1}{2(1 + \sqrt{n})} - \theta \right]^2 \\ &= \frac{n}{(1 + \sqrt{n})^2} V(\bar{X}) + \left(\frac{\sqrt{n}}{1 + \sqrt{n}} \theta + \frac{1}{2(1 + \sqrt{n})} - \theta \right)^2 \\ &= \frac{1}{(1 + \sqrt{n})^2} [V_F(X_1) + \left(\frac{1}{2} - \theta \right)^2] \end{aligned}$$

$$V_F(X_1) = E(X - \theta)^2 = E(X^2) - \theta^2 \leq E(X) - \theta^2 = \theta - \theta^2$$

زیرا $0 \leq X \leq 1 \Rightarrow X^2 \leq X$

$$\therefore R(\delta_0, \theta) \leq \frac{1}{(1 + \sqrt{n})^2} [(\theta - \theta^2) + \left(\frac{1}{2} - \theta \right)^2] = \frac{1}{4(1 + \sqrt{n})^2}$$

$$\Rightarrow \sup_{F_1} R(\theta, \delta_0) = \frac{1}{4(1 + \sqrt{n})^2} = \sup_{F_0} R(\theta, \delta_0)$$

پس طبق قضیه ۱-۳، $\delta_0(X)$ در F_1 برای θ مینیماکس است. ■

بخش ۲: تصمیم‌های مجاز^۱ (پذیرفتنی)

در یک مسئله تصمیم کلاس توابع تصمیم D معمولاً کلاس بزرگی است و انتخاب اینکه کدام تصمیم $\delta \in D$ را مورد استفاده قرار دهیم کار مشکلی می باشد. بنابراین بایستی یک زیر کلاس C از D پیدا کنیم که در این کلاس کلیه تصمیم‌های خوب قرار داشته باشند. حال چون این کلاس C یک کلاس کوچکتر است، پیدا کردن یک تصمیم خوب در آن راحت تر می باشد. در این قسمت کلاسی از توابع تصمیم خوب را با استفاده از محک مجاز بودن^۲ تعریف می کنیم و نحوه‌ی بدست آوردن این تصمیم‌ها را ارائه می دهیم.

مقایسه تصمیم‌ها:

تعریف ۱-۲: فرض کنید δ_1 و δ_2 دو تصمیم با توابع مخاطره $R(\theta, \delta_1)$ و $R(\theta, \delta_2)$ باشند.

الف- تصمیم δ_1 را به خوبی δ_2 گویند اگر: $R(\theta, \delta_1) \leq R(\theta, \delta_2) \quad \forall \theta \in \Theta$

ب- تصمیم δ_1 را بهتر از δ_2 گویند اگر: $R(\theta, \delta_1) < R(\theta, \delta_2) \quad \forall \theta \in \Theta$

و برای بعضی $\theta \in \Theta$ داشته باشیم که: $R(\theta, \delta_1) < R(\theta, \delta_2)$

ج- تصمیم δ_1 را هم ارز^۳ δ_2 گویند اگر: $R(\theta, \delta_1) = R(\theta, \delta_2) \quad \forall \theta \in \Theta$

د- تصمیم‌های δ_1 و δ_2 را غیرقابل مقایسه گویند هرگاه هیچکدام از سه حالت بالا برقرار نباشد.

مثال ۱-۲: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\mu, \sigma^2)$ باشد. می دانیم که S^2 برآوردگر UMVU برای پارامتر σ^2 است. نشان دهید که در کلاس برآوردگرهای CS^2 برآوردگری بهتر از S^2 وجود دارد.

حل: در فصل مربوط به برآوردگرهای نااریب نشان دادیم که برآوردگر

$$\tilde{S}^2 = \frac{n-1}{n+1} S^2 = \frac{1}{n+1} \sum_{i=1}^n (X_i - \bar{X})^2$$

که یک برآوردگر MRE است دارای مخاطره کمتری

^۱ Admissible Decisions

^۲ Admissibility

^۳ As Good As

^۴ Better Than

^۵ Equivalent

نسبت به S^2 است یعنی $\forall \sigma^2, R(\sigma^2, \tilde{S}^2) < R(\sigma^2, S^2)$, بنابراین برآوردگر $\tilde{S}^2$ در کلاس cS^2 بهترین برآوردگر است. ■

مثال ۲-۲: فرض کنید X_1, X_2 نمونه ای تصادفی از $N(\mu, \sigma^2)$ باشند برآوردگرهای زیر را برای پارامتر μ در نظر بگیرید.

$$\delta_1 = \frac{1}{2}X_1 + \frac{1}{2}X_2, \quad \delta_2 = \frac{1}{4}X_1 + \frac{3}{4}X_2, \quad \delta_3 = \frac{3}{4}X_1 + \frac{1}{4}X_2, \quad \delta_4 = \frac{1}{2}X_1 + \frac{1}{2}X_2 + 1$$

اگر تابع زیان درجه دوم باشد این ۴ برآوردگر را با یکدیگر مقایسه کنید.

حل:

$$R(\mu, \delta_1) = V(\delta_1) = \frac{1}{2}\sigma^2 \quad R(\mu, \delta_2) = V(\delta_2) = \left(\frac{1}{16} + \frac{9}{16}\right)\sigma^2 = \frac{5}{8}\sigma^2$$

$$R(\mu, \delta_3) = V(\delta_3) = \left(\frac{9}{16} + \frac{1}{16}\right)\sigma^2 = \frac{5}{8}\sigma^2 \quad R(\mu, \delta_4) = V(\delta_4) = \frac{1}{2}\sigma^2 + 1$$

بنابراین δ_1 بهتر از δ_2 و δ_3 و δ_4 است و δ_2 هم ارز δ_3 است و δ_4 با δ_2 و δ_3 قابل مقایسه نیست زیرا اگر $\sigma^2 \geq 8$ آنگاه $R(\mu, \delta_4) \leq R(\mu, \delta_2)$ و اگر $\sigma^2 < 8$ آنگاه $R(\mu, \delta_4) > R(\mu, \delta_2)$. ■

تعریف ۱-۲ مقایسه ای را بین تمام توابع تصمیم بوجود نمی آورد و تنها یک ترتیب را روی کلاس توابع تصمیم D تعریف می کند.

تعریف ۲-۲: یک تصمیم δ را مجاز^۱ گویند هرگاه تصمیم دیگری مانند $\delta' \in D$ وجود نداشته باشد که از δ بهتر باشد. یک تصمیم δ را غیر مجاز^۲ گویند اگر تصمیم دیگری مانند $\delta' \in D$ وجود داشته باشد که از δ بهتر باشد. ■

با توجه به مثال ۱-۲ در خانواده $N(\mu, \sigma^2)$ برآوردگر S^2 برای σ^2 غیر مجاز است زیرا برآوردگر $\tilde{S}^2 = \frac{n-1}{n+1}S^2$ از آن بهتر است. اما خود برآوردگر S^2 نیز مجاز نیست و توسط Brewster and Zidek (۱۹۷۴) برآوردگر بهتر از آن پیدا شده است.

^۱ admissible

^۲ inadmissible

با توجه به تعریف ۲-۳ یک تصمیم مجاز تصمیم خوبی است زیرا هیچ تصمیم دیگری بر آن ارجحیت ندارد و یک تصمیم غیرمجاز تصمیم بدی است زیرا تصمیم دیگری وجود دارد که بر آن ارجحیت دارد. با این وجود توجه کنید که مجاز بودن واقعاً یک خاصیت مثبت نیست و همچنین خاصیت منفی نیز نیست، یعنی یک تصمیم مجاز لزوماً به طور یکنواخت خوب نیست و همچنین به طور یکنواخت نیز بد نمی باشد. اگر بدانیم که یک تصمیم غیر مجاز است، در این صورت بایستی دنیال تصمیمی بگردیم که از آن بهتر باشد. اما اگر بدانیم که یک تصمیم مجاز است دلیل بر آن نیست که حتماً از این تصمیم استفاده کنیم زیرا در خیلی از موارد تصمیم‌های مجاز زیادی وجود دارد (مثال ۲-۴ را ملاحظه کنید) که بعضی معقول و بعضی غیرمعقول هستند و بعضی به سختی قابل محاسبه هستند. به مثال زیر توجه کنید.

مثال ۲-۳: فرض کنید $X \sim B(10, \theta)$ نشان دهید که $\delta(X) = \frac{1}{3}$ برآوردگری مجاز برای θ تحت

تابع زیان درجه دوم است و آن را با برآوردگر UMVU پارامتر θ یعنی $\delta^*(X) = \frac{X}{10}$ مقایسه کنید.

حل: تابع مخاطره برآورد δ در $p = \frac{1}{3}$ عبارت است از

$$\begin{aligned} R\left(\frac{1}{3}, \delta\right) &= E\left[\left(\delta(x) - \frac{1}{3}\right)^2\right] \\ &= \sum_{x=0}^{10} \left(\delta(x) - \frac{1}{3}\right)^2 P_{P=\frac{1}{3}}(X=x) = \sum_{x=0}^{10} \left(\frac{1}{3} - \frac{1}{3}\right)^2 P_{P=\frac{1}{3}}(X=x) = 0. \end{aligned}$$

حال فرض کنید $\delta'(x)$ برآوردگر دیگر به خوبی $\delta(x)$ باشد در این صورت

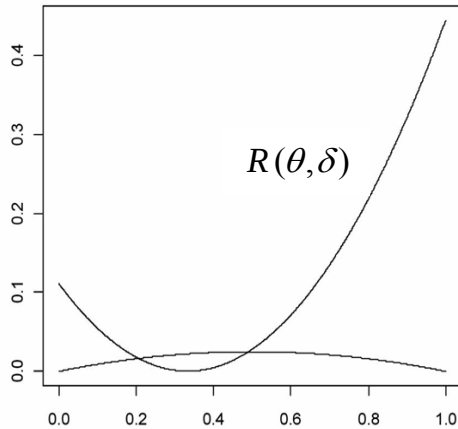
$$R\left(\frac{1}{3}, \delta'\right) \leq R\left(\frac{1}{3}, \delta\right) = 0 \xrightarrow{R\left(\frac{1}{3}, \delta'\right) \geq 0} R\left(\frac{1}{3}, \delta'\right) = 0 \Rightarrow E\left[\left(\delta'(x) - \frac{1}{3}\right)^2\right] = 0.$$

$$\Rightarrow P\left(\left(\delta'(x) - \frac{1}{3}\right)^2 = 0\right) = 1 \Rightarrow P\left(\delta'(x) = \frac{1}{3}\right) = 1$$

بنابراین اگر δ' بخواهد به خوبی δ باشد بایستی با احتمال یک با δ برابر باشد. حال اگر δ' بهتر از δ باشد آنگاه به خوبی δ است و در نتیجه باز هم بایستی δ' و δ با احتمال یک با هم برابر باشند و در نتیجه هیچ برآوردگری بهتر از δ وجود ندارد و در نتیجه δ یک برآوردگر مجاز است.

تنها مزیت برآوردگر مجاز $\delta(X) = \frac{1}{3}$ این است که در $\theta = \frac{1}{3}$ مخاطره آن از هر برآوردگری کمتر است و در مقایسه با برآوردگر UMVU $\delta^*(X) = \frac{X}{1.0}$ یک برآوردگر غیرمعقول است. مقایسه

مخاطره این دو برآوردگر نشان می دهد که در بیشتر نقاط



فاصله $\theta \in (0, 1)$ برآوردگر UMVU بهتر از $\delta(X) = \frac{1}{3}$

عمل می کند و برآوردگر δ تنها در نقطه ی $\theta = \frac{1}{3}$

خوب عمل می کند.

$$R(\theta, \delta) = \left(\frac{1}{3} - \theta\right)^2$$

$$R(\theta, \delta^*) = \frac{\theta(1-\theta)}{1.0}$$

■

با استفاده از تعریف ۱-۲ در زیر کلاسی از توابع تصمیم را معرفی می کنیم که شامل تمام توابع تصمیم خوب است.

تعریف ۳-۲: فرض کنید D کلاسی کلیه توابع تصمیم باشد و C کلاسی از توابع تصمیم و زیر کلاسی از D باشد $C \subseteq D$.

الف- کلاس C را یک کلاس کامل^۱ گویند اگر برای هر تصمیم $\delta' \notin C$ یک تصمیم $\delta \in C$ وجود داشته که δ بهتر از δ' باشد.

ب- کلاس C را یک کلاس اساسا کامل^۲ گویند اگر برای هر تصمیم $\delta' \notin C$ یک تصمیم $\delta \in C$ وجود داشته باشد که δ به خوبی δ' باشد. ■

قضیه های زیر رابطه تصمیم های مجاز و کلاس های کامل و اساسا کامل را نشان می دهد.

قضیه ۱-۲: اگر C یک کلاس کامل باشد آنگاه کلاس کلیه تصمیم های مجاز A درون C قرار دارد $(A \subseteq C)$.

^۱ complete

^۲ essentially complete

اثبات: فرض کنید δ' یک تصمیم مجاز باشد ($\delta' \in A$) ولی $\delta' \notin C$. در این صورت طبق تعریف بایستی یک $\delta \in C$ وجود داشته باشد که δ از δ' بهتر باشد و چون δ' مجاز است این امکان ندارد پس بایستی $\delta' \in C$ باشد و در نتیجه $\delta' \in A \Rightarrow \delta' \in C$ پس $A \subset C$.

قضیه ۲-۲: فرض کنید C یک کلاس اساسا کامل باشد. اگر $\delta' \notin C$ یک تصمیم مجاز باشد آنگاه یک تصمیم $\delta \in C$ وجود دارد بطوری که هم ارز δ' است.

اثبات: چون C اساسا کامل است و $\delta' \notin C$ پس یک تصمیم $\delta \in C$ وجود دارد بطوری که $R(\theta, \delta) \leq R(\theta, \delta'), \quad \forall \theta \in \Theta$ و چون δ' یک تصمیم مجاز است پس نامساوی فوق بطور اکید برای هیچ $\theta \in \Theta$ برقرار نیست پس $R(\theta, \delta) = R(\theta, \delta'), \quad \forall \theta \in \Theta$ یعنی δ و δ' هم ارز هستند. ■

با توجه به تعریف و قضایای بالا بایستی توجه خود را به کلاس توابع تصمیم کامل معطوف کنیم و در این کلاس برآوردهای مجاز را پیدا کنیم.

روش‌های بدست آوردن برآوردهای مجاز و تصمیم‌های مجاز

در زیر قضایایی را برای بدست آوردن تصمیم‌های مجاز و غیرمجاز می‌آوریم که عمده آنها در مورد برآوردهای مجاز است.

قضیه ۲-۳: اگر δ^π یک تصمیم بیز یکتا نسبت به چگالی پیشین $\pi(\theta)$ و تابع زیان $L(\theta, \delta)$ باشد آنگاه δ^π یک تصمیم مجاز است.

اثبات: فرض کنید δ^π یک تصمیم مجاز نباشد در این صورت یک تصمیم δ وجود دارد بطوری که:

$$R(\theta, \delta) \leq R(\theta, \delta^\pi) \quad \forall \theta \in \Theta, \quad R(\theta', \delta) < R(\theta', \delta^\pi) \quad \text{for some } \theta' \in \Theta \quad (۱)$$

بنابراین

$$r(\pi, \delta) = \int R(\theta, \delta) d\Pi(\theta) \stackrel{(۱)}{\leq} \int R(\theta, \delta^\pi) d\Pi(\theta) = r(\pi, \delta^\pi) \quad (۲)$$

از طرفی چون δ^π برآورده بیز است پس

$$r(\pi, \delta^\pi) \leq r(\pi, \delta) \quad (۳)$$

$$\text{بنابراین} \quad (۲), (۳) \Rightarrow r(\pi, \delta^\pi) = r(\pi, \delta)$$

یعنی δ نیز یک تصمیم بیز است که این با یکتا بودن تصمیم بیز متناقض است. پس δ^π یک تصمیم مجاز است. ■

قضیه ۲-۴: فرض کنید $\gamma(\theta)$ پارامتر مورد نظر ما باشد و $\gamma(\theta) \in [a, b]$ و تابع زیان $L(\theta, \delta)$ برای $\delta \neq \gamma(\theta)$ مثبت و برای $\delta = \gamma(\theta)$ صفر باشد و تابع زیان $L(\theta, \delta)$ با دور شدن δ از $\gamma(\theta)$ یک تابع صعودی باشد. در این صورت هر برآوردگر δ که مقادیری را در خارج از فاصله $[a, b]$ با احتمال مثبت بپذیرد غیرمجاز است.

اثبات: برآوردگر δ' را به صورت زیر تعریف می کنیم.

$$\delta' = \begin{cases} a & \delta < a \\ \delta & a \leq \delta \leq b \\ b & \delta > b \end{cases}$$

در این صورت بنابر فرض قضیه اگر $\delta < a$ آنگاه $L(\theta, a) < L(\theta, \delta)$ و اگر $\delta > b$ آنگاه $L(\theta, b) < L(\theta, \delta)$ بنابراین

$$\begin{aligned} R(\theta, \delta') &= E[L(\theta, a)I_{\delta < a}] + E[L(\theta, \delta)I_{a \leq \delta \leq b}] + E[L(\theta, b)I_{\delta > b}] \\ &\leq E[L(\theta, \delta)I_{\delta < a}] + E[L(\theta, \delta)I_{a \leq \delta \leq b}] + E[L(\theta, \delta)I_{\delta > b}] \\ &= E[L(\theta, \delta)] = R(\theta, \delta) \end{aligned}$$

پس δ برآوردگری غیرمجاز است. ■

قضیه ۲-۵: فرض کنید تابع زیان درجه دوم باشد و X یک متغیر تصادفی با میانگین θ و واریانس σ^2 باشد در این صورت $aX + b$ یک برآوردگر غیرمجاز برای θ است وقتی که

$$(i) \ a > 1 \quad \text{یا} \quad (ii) \ a < 0 \quad \text{یا} \quad (iii) \ a = 1, \ b \neq 0$$

$$\text{اثبات:} \quad f(a, b) = R(\theta, aX + b) = E[(aX + b - \theta)^2] = \text{Var}(aX + b) + [E(aX + b) - \theta]^2$$

$$= a^2 \text{Var}(X) + (a\theta + b - \theta)^2 = a^2 \sigma^2 + [(a-1)\theta + b]^2$$

$$(i) \quad a > 1 \Rightarrow f(a, b) \geq a^2 \sigma^2 > \sigma^2 = f(1, 0) = R(\theta, X)$$

پس $aX + b$ به وسیله X مغلوب^۱ می شود (یعنی $(R(\theta, X) < R(\theta, aX + b))$)

$$(ii) \quad a < 0 \Rightarrow (a-1)^r > 1 \Rightarrow f(a, b) \geq [(a-1)\theta + b]^r = (a-1)^r [\theta + \frac{b}{a-1}]^r \\ > [\theta + \frac{b}{a-1}]^r = f(\cdot, \frac{-b}{a-1}) = R(\theta, \frac{-b}{a-1})$$

بنابراین $aX + b$ به وسیله برآوردگر ثابت $\delta = \frac{-b}{a-1}$ مغلوب می شود.

$$(iii) \quad a = 1, b \neq 0 \Rightarrow f(a, b) = f(1, b) = \sigma^r + b^r > \sigma^r = f(1, 0) = R(\theta, X)$$

پس $aX + b$ به وسیله X مغلوب می شود. ■

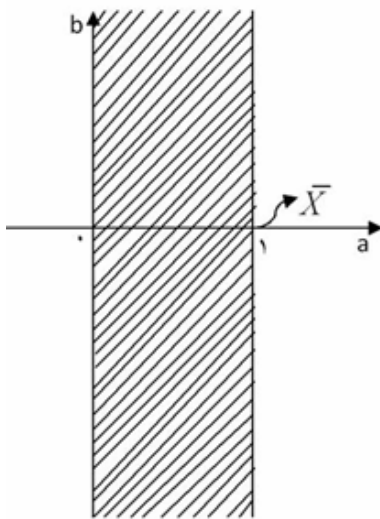
مثال ۲-۴: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\theta, \sigma^2)$ باشند که در آن σ^2 معلوم است. برای چه مقادیر از a و b کلاس برآوردگرهای خطی $a\bar{X} + b$ مجاز می باشد.

حل: در فصل قبل دیدیم که برآوردگر بیز یکتای θ تحت توزیع پیشین $N(\mu, \tau^2)$ عبارت است از:

$$\delta^\pi(\bar{X}) = \frac{n\tau^2}{\sigma^2 + n\tau^2} \bar{X} + \frac{\sigma^2}{\sigma^2 + n\tau^2} \mu$$

که مخاطره بیز این برآوردگر متناهی است زیرا

$$R(\theta, \delta^\pi) = \left(\frac{n\tau^2}{\sigma^2 + n\tau^2} \right)^r \frac{\sigma^2}{n} + \left[\frac{n\tau^2\theta + \mu\sigma^2}{\sigma^2 + n\tau^2} - \theta \right]^r \Rightarrow r(\pi, \delta^\pi) = \int_{\Theta} R(\theta, \delta^\pi) d\Pi(\theta) < \infty$$



حال چون $0 < \frac{n\tau^2}{\sigma^2 + n\tau^2} < 1$ و $-\infty < \frac{\mu\sigma^2}{\sigma^2 + n\tau^2} < +\infty$ پس

برآوردگر $a\bar{X} + b$ برای $b \in R, 0 < a < 1$ یک برآوردگر مجاز است.

چون $E(\bar{X}) = \theta$ پس طبق قضیه ۵-۲ برآوردگر $a\bar{X} + b$ برای

$a > 1$ (حالت (i)) و $a < 0$ (حالت (ii)) و $a = 1, b \neq 0$ (حالت

(iii)) مجاز نیست.

^۱ Dominate

برای حالت $a = 0$ برآوردگر $\delta(X) = b$ مجاز است زیرا این برآوردگر تنها برآوردگری است که در $\theta = b$ دارای مخاطره صفر است، $R(b, \delta) = 0$ (مثال ۲-۳ را ملاحظه کنید).
تنها نقطه‌ای از صفحه مختصات (a, b) که باقی مانده است نقطه $(0, 1)$ مربوط به برآوردگر $ML, UMVU, MRE$ و مینیماکس $\bar{X}$ است. که از دو روش، یکی روش بلیت^۱ (۱۹۵۱) یا روش بیزی حدی^۲ و دیگری روش نامساوی اطلاع مجاز بودن آن را ثابت می‌کنیم.

الف- روش بیز حدی

قضیه ۲-۶: فرض کنید $\theta \in \Theta \subset R$ و δ یک تابع تصمیم باشد. فرض کنید که $R(\theta, \delta)$ برای هر δ ثابت یک تابع پیوسته از θ باشد. فرض کنید دنباله‌ای از توزیع‌های پیشین π_m با تصمیم‌های بیز متناظر $\delta^m = \delta^{\pi_m}$ و مخاطره‌های بیز متناظر $r_m = r(\pi_m, \delta^m)$ موجود باشند. همچنین فرض کنید r_m^* مخاطره بیز تصمیم δ نسبت به توزیع پیشین π_m باشد. اگر برای هر $\eta > 0$ و $\theta \in \Theta$ داشته باشیم که:

$$\frac{\Pi_m(\theta + \eta) - \Pi_m(\theta - \eta)}{r_m^* - r_m} \rightarrow \infty \quad \text{as} \quad m \rightarrow \infty$$

(که Π_m تابع توزیع پیشین با چگالی پیشین π_m است) آنگاه δ یک تصمیم مجاز است.

اثبات: فرض کنید δ یک تصمیم مجاز نباشد در این صورت یک تصمیم δ^* وجود دارد به طوری که

$$R(\theta, \delta^*) \leq R(\theta, \delta) \quad \forall \theta \in \Theta$$

$$R(\theta, \delta^*) < R(\theta, \delta) \quad \text{for some } \theta \in \Theta$$

$$\varepsilon = R(\theta, \delta) - R(\theta, \delta^*) > 0 \quad \text{قرار می‌دهیم}$$

چون تابع مخاطره $R(\theta, \delta)$ برای δ ثابت یک تابع پیوسته از θ است پس

$$\forall \varepsilon > 0 \quad \exists \eta > 0 \quad \exists |\theta - \theta_0| \leq \eta \Rightarrow \begin{cases} |R(\theta, \delta^*) - R(\theta_0, \delta^*)| < \varepsilon/4 \\ |R(\theta, \delta) - R(\theta_0, \delta)| < \varepsilon/4 \end{cases}$$

^۱ Blyth

^۲ Limiting Bayes Method

$$\begin{aligned}\Rightarrow R(\theta, \delta) - R(\theta, \delta^*) &> \left(R(\theta, \delta) - \frac{\varepsilon}{4} \right) - \left(R(\theta, \delta^*) + \frac{\varepsilon}{4} \right) \\ &= R(\theta, \delta) - R(\theta, \delta^*) - \frac{\varepsilon}{2} = \varepsilon - \frac{\varepsilon}{2} = \frac{\varepsilon}{2}\end{aligned}$$

$$R(\theta, \delta) - R(\theta, \delta^*) > \frac{\varepsilon}{2} \quad \forall \theta \in [\theta - \eta, \theta + \eta] \quad \text{در نتیجه}$$

بنابراین اگر r_m^{**} مخاطره بیز تصمیم δ^* باشد یعنی $r_m^{**} = \int R(\theta, \delta^*) d\Pi_m(\theta)$ در این صورت

$$\begin{aligned}r_m^* - r_m^{**} &= \int_{-\infty}^{+\infty} [R(\theta, \delta) - R(\theta, \delta^*)] d\Pi_m(\theta) \\ &\geq \int_{\theta-\eta}^{\theta+\eta} [R(\theta, \delta) - R(\theta, \delta^*)] d\Pi_m(\theta) \\ &\geq \frac{\varepsilon}{2} \int_{\theta-\eta}^{\theta+\eta} d\Pi_m(\theta) = \frac{\varepsilon}{2} [\Pi_m(\theta + \eta) - \Pi_m(\theta - \eta)] \\ \Rightarrow \frac{r_m^* - r_m^{**}}{r_m^* - r_m} &\geq \frac{\frac{\varepsilon}{2} [\Pi_m(\theta + \eta) - \Pi_m(\theta - \eta)]}{r_m^* - r_m} \rightarrow \infty \quad \text{as } m \rightarrow \infty\end{aligned}$$

بنابراین m وجود دارد که $\frac{r_m^* - r_m^{**}}{r_m^* - r_m} > 1$ و یا معادلاً $r_m^* > r_m^{**}$ که این با تصمیم بیز بودن δ^m نسبت به توزیع پیشین π_m متناقض است. بنابراین δ یک تصمیم مجاز است. ■

مثال ۲-۵: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\theta, \sigma^2)$ باشند نشان دهید که $\bar{X}$ یک برآوردگر مجاز برای θ است.

حل: اگر توزیع پیشین $N(0, m^2)$ را برای θ در نظر بگیریم آنگاه

$$\begin{aligned}\theta | \underline{X} &\sim N \left(\frac{nm^2 \bar{X}}{\sigma^2 + nm^2}, \frac{1}{\frac{n}{\sigma^2} + \frac{1}{m^2}} \right) \quad \delta^m = \delta^{\pi_m} = \frac{nm^2}{\sigma^2 + nm^2} \bar{X} \\ r_m &= r(\pi_m, \delta^{\pi_m}) = E_{\underline{X}} [V(\theta | \underline{X})] = \frac{1}{\frac{n}{\sigma^2} + \frac{1}{m^2}}\end{aligned}$$

همچنین مخاطره بیز برآوردگر $\delta.(X) = \bar{X}$ عبارت است از

$$r_m^* = r(\pi_m, \delta.) = E[E(\bar{X} - \theta) | \theta] = E\left(\frac{\sigma^2}{n}\right) = \frac{\sigma^2}{n}$$

بنابراین

$$\begin{aligned} \frac{\Pi_m(\theta + \eta) - \Pi_m(\theta - \eta)}{r_m^* - r_m} &= \frac{\int_{\theta-\eta}^{\theta+\eta} \frac{1}{\sqrt{2\pi m^2}} e^{-\frac{1}{2m^2}t^2} dt}{\frac{\sigma^2}{n} - \left[\frac{n}{\sigma^2} + \frac{1}{m^2}\right]^{-1}} \\ &= \frac{\frac{1}{m\sqrt{2\pi}} \int_{\theta-\eta}^{\theta+\eta} e^{-\frac{1}{2m^2}t^2} dt}{\frac{\sigma^2}{n(nm^2 + \sigma^2)}} = \frac{n(nm^2 + \sigma^2)}{m\sqrt{2\pi}\sigma^2} \int_{\theta-\eta}^{\theta+\eta} e^{-\frac{1}{2m^2}t^2} dt \\ \lim_{m \rightarrow +\infty} e^{-\frac{1}{2m^2}t^2} &= 1 \Rightarrow \lim_{m \rightarrow +\infty} \int_{\theta-\eta}^{\theta+\eta} e^{-\frac{1}{2m^2}t^2} dt = \int_{\theta-\eta}^{\theta+\eta} \lim_{m \rightarrow +\infty} e^{-\frac{1}{2m^2}t^2} dt = (\theta + \eta) - (\theta - \eta) = 2\eta \end{aligned}$$

$$\therefore \frac{\Pi_m(\theta + \eta) - \Pi_m(\theta - \eta)}{r_m^* - r_m} \rightarrow +\infty \quad \text{as } m \rightarrow +\infty$$

■ بنابراین طبق قضیه ۲-۶ (روش بلیت) $\delta.(X) = \bar{X}$ یک برآوردگر مجاز است.

نتیجه: از مثال های ۲-۴ و ۲-۵ نتیجه می شود که در توزیع $N(\theta, \sigma^2)$ برآوردگر $a\bar{X} + b$ برای θ مجاز است اگر و فقط اگر

$$\blacksquare \quad (i) \quad 0 \leq a < 1 \quad \text{یا} \quad (ii) \quad a = 1, b = 0.$$

ب- روش نامساوی اطلاع^۱:

از این روش می توان برای اثبات مجاز بودن یا مینیماکس بودن برآوردگر استفاده کرد (Hodges & Lehmann (۱۹۵۱)

$$R(\theta, \delta) = E[(\delta - \theta)^2] = V_\theta(\delta) + b^2(\theta) \quad \text{با توجه به اینکه}$$

که در آن $b(\theta) = Bias(\delta)$ بنابراین از نامساوی اطلاع با $\gamma(\theta) = \theta$ داریم که

^۱ The Information Inequality Method

$$R(\theta, \delta) \geq \frac{[1+b'(\theta)]^{\tau}}{nI(\theta)} + b^{\tau}(\theta) \quad (۱)$$

که در آن $I(\theta) = \frac{1}{\sigma^2}$ می باشد. حال فرض کنید $\bar{X}$ برآوردگری مجاز نباشد در این صورت برآوردگری مانند δ وجود دارد که

$$R(\theta, \delta) \leq R(\theta, \bar{X}) = \frac{\sigma^2}{n} \quad \forall \theta \in \Theta \quad (۲)$$

و بنابراین

$$\frac{\sigma^2[1+b'(\theta)]^{\tau}}{n} + b^{\tau}(\theta) \leq R(\theta, \delta) \leq \frac{\sigma^2}{n} \quad \forall \theta \quad (۳)$$

می خواهیم نشان دهیم که $b(\theta) \equiv 0$ یعنی δ یک برآوردگر ناریب است.

$$(۳) \Rightarrow b^{\tau}(\theta) \leq \frac{\sigma^2}{n} \Rightarrow |b(\theta)| \leq \frac{\sigma}{\sqrt{n}} \Rightarrow b \text{ کراندار است}$$

$$(۳) \Rightarrow \frac{\sigma^2[1+b'(\theta)]^{\tau}}{n} \leq \frac{\sigma^2}{n} \Rightarrow [1+b'(\theta)]^{\tau} \leq 1 \Rightarrow b'(\theta) \leq 0.$$

پس b یک تابع غیر صعودی است.

حال فرض کنید $b(\theta) \not\equiv 0$. در این صورت دو حالت زیر را داریم:

(i) فرض کنید θ وجود داشته باشد که $b(\theta) < 0$ در این صورت

$$b(\theta) < 0 \quad \forall \theta \geq \theta.$$

$$(۳) \Rightarrow nb^{\tau}(\theta) + \sigma^2 b'^{\tau}(\theta) + 2\sigma^2 b'(\theta) \leq 0.$$

$$\Rightarrow nb^{\tau}(\theta) + 2\sigma^2 b'(\theta) \leq 0 \Rightarrow \frac{-b'(\theta)}{b^{\tau}(\theta)} \geq \frac{n}{2\sigma^2} \quad \forall \theta \geq \theta.$$

$$\Rightarrow \int_{\theta}^{\theta} \frac{-b'(t)}{b^{\tau}(t)} dt \geq \frac{n}{2\sigma^2} \int_{\theta}^{\theta} dt \Rightarrow \frac{1}{b(t)} \Big|_{\theta}^{\theta} \geq \frac{n}{2\sigma^2} (\theta - \theta.)$$

$$\Rightarrow \frac{1}{b(\theta)} - \frac{1}{b(\theta.)} \geq \frac{n}{2\sigma^2} (\theta - \theta.)$$

$$\stackrel{b(\theta) < 0}{\Rightarrow} -\frac{1}{b(\theta)} > \frac{1}{b(\theta)} - \frac{1}{b(\theta.)} \geq \frac{n}{2\sigma^2} (\theta - \theta.)$$

$$\theta \rightarrow \infty \Rightarrow -\frac{1}{b(\theta)} \rightarrow \infty \# \Rightarrow b(\theta.) = 0.$$

(ii) فرض کنید θ ی وجود داشته باشد که $b(\theta) > 0$ در این صورت $b(\theta) > 0, \forall \theta \leq \theta$

و همانند قسمت (i) می توان به یک تناقض رسید پس $b(\theta) = 0$ است.

بنابراین $b(\theta) = b'(\theta) = 0 \quad \forall \theta$ و (۱) نتیجه می دهد که $R(\theta, \delta) \geq \frac{\sigma^2}{n}$ و بنابراین

$$R(\theta, \delta) = \frac{\sigma^2}{n} = R(\theta, \bar{X}) \quad \text{یعنی } \bar{X} \text{ یک برآوردگر مجاز و همچنین مینیماکس برای } \theta \text{ است.}$$

(مینیماکس بودن به این دلیل است که اگر برآوردگری مانند δ موجود باشد که

$$R(\theta, \delta) \leq \frac{\sigma^2}{n}, \quad \forall \theta \quad \text{آنگاه} \quad \sup_{\Theta} R(\theta, \delta) \leq \sup_{\Theta} R(\theta, \bar{X}) = \frac{\sigma^2}{n}$$

$R(\theta, \delta) = R(\theta, \bar{X})$ پس $\bar{X}$ برآوردگر مینیماکس است.)

همچنین $\bar{X}$ تنها برآوردگر مینیماکس برای θ است که این موضوع از لم زیر نتیجه می شود.

لم ۱-۲: اگر تابع زیان $L(\theta, \delta)$ اکیدا محدب باشد و δ برآوردگری مجاز برای $\gamma(\theta)$ باشد و δ' هر

برآوردگر دیگری باشد که $R(\theta, \delta) = R(\theta, \delta') \quad \forall \theta$ آنگاه $P(\delta' = \delta) = 1$.

$$\text{اثبات: اگر } \delta^* = \frac{\delta + \delta'}{2} \text{ آنگاه } R(\theta, \delta^*) \stackrel{(*)}{<} \frac{1}{2} [R(\theta, \delta) + R(\theta, \delta')] = R(\theta, \delta)$$

بجز اینکه با احتمال یک $\delta = \delta'$ باشد. (نامساوی $(*)$ از اکیدا محدب بودن تابع زیان نتیجه می گردد.)

بنابراین رابطه فوق متناقض با مجاز بودن δ است بجز اینکه $P(\delta = \delta') = 1$ ■

مثال ۶-۲: میانگین نرمال بریده شده^۱

فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\theta, \sigma^2)$ باشند که σ^2 معلوم و حالات زیر را

برای θ در نظر بگیرید.

الف- اگر $\theta > \theta$ باشد آنگاه طبق قضیه ۲-۴، $\bar{X}$ غیرمجاز است و به وسیله برآوردگر

$$\delta(\bar{X}) = \begin{cases} \theta & \bar{X} < \theta \\ \bar{X} & \bar{X} \geq \theta \end{cases} = \max(\theta, \bar{X})$$

اطلاع نشان می دهیم که $\bar{X}$ برآوردگری مینیماکس برای θ است.

^۱ Truncated Normal mean

اگر $\bar{X}$ مینیمکس نباشد آنگاه برآوردگری مانند δ و $\varepsilon > 0$ وجود دارند به طوری که

$$R(\theta, \delta) \leq R(\theta, \bar{X}) - \varepsilon = \frac{\sigma^2}{n} - \varepsilon \quad \forall \theta > \theta_0$$

$$\stackrel{(۳)}{\Rightarrow} \frac{\sigma^2 [1 + b'(\theta)]^2}{n} + b^2(\theta) \leq \frac{\sigma^2}{n} - \varepsilon \quad \forall \theta > \theta_0 \quad (۴)$$

بنابراین b کراندار است $\Rightarrow b^2(\theta) \leq \frac{\sigma^2}{n} - \varepsilon \Rightarrow$ (۴)

$$\begin{aligned} (۴) \Rightarrow \frac{\sigma^2 [1 + b'(\theta)]^2}{n} &\leq \frac{\sigma^2}{n} - \varepsilon \Rightarrow [1 + b'(\theta)]^2 \leq 1 - \frac{\varepsilon n}{\sigma^2} \\ \Rightarrow 1 + b'(\theta) + b'(\theta) &\leq 1 - \frac{\varepsilon n}{\sigma^2} \Rightarrow 2b'(\theta) \leq -\frac{\varepsilon n}{\sigma^2} \\ \Rightarrow b'(\theta) &\leq \frac{-\varepsilon n}{2\sigma^2} < 0 \quad \forall \theta > \theta_0. \end{aligned}$$

کراندار بودن b و اکیداً نزولی بودن b در تناقض می‌باشد، پس $\bar{X}$ مینیمکس است. توجه کنید برآوردگر $\delta(X) = \max(\theta_0, \bar{X})$ برآوردگری بهتر از $\bar{X}$ است و می‌توان نشان داد که $\delta(X)$ نیز مینیمکس است (اثبات: تمرین) و بنابراین برآوردگر مینیمکس یکتا نیست.

$$\left(\begin{aligned} R(\theta, \delta) \leq R(\theta, \bar{X}) = \frac{\sigma^2}{n}, \quad \bar{X} > \theta_0 \Rightarrow R(\theta, \delta) = R(\theta, \bar{X}) = \frac{\sigma^2}{n} \\ \Rightarrow \sup_{\theta > \theta_0} R(\theta, \delta) = \frac{\sigma^2}{n} \end{aligned} \right)$$

ب- اگر $a \leq \theta \leq b$ باشد آنگاه $\bar{X}$ طبق قضیه ۲-۴ به وسیله برآوردگر

$$\delta^*(X) = \begin{cases} a & \bar{X} < a \\ \bar{X} & a \leq \bar{X} \leq b \\ b & \bar{X} > b \end{cases}$$

مینیمکس باشد آنگاه بایستی مغلوب کننده آن نیز مینیمکس باشد یعنی

$$\sup_{a \leq \theta \leq b} R(\theta, \delta^*) = \sup_{a \leq \theta \leq b} R(\theta, \bar{X}) = \frac{\sigma^2}{n}$$

اما $R(\theta, \delta^*) < R(\theta, \bar{X}) = \frac{\sigma^2}{n}, \forall \theta \in [a, b]$ به علاوه $R(\theta, \delta^*)$ یک تابع پیوسته از θ است، بنابراین $R(\theta, \delta^*)$ ماکسیمم خود را در یک نقطه $a \leq \theta \leq b$ اختیار می‌کند. در نتیجه $\sup_{a \leq \theta \leq b} R(\theta, \delta^*) = R(\theta, \delta^*) < \frac{\sigma^2}{n}$ که با رابطه قبل در تناقض است پس $\bar{X}$ نمی‌تواند مینیماکس باشد.

ج- اگر $-m \leq \theta \leq m$ آنگاه در بخش قبل نشان دادیم که برآوردگر

$$\delta_m(x) = m \frac{e^{+mx} - e^{-mx}}{e^{mx} + e^{-mx}} = mtghmx \quad \blacksquare$$

مثال ۷-۲: (واریانس توزیع نرمال) فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\cdot, \sigma^2)$

باشند. اگر $\tau = \frac{1}{2\sigma^2} \sim \Gamma(g, \frac{1}{\alpha})$ آنگاه $\tau | x \sim \Gamma(g + \frac{n}{2}, \frac{1}{\alpha + y})$ که $Y = \sum_{i=1}^n X_i^2$ اگر قرار

دهیم $r = \frac{n}{2}$ و تابع زیان درجه دوم باشد آنگاه برآوردگر بیز $\frac{1}{\tau} = 2\sigma^2$ عبارت است از

$$\delta^\pi(x) = E\left(\frac{1}{\tau} | x\right) = \frac{\alpha + y}{r + g - 1}$$

می‌خواهیم برآوردگرهای مجاز و غیرمجاز خطی $\tau = \frac{1}{2\sigma^2}$ به فرم $aY + b$ را تعیین کنیم.

چون $\delta^\pi(X) = \frac{1}{r + g - 1}Y + \frac{\alpha}{r + g - 1}$ پس طبق قضیه ۳-۲ به نظر می‌رسد که $aY + b$ برای

مقادیر $b > 0$ و $0 < a < \frac{1}{r-1}$ ($g > 0$) مجاز باشد. یکی از حالت‌های این برآوردگرهای

مجاز $\frac{1}{r}Y + b$ می‌باشد. از طرفی $E(Y) = n\sigma^2 = \frac{r}{\tau}$ پس $\frac{1}{r}Y$ یک برآوردگر نااریب برای $\frac{1}{\tau}$ است.

بنابراین چون $\frac{1}{r}Y + b$ که برآوردگری اریب‌دار است در رابطه

$$R(2\sigma^2, \frac{1}{r}Y + b) = R(2\sigma^2, \frac{1}{r}Y) + b^2 > R(2\sigma^2, \frac{1}{r}Y)$$

صدق می‌کند پس برآوردگری غیرمجاز است. به نظر می‌رسد که اشکالی در این خصوص وجود دارد.

قضیه ۳-۲ بیان می‌کرد که برآوردگر بیز یکتا، برآوردگری مجاز است. و دو شرط برای یکتا بودن

برآوردگر بیز وجود داشت که در این جا شرط اول یعنی متناهی بودن مخاطره بیز به ازای هر $b > 0$ و

$$0 < a < \frac{1}{r-1}$$

برقرار نیست زیرا

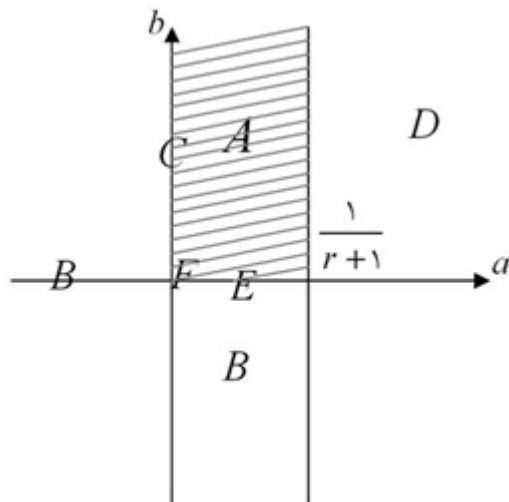
$$\begin{aligned} R(\tau\sigma^2, \frac{Y+\alpha}{r+g-1}) &= E\left(\frac{Y+\alpha}{r+g-1} - \frac{1}{\tau}\right)^2 = V\left(\frac{Y+\alpha}{r+g-1}\right) - \left[E\left(\frac{Y+\alpha}{r+g-1} - \frac{1}{\tau}\right)\right]^2 \\ &= \frac{1}{(r+g-1)^2} \left[\frac{r}{\tau^2}\right] - \left[\frac{\left(\frac{r}{\tau}\right) + \alpha}{r+g-1} - \frac{1}{\tau}\right]^2 = \frac{1}{(r+g-1)^2} \left\{ \frac{r}{\tau^2} - \left(\alpha - \frac{g-1}{\tau}\right)^2 \right\} \end{aligned}$$

بنابراین مخاطره بیز متناهی خواهد بود اگر $E\left(\frac{1}{\tau^2}\right) < \infty$ و این امید ریاضی موقعی متناهی است که

$g > 2$ باشد $\left(E\left(\frac{1}{\tau^2}\right) = \frac{\alpha^2}{(g-1)(g-2)}\right)$. با اعمال این شرط روی برآوردگر بیز $\delta^\pi(X)$ دیده می‌شود که برآوردگر بیز یکتا در شرط زیر صدق می‌کند.

$$b > 0 \text{ و } 0 < a < \frac{1}{r+1}$$

و بنابراین برای این مقادیر (a, b) برآوردگر $aY + b$ برای $\tau\sigma^2 = \frac{1}{r}$ مجاز است. (A) برای $a < 0$ و



همچنین برای $b < 0$ توسط قضیه ۲-۴ برآوردگر

$aY + b$ غیرمجاز است. (B) (زیرا مقادیری خارج

از دامنه پارامتر $(0, +\infty)$ را می‌گیرد.) حالت $a = 0$

بر آوردگر $\delta(X) = b$, $b > 0$ مجاز است زیرا

این برآوردگر تنها برآوردگری است که در $\frac{1}{\tau} = b$

دارای مخاطره صفر است $R(b, \delta) = 0$. (مثال ۲-۳)

را ملاحظه کنید. (C)

برای حالت $a > \frac{1}{r+1}$ (مسئله ۱۲، ۲) می‌توان نشان

داد که برآوردگر فوق غیرمجاز است. (D) همچنین

$$R(\tau\sigma^\tau, aY + b) = E[(aY + b - \frac{1}{\tau})^\tau] = a^\tau V(Y) + \left[aE(Y) + b - \frac{1}{\tau} \right]^\tau$$

$$= \frac{a^\tau r}{\tau^\tau} + \left[\frac{ar}{\tau} + b - \frac{1}{\tau} \right]^\tau = a^\tau r \left(\frac{1}{\tau^\tau} \right) + \left[(ar - 1) \frac{1}{\tau} + b \right]^\tau$$

پس برای $b = 0$ داریم که $a \neq 0$ ، $R(\tau\sigma^\tau, aY) = [a^\tau r + (ar - 1)^\tau] \frac{1}{\tau^\tau} = f(a)$ ،

$$f'(a) = [\tau ar + \tau r(ar - 1)] \frac{1}{\tau^\tau} = 0 \Rightarrow a(r + 1) - 1 = 0 \Rightarrow a = \frac{1}{r + 1}$$

$$f''(a) = (\tau r + \tau r^\tau) \frac{1}{\tau^\tau} > 0$$

پس $R(\frac{1}{\tau}, \frac{1}{r+1}Y) < R(\frac{1}{\tau}, aY)$ برای هر $a \neq 0$ یعنی $aY + b$ برای $b = 0$ و $a \neq \frac{1}{r+1}$ و $a \neq 0$ برآوردگری غیرمجاز است. (E) برای $b = 0$ و $a = 0$ داریم که

$$R(\frac{1}{\tau}, aY) < R(\frac{1}{\tau}, 0) \Leftrightarrow a^\tau r + (ar - 1)^\tau < 1 \Leftrightarrow a^\tau r(1 + r) - \tau ar < 0$$

$$\Leftrightarrow a(1 + r) - \tau < 0 \Leftrightarrow a < \frac{\tau}{1 + r}$$

پس برآوردگر $\delta = 0$ به وسیله برآوردگر $\delta^* = aY$ که $a < \frac{\tau}{1 + r}$ مغلوب می‌شود پس غیرمجاز

است (F). تنها حالت $b \geq 0$ و $\frac{1}{r+1}Y + b$ باقی مانده است که در زیر توسط قضیه کارلین^۱ مجاز

بودن آن را ثابت می‌کنیم. ■

فرض کنید X دارای توزیعی از خانواده نمایی با تابع چگالی

$$p_\theta(x) = \beta(\theta)e^{\theta T(x)} \quad (1)$$

نسبت به اندازه μ باشد و Θ فضای پارامتری طبیعی باشد. در این صورت Θ دارای نقاط انتهایی مانند $\underline{\theta}$ و $\bar{\theta}$ است $(-\infty \leq \underline{\theta} \leq \bar{\theta} \leq +\infty)$. برای برآورد $E_\theta(T)$ تحت تابع زیان درجه دوم، برآوردگر $aT + b$ برای $a < 0$ یا $a > 1$ غیر مجاز است و برای $a = 0$ مقداری ثابت دارد پس مجاز

^۱ Karlin (۱۹۵۸)

است. برای اینکه بتوانیم شرط قضیه کارلین را بیان کنیم راحت تر است که این برآوردگر را به صورت زیر بنویسیم.

$$\delta_{\lambda,\gamma}(x) = \frac{1}{1+\lambda}T + \frac{\gamma\lambda}{1+\lambda} \quad (2)$$

که در آن $0 < \lambda < \infty$ معادل $0 < a \leq 1$ است.

قضیه ۲-۷: (قضیه کارلین) تحت فرضیات بالا، شرط کافی برای آنکه برآوردگر (۲) برای

$g(\theta) = E_\theta(T)$ تحت زیان درجه دوم مجاز باشد این است که انتگرال تابع $e^{-\gamma\lambda\theta}[\beta(\theta)]^{-\lambda}$ در $\underline{\theta}$ و $\bar{\theta}$ واگرا باشد. یعنی برای یک $\underline{\theta} < \bar{\theta}$ انتگرالهای

$$\int_{\underline{\theta}}^{\theta^*} \frac{e^{-\gamma\lambda\theta}}{[\beta(\theta)]^{-\lambda}} d\theta \quad \text{and} \quad \int_{\theta^*}^{\bar{\theta}} \frac{e^{-\gamma\lambda\theta}}{[\beta(\theta)]^{-\lambda}} d\theta$$

موقعی که به ترتیب $\theta^* \rightarrow \underline{\theta}$ و $\theta^* \rightarrow \bar{\theta}$ به سمت بی نهایت میل کنند.

اثبات: از روش نامساوی اطلاع قضیه را ثابت می کنیم. با توجه به اینکه در خانواده نمایی (۱) داریم که

$$g(\theta) = E_\theta(T) = -\frac{\beta'(\theta)}{\beta(\theta)}, \quad g'(\theta) = V_\theta(T) = I(\theta)$$

بنابراین طبق نامساوی اطلاع برای هر برآوردگر δ داریم که

$$\begin{aligned} E_\theta \left[(\delta(X) - g(\theta))^2 \right] &= V_\theta(\delta(X)) + b^2(\theta) \geq \frac{[g'(\theta) + b'(\theta)]^2}{I(\theta)} + b^2(\theta) \\ &= \frac{[I(\theta) + b'(\theta)]^2}{I(\theta)} + b^2(\theta) \end{aligned} \quad (3)$$

که در آن $b(\theta) = E_\theta[\delta(X)] - \gamma(\theta)$. چون در خانواده نمایی یک پارامتری هستیم پس برای

برآوردگر $\delta = \delta_{\lambda,\gamma}$ نامساوی فوق به تساوی تبدیل می شود و اگر $b_{\lambda,\gamma}(\theta)$ اربیی برآوردگر $\delta_{\lambda,\gamma}$ باشد آنگاه

$$\begin{cases} b_{\lambda,\gamma}(\theta) = \frac{\lambda}{1+\lambda}[\gamma - g(\theta)] \\ b'_{\lambda,\gamma}(\theta) = -\frac{\lambda}{1+\lambda}g'(\theta) \end{cases}$$

$$E_{\theta} \left[\left(b_{\lambda, \gamma}(X) - g(\theta) \right)^{\gamma} \right] = \frac{[I(\theta) + b'_{\lambda, \gamma}(\theta)]^{\gamma}}{I(\theta)} + [b_{\lambda, \gamma}(\theta)]^{\gamma} \quad (۴)$$

حال اگر δ برآوردگری باشد که

$$E \left[\left(\delta(X) - g(\theta) \right)^{\gamma} \right] \leq E_{\theta} \left[\left(\delta_{\lambda, \gamma}(X) - g(\theta) \right)^{\gamma} \right] \quad (۵)$$

و $b.(\theta)$ اریبی آن باشد آن گاه از روابط (۳) و (۴) داریم که

$$\begin{aligned} \frac{[I(\theta) + b'(\theta)]^{\gamma}}{I(\theta)} + b.(\theta)^{\gamma} &\leq E \left[\left(\delta(X) - g(\theta) \right)^{\gamma} \right] \leq E_{\theta} \left[\left(\delta_{\lambda, \gamma}(X) - g(\theta) \right)^{\gamma} \right] \\ &= \frac{[I(\theta) + b'_{\lambda, \gamma}(\theta)]^{\gamma}}{I(\theta)} + [b_{\lambda, \gamma}(\theta)]^{\gamma} \end{aligned}$$

اگر قرار دهیم $h(\theta) = b.(\theta) - b_{\lambda, \gamma}(\theta)$ از نامساوی فوق داریم که

$$\begin{aligned} \{b.(\theta)^{\gamma} - b_{\lambda, \gamma}^{\gamma}(\theta)\} + \frac{(b'(\theta) - b'_{\lambda, \gamma}(\theta))(I(\theta) + b'(\theta) - b'_{\lambda, \gamma}(\theta))}{I(\theta)} &\leq . \\ \Rightarrow [h(\theta)]^{\gamma} + \gamma b_{\lambda, \gamma}(\theta) [b.(\theta) - b_{\lambda, \gamma}(\theta)] + \frac{h'(\theta) [h'(\theta) + \gamma(I(\theta) + b'_{\lambda, \gamma}(\theta))]}{I(\theta)} &\leq . \\ \Rightarrow h^{\gamma}(\theta) - \frac{\gamma \lambda}{1 + \lambda} (\gamma - g(\theta)) h(\theta) + \frac{\gamma}{1 + \lambda} h'(\theta) + \frac{[h'(\theta)]^{\gamma}}{I(\theta)} &\leq . \\ \Rightarrow h^{\gamma}(\theta) - \frac{\gamma \lambda}{1 + \lambda} (\gamma - g(\theta)) h(\theta) + \frac{\gamma}{1 + \lambda} h'(\theta) &\leq . \end{aligned} \quad (۶)$$

قرار می دهیم

$$k(\theta) = h(\theta) \beta^{+\lambda}(\theta) e^{\gamma \lambda \theta} \quad (۷)$$

در این صورت

$$\begin{aligned} k'(\theta) &= h'(\theta) \beta^{\lambda}(\theta) e^{\gamma \lambda \theta} + \lambda h(\theta) \beta^{\lambda-1}(\theta) \beta'(\theta) e^{\gamma \lambda \theta} + \gamma \lambda h(\theta) \beta^{\lambda}(\theta) e^{\gamma \lambda \theta} \\ &= k(\theta) \left\{ \frac{h'(\theta)}{h(\theta)} + \lambda \frac{\beta'(\theta)}{\beta(\theta)} + \gamma \lambda \right\} \\ &= k(\theta) \left\{ \frac{h'(\theta)}{h(\theta)} + \lambda (g(\theta) - \gamma) \right\} = \frac{k(\theta)}{h(\theta)} \{h'(\theta) - \lambda h(\theta) (g(\theta) - \gamma)\} \\ (۶) \Rightarrow h^{\gamma}(\theta) + \frac{\gamma}{1 + \lambda} \frac{h(\theta) k'(\theta)}{k(\theta)} &\leq . \Rightarrow \frac{h(\theta)}{k(\theta)} [h(\theta) k(\theta) + \frac{\gamma}{1 + \lambda} k'(\theta)] \leq . \end{aligned}$$

$$\frac{h(\theta)}{k(\theta)} = \beta^{-\lambda}(\theta) e^{-\gamma\lambda\theta} > 0, \\ \Rightarrow k^{\gamma}(\theta) \beta^{-\lambda}(\theta) e^{-\gamma\lambda\theta} + \frac{\gamma}{1+\lambda} k'(\theta) \leq 0. \quad (۸)$$

از این رابطه نتیجه می شود که $k'(\theta) < 0$ یعنی k یک تابع نزولی از θ است. حال نشان می دهیم که $k(\theta) = 0 \quad \forall \theta$.

فرض کنید که مقداری از θ وجود داشته باشد که $k(\theta) \neq 0$ باشد. یعنی یک مقدار θ وجود دارد که $k(\theta) < 0$ (i) یا $k(\theta) > 0$ (ii) است.

حالت (i): اگر $k(\theta) < 0$ آنگاه $\forall \theta \geq \theta, k(\theta) < 0$ حال از رابطه (۸) داریم که

$$\frac{d}{d\theta} \left[\frac{1}{k(\theta)} \right] = \frac{-k'(\theta)}{k^{\gamma}(\theta)} \geq \frac{1+\lambda}{\gamma} \beta^{-\lambda}(\theta) e^{-\gamma\lambda\theta} \\ \Rightarrow \int_{\theta}^{\theta^*} \frac{d}{d\theta} \left[\frac{1}{k(t)} \right] dt \geq \frac{1+\lambda}{\gamma} \int_{\theta}^{\theta^*} \beta^{-\lambda}(t) e^{-\gamma\lambda t} dt \\ \Rightarrow -\frac{1}{k(\theta)} > \frac{1}{k(\theta^*)} - \frac{1}{k(\theta)} \geq \frac{1+\lambda}{\gamma} \int_{\theta}^{\theta^*} \beta^{-\lambda}(t) e^{\gamma\lambda\theta} dt$$

حال اگر $\theta^* \rightarrow \infty$ آنگاه سمت راست به ∞ میل می کند که با کمتر بودن آن از $-\frac{1}{k(\theta)}$ در تناقض

است بجز اینکه $k(\theta) = 0$.

(ii) به طور مشابه می توان نشان داد که نمی تواند $k(\theta) > 0$ باشد.

بنابراین برای هر θ ، $k(\theta) = 0$ و بنابراین برای هر θ ، $h(\theta) = 0$ خواهد بود، یعنی در نامساوی (۶) حالت تساوی برقرار می شود و در نتیجه در نامساوی (۵) نیز حالت تساوی خواهیم داشت و این نشان

می دهد که برآوردگر $\delta_{\lambda, \gamma}(X) = \frac{T + \gamma\lambda}{1 + \lambda}$ برآوردگری مجاز است. ■

مثال ۷-۲ (ادامه): با توجه به اینکه در توزیع $N(\cdot, \sigma^2)$ داریم که

$$f_{\sigma^2}(x) = (\gamma\pi\sigma^2)^{-\frac{n}{2}} e^{-\frac{1}{2\sigma^2} \sum x_i^2} = \pi^{-\frac{n}{2}} \tau^r e^{-\tau y} \quad y = \sum x_i^2, \tau = \frac{1}{2\sigma^2}$$

پس این توزیع به فرم خانواده نمایی (۱) با مقادیر زیر است

$$T(X) = \frac{Y}{r}, \quad \theta = -r\tau, \quad \beta(\theta) = \left(\frac{-\theta}{r}\right)^r, \quad \underline{\theta} = -\infty, \quad \bar{\theta} = 0.$$

(توجه کنید که پارامترهای بگونه‌ای انتخاب شده‌اند که $E(T(X)) = \frac{1}{\tau}$ باشد). حال طبق قضیه کارلین

برآوردگر $\frac{1}{1+\lambda} \frac{Y}{r} + \frac{\gamma\lambda}{1+\lambda}$ برای $\frac{1}{\tau}$ مجاز است اگر انتگرال‌های

$$\int_0^c e^{-\gamma\lambda\theta} \theta^{-r\lambda} d\theta \quad \text{و} \quad \int_{-\infty}^{-c} e^{-\gamma\lambda\theta} \left(\frac{-\theta}{r}\right)^{-r\lambda} d\theta = k \int_c^{\infty} e^{\gamma\lambda\theta} \theta^{-r\lambda} d\theta$$

هر دو واگرا باشند. انتگرال اولی موقعی واگرا است که $r\lambda \leq 1$ و $\gamma = 0$ یا $\gamma\lambda > 0$ باشد. و در انتگرال

دوم عامل $e^{\gamma\lambda\theta}$ نقشی در انتگرال ندارد و موقعی واگرا است که $r\lambda \geq 1$ باشد. اگر این شرایط را با

هم ادغام کنیم آنگاه برآوردگر $\frac{1}{1+\lambda} \frac{Y}{r} + \frac{\gamma\lambda}{1+\lambda}$ مجاز است اگر

$$(i) \quad \gamma = 0, \lambda = \frac{1}{r} \quad \text{یا} \quad (ii) \quad \lambda \geq \frac{1}{r}, \gamma > 0.$$

اگر قرار دهیم $a = \frac{1}{(1+\lambda)r}$ و $b = \frac{\gamma\lambda}{1+\lambda}$ آنگاه برآوردگر $aY + b$ برای $\frac{1}{\tau}$ مجاز است اگر

$$(i) \quad a = \frac{1}{1+r}, b = 0 \quad \text{یا} \quad (ii) \quad 0 < a \leq \frac{1}{1+r}, b > 0.$$

یعنی برآوردگر $aY + b$, $b \geq 0$ برای $\frac{1}{1+r}Y + b$ مجاز است. با توجه به نواحی بدست آمده در این مثال

می‌توان گفت که برآوردگر $aY + b$ برای $\frac{1}{\tau} \sigma^2$ مجاز است اگر و فقط اگر

$$(i) \quad 0 \leq a \leq \frac{1}{r+1}, b > 0 \quad (ii) \quad a = \frac{1}{r+1}, b = 0.$$

بنابراین برآوردگر $cY + d$ برای $\frac{1}{\tau} \sigma^2$ مجاز است اگر و فقط اگر

$$(i) \quad 0 \leq c \leq \frac{1}{n+2}, d > 0 \quad (ii) \quad c = \frac{1}{n+2}, d = 0 \quad (r = \frac{n}{2}) \quad (c = \frac{a}{2}, d = \frac{b}{2})$$

یعنی برآوردگر MRE , $\delta^*(X) = \frac{1}{n+2} \sum X_i^2$ برای σ^2 مجاز است. ■

مثال ۲-۸: (واریانس نرمال با میانگین مجهول)

فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از $N(\mu, \sigma^2)$ باشند. در مثال (۷) دیدیم که

اگر $\mu = 0$ باشد، آنگاه برآوردگر MRE، $\frac{1}{n+2} \sum_{i=1}^n X_i^2$ ، برای σ^2 مجاز است. حال اگر μ نامعلوم

باشد آیا برآوردگر MRE پارامتر σ^2 یعنی $\frac{1}{n+1} \sum_{i=1}^n (X_i - \bar{X})^2$ برای σ^2 مجاز است.

حل: اشتاین^۱ (۱۹۶۴) نشان داد که برآوردگر $\frac{1}{n+1} \sum_{i=1}^n (X_i - \bar{X})^2$ مجاز نیست و بوسیله برآوردگر

$$\delta^* = \min\left(\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n+1}, \frac{\sum_{i=1}^n X_i^2}{n+2}\right)$$
 مغلوب می‌شود (اثبات تمرین) دلیل توجیهی این برآوردگر

این است که فرض کنید بخواهیم آزمون $H_0: \mu = 0$ vs $H_1: \mu \neq 0$ را آزمون کنیم در این صورت

$$\text{فرض } H_0 \text{ را قبول می‌کنیم اگر } \frac{|\sqrt{n}\bar{X}|}{\sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}} < C \text{ و در نتیجه } \sigma^2 \text{ را با}$$

$$\frac{1}{n+2} \sum_{i=1}^n X_i^2 \text{ برآورد می‌کنیم و در غیر این صورت فرض } H_1 \text{ را قبول می‌کنیم و } \sigma^2 \text{ را با}$$

$$\frac{1}{n+1} \sum_{i=1}^n (X_i - \bar{X})^2 \text{ برآورد می‌کنیم. اگر قرار دهیم } C = \sqrt{\frac{n-1}{n+1}} \text{ آنگاه نامساوی فوق به نامساوی}$$

$$\frac{1}{n+2} \sum_{i=1}^n X_i^2 < \frac{1}{n+1} \sum_{i=1}^n (X_i - \bar{X})^2 \text{ تبدیل می‌شود و بنابراین با توجه به آزمون فوق}$$

برآوردگر δ^* حاصل می‌گردد.

$$n\bar{X}^2 < \frac{1}{n+1} \sum_{i=1}^n (X_i - \bar{X})^2 \Rightarrow n\bar{X}^2 < \frac{1}{n+1} \sum_{i=1}^n X_i^2 - \frac{n}{n+1} \bar{X}^2$$

$$\Rightarrow \frac{n(n+2)}{n+1} \bar{X}^2 < \frac{1}{n+1} \sum_{i=1}^n X_i^2$$

$$\Rightarrow \frac{n}{n+1} \bar{X}^2 < \frac{1}{(n+1)(n+2)} \sum_{i=1}^n X_i^2 = \left(\frac{1}{n+1} - \frac{1}{n+2}\right) \sum_{i=1}^n X_i^2$$

^۱ Stein

$$\Rightarrow \frac{1}{n+2} \sum_{i=1}^n X_i^2 < \frac{1}{n+1} [\sum X_i^2 - n\bar{X}^2] = \frac{1}{n+1} \sum (X_i - \bar{X})^2$$

برآوردگر δ^* در رده بزرگتری از کلاس‌ها به صورت $\delta(\bar{X}, S) = \varphi(\frac{\bar{X}}{S}) S^2$ قرار دارد زیرا

$$\begin{aligned} \delta^* &= \min\left(\frac{S^2}{n+1}, \frac{\sum X_i^2}{n+2}\right) = \min\left(\frac{1}{n+1}, \frac{S^2 + n\bar{X}^2}{(n+2)S^2}\right) S^2 \\ &= \min\left(\frac{1}{n+1}, \frac{1}{n+2} + \frac{n}{n+2} \left(\frac{\bar{X}}{S}\right)^2\right) S^2 = \varphi\left(\frac{\bar{X}}{S}\right) S^2 \end{aligned}$$

برآوردگرهای به فرم $\delta(\bar{X}, S)$ نیز بر برآوردگر $\frac{1}{n+1} \sum (X_i - \bar{X})^2$ غلبه دارند (اثبات: تمرین) ■

مثال ۹-۲: (توزیع دو جمله‌ای) فرض کنید $X \sim B(n, p)$ باشد. برای چه مقادیر a و b برآوردگر

$a\left(\frac{X}{n}\right) + b$ تحت تابع زیان درجه دوم مجاز و برای چه مقادیر غیرمجاز است.

$$f_X(x) = p(X=x) = \binom{n}{x} p^x (1-p)^{n-x} = \binom{n}{x} (1-p)^n e^{\frac{x}{n} \ln \frac{p}{1-p}}$$

قرار می‌دهیم

$$\theta = n \ln \frac{p}{1-p}, \quad T(X) = \frac{X}{n}, \quad \beta(\theta) = (1-p)^n = \left(1 + e^{\frac{\theta}{n}}\right)^{-n}$$

$$g(\theta) = E_\theta(T(X)) = E_\theta\left(\frac{X}{n}\right) = p = \frac{e^{\frac{\theta}{n}}}{1 + e^{\frac{\theta}{n}}}$$

$$0 < p < 1 \Rightarrow \underline{\theta} = -\infty, \quad \bar{\theta} = +\infty$$

طبق قضیه کارلین برآوردگر $\delta_{\lambda, \gamma}(x) = \frac{X}{n(1+\lambda)} + \frac{\gamma\lambda}{1+\lambda}$ برای p مجاز است اگر دو انتگرال زیر

واگرا باشند

$$\int_{-\infty}^c e^{-\gamma\lambda\theta} \left(1 + e^{\frac{\theta}{n}}\right)^{\lambda n} d\theta = \int_c^{\infty} e^{\gamma\lambda\theta} \left(1 + \frac{1}{e^{\frac{\theta}{n}}}\right)^{\lambda n} d\theta \quad (1)$$

$$\int_c^\infty e^{-\gamma\lambda\theta} \left(1 + e^{\frac{\theta}{n}}\right)^{\lambda n} d\theta \quad (2)$$

اگر $\lambda < 0$ باشد، آنگاه $e^{-\gamma\lambda\theta} \left(1 + e^{\frac{\theta}{n}}\right)^{\lambda n} \leq e^{-\gamma\lambda\theta}$ و بنابراین هر دو انتگرال نمی‌توانند همزمان واگرا باشند. اگر $\lambda = 0$ آنگاه هر دو انتگرال واگرا هستند. اگر $\lambda > 0$ باشد آنگاه

$$1 + e^{\frac{\theta}{n}} > e^{\frac{\theta}{n}} \Rightarrow \left(1 + e^{\frac{\theta}{n}}\right)^{\lambda n} > \left(e^{\frac{\theta}{n}}\right)^{\lambda n} \\ \Rightarrow e^{-\gamma\lambda\theta} \left(1 + e^{\frac{\theta}{n}}\right)^{\lambda n} > e^{\lambda\theta} e^{-\gamma\lambda\theta} = e^{(1-\gamma)\lambda\theta}$$

در این صورت انتگرال دوم به ازای $1 - \gamma \geq 0$ یا $\gamma \leq 1$ واگراست و در انتگرال اول عامل $\left(1 + e^{\frac{\theta}{n}}\right)^{\lambda n}$ در همگرایی یا واگرایی تاثیر نمی‌گذارد پس این انتگرال به ازای $\gamma\lambda \geq 0$ یا $\gamma \geq 0$ واگرا است.

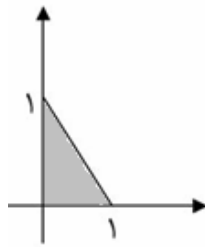
بنابراین در انتگرال در صورتی همزمان واگرا هستند که $(\lambda > 0, 0 \leq \gamma \leq 1)$ یا $\lambda = 0$ با $a = \frac{1}{1+\lambda}$ و

$b = \frac{\gamma\lambda}{1+\lambda}$ نتیجه می‌شود که برآوردگر $a\left(\frac{X}{n}\right) + b$ برای p مجاز است اگر

$$a = 1, b = 0 \quad \text{یا} \quad (0 < a < 1, b \geq 0, a + b = \frac{1+\gamma\lambda}{1+\lambda} \leq 1)$$

که با ادغام این دو حالت داریم که $0 < a \leq 1, b \geq 0, a + b \leq 1$

برای $0 \leq b \leq 1$ و $a = 0$ برآوردگر $\delta = b$ مجاز است (زیرا این برآوردگر تنها برآوردگری است که در $p = b$ دارای مخاطره صفر است). به راحتی می‌توان



نشان داد که برای سایر مقادیر a و b برآوردگر $a\left(\frac{X}{n}\right) + b$ غیرمجاز است (مسئله ۲-۴). بنابراین

برآوردگر $a\left(\frac{X}{n}\right) + b$ برای p مجاز است اگر و فقط اگر (a, b) در مثلث بسته زیر قرار گیرند:

$$\{(a, b) | a \geq 0, b \geq 0, a + b \leq 1\}$$

■

از قضیه ۷-۲ نتیجه زیر حاصل می‌شود.

نتیجه ۱-۲: اگر فضای پارامتری طبیعی در خانواده نمایی $\beta(\theta)e^{\theta T(x)}$ تمام اعداد حقیقی باشد، یعنی $\bar{\theta} = +\infty$ و $\underline{\theta} = -\infty$ آنگاه T برای برآورد $E_{\theta}(T)$ تحت تابع زیان درجه دوم مجاز است.

اثبات: با $\lambda = 0$ و $\gamma = 1$ هر دو انتگرال قضیه کارلین واگرا می‌شوند موقعی که $\theta^* \rightarrow \pm\infty$ ■

با توجه به نتیجه فوق در توزیع نرمال $N(\mu, \sigma^2)$ با σ^2 معلوم

$$f_{\mu}(x) \propto e^{-\frac{n\mu\bar{x} + \frac{n\mu^2}{2\sigma^2}}{\sigma^2}}$$

$\bar{X}$ برای μ مجاز است و در توزیع دوجمله‌ای (مثال ۹-۲) $\frac{X}{n}$ برای p مجاز است و در توزیع پواسون $f_{\theta}(x) = e^{-\mu} e^{n \ln \mu x}$ که $\theta = n \ln \mu \in (-\infty, +\infty)$ برآوردگر $\bar{X}$ برای μ مجاز است.

رابطه برآوردگرهای مجاز و مینیماکس

قضیه ۸-۲: اگر δ یک تصمیم مجاز با مخاطره ثابت باشد آنگاه δ یک تصمیم مینیماکس است.

اثبات: فرض کنید δ یک تصمیم مینیماکس نباشد در این صورت تصمیم دیگری مانند δ^* وجود دارد

$$\sup_{\Theta} R(\theta, \delta^*) < \sup_{\Theta} R(\theta, \delta) = c \quad \text{که}$$

$$R(\theta, \delta^*) \leq \sup_{\Theta} R(\theta, \delta^*) < c = R(\theta, \delta) \quad \forall \theta \in \Theta \quad \text{بنابراین}$$

و این با مجاز بودن δ در تناقض است پس δ یک تصمیم مینیماکس است. ■

نتیجه ۲-۲: تحت فرضیات نتیجه ۱-۲ و تحت تابع زیان $\frac{(\delta - g(\theta))^2}{V_{\theta}(T)}$ برآوردگر T برای

$$g(\theta) = E_{\theta}(T)$$

اثبات: طبق نتیجه ۱-۲، T برای $g(\theta)$ تحت تابع زیان $(\delta - g(\theta))^2$ مجاز است پس تحت تابع

زیان $\frac{(\delta - g(\theta))^2}{V_{\theta}(T)}$ نیز مجاز است. حال چون تحت این تابع زیان T دارای مخاطره ثابت است پس

طبق قضیه ۸-۲، T برای $g(\theta)$ مینیماکس است. یکتایی این برآورد از لم ۱-۲ حاصل می‌شود (تابع

زیان اکیدا محدب است). ■

قضیه ۹-۲: اگر δ_m یک تصمیم مینیماکس یکتا باشد آنگاه δ_m یک تصمیم مجاز است.

اثبات: اگر δ_m مجاز نباشد آنگاه تصمیم دیگری مانند δ وجود دارد که بر آن غلبه می‌کند و در نتیجه $\underset{\Theta}{Sup} R(\theta, \delta) \leq \underset{\Theta}{Sup} R(\theta, \delta_m)$. یعنی δ نیز یک تصمیم مینیماکس است و این با یکتا بودن

برآوردگر مینیماکس δ_m در تناقض است. پس δ_m یک تصمیم مجاز است. ■

مثال ۲-۱۰: اگر $X \sim B(n, p)$ آنگاه طبق نتیجه ۲-۲ تحت تابع زیان $\frac{(\delta - p)^2}{pq}$ برآوردگر $\frac{X}{n}$ یک

برآوردگر یکتای مینیماکس است پس طبق قضیه ۲-۹ یک برآوردگر مجاز است. ■

مثال ۲-۱۱: فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیع $N(\theta, \sigma^2)$ با σ^2 معلوم باشند. تحت تابع زیان درجه دوم

الف- در مثال‌های قبل دیدیم که $\bar{X}$ برآوردگر یکتای مینیماکس برای θ است پس $\bar{X}$ برای θ مجاز است.

ب- در مثال‌های قبل دیدیم که $\bar{X}$ برآوردگر مجاز با مخاطره ثابت $R(\theta, \bar{X}) = \frac{\sigma^2}{n}$ است پس برای θ مینیماکس است. ■

مثال ۲-۱۲: (دو توزیع دو جمله‌ای) فرض کنید $X \sim B(m, p)$ و $Y \sim B(n, \pi)$ از یکدیگر مستقل باشند. نشان دهید که برآوردگر $a\frac{X}{m} + b\frac{Y}{n} + c$ تحت تابع زیان درجه دوم و تحت شرایط زیر برای p مجاز است.

$$(0 \leq a < 1, 0 \leq c \leq 1, 0 \leq a+c \leq 1, 0 \leq b+c \leq 1, 0 \leq a+b+c \leq 1) \text{ یا } (a=1, b=c=0)$$

اثبات: فرض کنید برآوردگر دیگری مانند $\delta(X, Y)$ موجود باشد که

$$\begin{aligned} E\left(a\frac{X}{m} + b\frac{Y}{n} + c - p\right)^2 &\geq E[(\delta(X, Y) - p)^2] \quad \forall p \\ &\Rightarrow \sum_{x=0}^m \sum_{k=0}^n \left(a\frac{x}{m} + b\frac{k}{n} + c - p\right)^2 P(X=x, Y=k) \\ &\geq \sum_{x=0}^m \sum_{k=0}^n (\delta(x, k) - p)^2 P(X=x, Y=k) \end{aligned} \quad (1)$$

حال اگر $\pi \rightarrow 0$ آنگاه (۱) نتیجه می‌دهد که

$p(X = x, Y = k) = \binom{m}{x} p^x q^{m-x} \binom{n}{k} \pi^k (1-\pi)^{n-k}$ که به ازای $k = 0$ برابر $p(X = x)$ است.

$$\sum_{x=0}^m (a \frac{x}{m} + c - p)^r P(X = x) \geq \sum_{x=0}^m (\delta(x, \cdot) - p)^r P(X = x)$$

طبق مثال ۹-۲، $a \frac{x}{m} + c$ تحت شرایط این مثال مجاز است پس بایستی $\delta(x, \cdot) = a \frac{x}{m} + c$ برای هر $x = 0, 1, \dots, m$ بنابرین عبارت $k = 0$ در (۱) حذف می‌شود. عبارات دیگر شامل یک عامل π هستند که می‌تواند حذف شود و بنابرین (۱) نتیجه می‌دهد که

$$\sum_{x=0}^m \sum_{k=1}^n (a \frac{x}{m} + b \frac{k}{n} + c - p)^r \binom{m}{x} p^x q^{m-x} \binom{n}{k} \pi^{k-1} (1-\pi)^{n-k} \geq \sum_{x=0}^m \sum_{k=1}^n (\delta(x, k) - p)^r \binom{m}{x} p^x q^{m-x} \binom{n}{k} \pi^{k-1} (1-\pi)^{n-k} \quad (2)$$

دو مرتبه با $\pi \rightarrow 0$ رابطه (۲) نتیجه می‌دهد که

$$\sum_{x=0}^m (a \frac{x}{m} + \frac{b}{n} + c - p)^r P(X = x) \geq \sum_{x=0}^m (\delta(x, 1) - p)^r P(X = x)$$

و چون تحت شرایط این مسئله $0 \leq a + \frac{b}{n} + c \leq 1$ پس طبق مثال ۹-۲ برآوردگر $a \frac{x}{m} + \frac{b}{n} + c$ برای

p مجاز است. بنابرین (۲) نتیجه می‌دهد که برای هر $x = 0, 1, \dots, m$ ، $\delta(x, 1) = a \frac{x}{m} + \frac{b}{n} + c$ به

همین ترتیب اگر با استقراء روی k پیش برویم نتیجه می‌شود که

$$\delta(x, k) = a \frac{x}{m} + b \frac{k}{n} + c \quad \forall x=0, \dots, m, k=0, \dots, n$$

شرایط مسئله برآوردگری مجاز برای p است.

در خارج از ناحیه داده شده در این مثال می‌توان نشان داد که برآوردگر $\delta(X, Y)$ غیرمجاز

است (مسئله ۲۱، ۲). پس $\delta(X, Y)$ مجاز است اگر و فقط اگر در شرایط این مسئله صدق کند. ■

- 1- Berger, J. O. (1985). Statistical Decision Theory and Bayesian Analysis, Second Edition, Springer-Verlag, New York.
 - 2- Lehman E. L. and Casella G. (1998). Theory of Point Estimation, Second Edition, Springer, New York.
 - 3- Casella, G. and Berger, R. L. (1990). Statistical Inference, Wodsworth & Brooks, California.
- ۴- بهبودیان، جواد (۱۳۷۴). تصمیم آماری، انتشارات دانشگاه شیراز.