

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

دعای مطالعه

اَللّٰهُمَّ اَخْرِجْنِيْ مِنْ ظُلُمَاتِ الْوَهْمِ وَاكْرِمْ نِيْ بِنُورِ الْفَهْمِ

اَللّٰهُمَّ افْتَحْ عَلَيْنَا اَبْوَابَ رَحْمَتِكَ وَ انْشُرْ عَلَيْنَا خَزَائِنَ عُلُوْمِكَ

بِرَحْمَتِكَ يَا اَرْحَمَ الرَّاحِمِيْنَ

پروردگارا، خارج کن مرا از تاریکی های فکر و گرامی بدار به نور فهم

پروردگارا، بکشی بر ما درهای رحمت را و بگستران کنج های دانشت را به امید رحمت

تو ای مهربان ترین مهربانان

بیایید به حقوق دیگران احترام بگذاریم

دوست عزیز، این کتاب حاصل دسترنج چندین ساله ی مؤلف، مترجم و ناشر آن است. تکثیر و فروش آن به هر شکلی بدون اجازه از پدیدآورنده کاری غیراخلاقی، غیرقانونی، غیرشرعی و کسب درآمد از دسترنج دیگران است، نتیجه ی این عمل نادرست، موجب رواج بی اعتمادی در جامعه و بروز پی آمدهای ناگوار در زندگی و محیط ناسالم برای خود و فرزندانمان می گردد.



بانک سوالات ده سالانه

آمار زیستی

«دکتری تخصصی»

(همراه با پاسخنامه تشریحی)

مولفین و گردآورندگان:

دکتر رسول علیمی

«استادیار آمار زیستی، دانشگاه علوم پزشکی تربت حیدریه»

دکتر فاطمه خراشادیزاده

«استادیار آمار زیستی، دانشگاه علوم پزشکی نیشابور»

دکتر الهام شرباف عیدگاهی

«دکتری تخصصی آمار زیستی، معاونت پژوهشی و فناوری دانشگاه علوم پزشکی مشهد»

سمیه برزنونی

«کارشناسی ارشد آمار زیستی، معاونت آموزش، تحقیقات و فناوری دانشگاه علوم پزشکی تربت حیدریه»

عنوان و نام پدیدآور	AGK بانک سوالات ده سالانه آمار زیستی «دکتری تخصصی» (همراه با پاسخنامه تشریحی) / مولفین و گردآوردندگان رسول علیمی ... [و دیگران].
مشخصات نشر	تهران: آناهید گستر خلیلی، ۱۴۰۳.
مشخصات ظاهری	۶۸۰ ص. : جدول، نمودار.
شابک	978-622-4907-01-1
وضعیت فهرست نویسی	فیپا :
یادداشت	مولفین و گردآوردندگان رسول علیمی، فاطمه خراشادیزاده، الهام شعر باف عیدگاهی، سمیه برزنونی.
موضوع	زیست سنجی -- آزمون ها و تمرین ها (عالی) Biometry -- Examinations, questions, etc (Higher) دانشگاه ها و مدارس عالی -- ایران -- آزمون دوره های تحصیلات تکمیلی Universities and colleges -- Iran -- Graduate Record Examination
شناسه افزوده	علیمی، رسول، ۱۳۶۵-، گردآورنده
رده بندی کنگره	QH۳۲۳/۵ :
رده بندی دیویی	۵۷۴/۱۵۱۹۵ :
شماره کتابشناسی ملی	۹۷۴۶۱۲۲ :



انتشارات آناهید گستر خلیلی

نام کتاب: AGK بانک سوالات ده سالانه آمار زیستی «دکتری تخصصی» (همراه با پاسخنامه تشریحی)

مولفین و گردآوردندگان: دکتر رسول علیمی

دکتر فاطمه خراشادیزاده . دکتر الهام شعر باف عیدگاهی . سمیه برزنونی

ناشر: آناهید گستر خلیلی

نوبت و سال چاپ: اول . ۱۴۰۳

شمارگان: ۱۰۰۰

چاپ و صحافی: گیلان

مدیر تولید: اقبال شرقی

مدیر فنی و هنری: مریم آرده

تایپ و صفحه آرایی: بیتا اندوژفر

بهاء: ۵۵۰۰۰۰ تومان

آموزشگاه / فروشگاه: ۰۲۱۶۶۵۶۸۶۲۱

مرکز پخش: ضلع جنوب غربی میدان انقلاب . جنب بانک پارسیان . مجتمع تجاری پارس . طبقه اول

مدیر فروش: ۰۹۹۰۲۱۰۵۲۹۲

مرکز فروش: ۰۲۱۶۶۵۶۹۲۱۶

 agkpublisher@gmail.com

 www.agkins.com

طلیحه سخن مؤلفین:

موفقیت در آزمون‌های کنکور و رقابت فشرده برای ورود به مقاطع تحصیلی بالاتر، نیازمند تلاش مستمر، مطالعه بیش‌تر و استفاده از منابع مناسب جهت دستیابی به این هدف می‌باشد. مرور و یادگیری سوالات و آزمون‌های سال‌های گذشته به‌عنوان منبعی مناسب برای آمادگی جهت آزمون‌های آینده می‌تواند بسیار مفید و راهگشا باشد.

این مجموعه به‌عنوان یکی از اولین کتاب‌ها با هدف کمک به متقاضیان شرکت‌کننده در کنکور دکتری رشته آمار زیستی تهیه و تدوین شده است.

در نگارش این مجموعه سعی شده است تا سوالات مربوط به آزمون‌های سال‌های ۱۳۸۷ الی ۱۴۰۳ در مباحث آمار زیستی، روش‌های آماری، تحلیل بقاء و کارآزمایی بالینی با راه حل تقریباً کامل و همراه با نکات ضروری ارائه شوند تا دانشجویان عزیز بتوانند با بهره‌گیری از آن درصد بالایی از نمرات این درس را در کنکور کسب نمایند. امید که این مجموعه توانسته باشد سهم کوچکی در ارتقای علمی خوانندگان محترم داشته باشد.

بی‌شک در این‌گونه مجموعه‌ها که در برگیرنده حجم بالایی از مطالب هستند بروز برخی از نواقص و کاستی‌ها اجتناب‌ناپذیر است. ارائه هر گونه راهکار و پیشنهاد ارزشمندی از سوی خوانندگان محترم این مجموعه در برطرف کردن نقایص احتمالی و بهبود آن می‌تواند ما را در ارتقای این مجموعه یاری دهد. در این راستا شما می‌توانید از طریق پست الکترونیکی rasulalimi@yahoo.com ما را از نظرات سازنده خود بهره‌مند سازید.

با آرزوی توفیقات روزافزون

طلیحه سخن مؤلفین:

در سال‌های اخیر، کتاب‌های متفاوتی در زمینه‌ی حل سوالات کنکور در رشته‌های مختلف به چاپ رسیده است. اما کتابی با پاسخ کاملاً تشریحی و به‌روز در رشته آمار زیستی و در مقطع دکتری موجود نمی‌باشد. از آن‌جا که هر سال تعداد کثیری از دانشجویان متقاضی ادامه تحصیل در رشته آمار زیستی و به طبع آن؛ شرکت در کنکور وزارت بهداشت هستند، ما بر آن شدیم که در کتاب حاضر با پاسخ تشریحی سوالات کنکور دکتری آمار زیستی وزارت بهداشت از سال ۱۴۰۳-۱۳۸۷ راه‌گشایی برای این عزیزان باشیم. این کتاب شامل حل سوالات دروس استنباط آماری، روش‌های آمار زیستی، تحلیل چند متغیره، تحلیل بقا و کارآزمایی بالینی است. کتب مرجع حل سوالات به شرح زیر است:

Applied Linear Statistical Models- Micheal H. Kutner, Christopher J. Nachtsheim, John Neter, William Li-Fifth Edition.

Applied Multivariate Techniques. Subhash Sharma- Wiley.

تحلیل آماری چند متغیری کاربردی، ریچارد آ جانسون، ترجمه دکتر حسینعلی نیرومند، انتشارات: دانشگاه فردوسی مشهد.

Survival Analysis- David G. Kleinbaum, Mitchel Klein- Third Edition.

کارآزمایی‌های بالینی، استوارت پوکاک، ترجمه دکتر آیت‌اللهی.

آمار نظری پردازش، تالیف عبدالرضا محمدی (جلد اول و جلد دوم).

امید است که کتاب حاضر گامی مثبت در جهت پیشرفت شما عزیزان برداشته باشد. در جهت هر آن‌چه بهتر کردن این کتاب خواهم نمودیم نظرات و پیشنهادات خود را با در میان بگذارید.

با تشکر

الهام شعرباف عیدگاهی
elham.Shaarbaf@gmail.com

سمیه برزنونی
s.barzanouni@gmail.com

فهرست مطالب

صفحه

عنوان

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۸۷-۸۸

سوالات ۹

پاسخنامه تشریحی ۱۹

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۸۸-۸۹

سوالات ۵۶

پاسخنامه تشریحی ۷۱

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۸۹-۹۰

سوالات ۱۱۶

پاسخنامه تشریحی ۱۳۰

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۰-۹۱

سوالات ۱۶۶

پاسخنامه تشریحی ۱۸۲

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۱-۹۲

سوالات ۲۱۶

پاسخنامه تشریحی ۲۳۰

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۲-۹۳

سوالات ۲۶۲

پاسخنامه تشریحی ۲۷۶

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۳-۹۴

سوالات ۳۰۷

پاسخنامه تشریحی ۳۲۱

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۴-۹۵

سوالات ۳۴۸

پاسخنامه تشریحی ۳۶۲

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۵-۹۶

سوالات ۳۹۰

پاسخنامه تشریحی ۴۰۴

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۶-۹۷

سوالات ۴۳۱

پاسخنامه تشریحی ۴۴۴

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۷-۹۸

سوالات ۴۶۶

پاسخنامه تشریحی ۴۸۰

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۸-۹۹

سوالات ۵۰۵

پاسخنامه تشریحی ۵۱۹

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۹۹-۴۰۰

سوالات ۵۴۱

پاسخنامه تشریحی ۵۵۶

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۴۰۰-۴۰۱

سوالات ۵۷۳

پاسخنامه تشریحی ۵۸۶

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۴۰۱-۴۰۲

سوالات ۶۰۵

پاسخنامه تشریحی ۶۱۹

دکتری تخصصی (Ph.D) رشته آمار زیستی سال ۴۰۲-۴۰۳

سوالات ۶۴۱

پاسخنامه تشریحی ۶۵۵

سال

سوالات

۸۷-۸۸

دکتری تخصصی (Ph.D) آمارزیستی

رسول علیمی - فاطمه خراشادیزاده

آمار ریاضی و احتمالات

۱. نمونه تصادفی $X_1, X_2, \dots, X_n$ را از توزیعی با واریانس σ^2 در نظر می‌گیریم. فرض کنیم $\bar{X} = \sum_{i=1}^n X_i / n$. امید

ریاضی $T = \sum_{i=1}^n (X_i - \bar{X})^2$ برابر است با:

$$(1) \sigma^2 \quad (2) n\sigma^2$$

$$(3) (n-1)\sigma^2 \quad (4) \sigma^2 + n$$

۲. X_1, X_2, X_3 سه متغیر تصادفی مثبت و هم توزیع می‌باشند. فرض کنید، $\bar{X} = \frac{X_1 + X_2 + X_3}{3}$.

$$E\left(\frac{X_1}{\bar{X}}\right)$$

(۱) وجود ندارد. (۲) برابر $\frac{1}{3}$ است.

(۳) برابر ۳ است. (۴) برابر یک است.

۳. جعبه‌ای شامل ۶ مهره به شماره‌های ۱ تا ۶ می‌باشد. یک نمونه تصادفی چهارتایی بدون جایگذاری از این جعبه بیرون می‌آوریم. احتمال این که ۳ مهره در این نمونه دارای شماره‌های زوج باشند، برابر است با:

$$(1) \frac{1}{10} \quad (2) \frac{8}{15} \quad (3) \frac{2}{10} \quad (4) \frac{2}{5}$$

۴. متغیر تصادفی X دارای چگالی $f(x) = \frac{1}{\pi(1+x^2)}$ ، $-\infty < x < +\infty$ می‌باشد. تابع توزیع این متغیر تصادفی برابر

است با:

(۱) $\frac{1}{2} + \frac{1}{\pi} \arctan x$ (۲) $\frac{1}{\pi} \arctan x$

(۳) $\frac{1}{2} - \arctan x$ (۴) $\frac{1}{2} + \arctan x$

۵. X و Y دو متغیر تصادفی مستقل و هر یک دارای توزیع پواسن با میانگین ۲ و ۳ می‌باشد. $P(X+Y=0)$ برابر است با:

(۱) $e^{-2} + e^{-3}$ (۲) e^{-6}

(۳) e^{Δ} (۴) $e^{-\Delta}$

۶. فرض کنید X دارای میانگین $\mu = 1$ و واریانس $\sigma^2 = \frac{1}{4}$ باشد. کران بالا، طبق نامساوی چیبیوف،

برای $P(|x-1| \geq 2)$ کدام عدد است؟

(۱) $\frac{1}{4}$ (۲) $\frac{1}{2}$ (۳) $\frac{1}{16}$ (۴) $\frac{1}{8}$

۷. متغیر تصادفی X دارای چگالی یکنواخت در فاصله (۱ و -۱) می‌باشد. $P(|x| > \sqrt{3 \text{Var}(x)})$ برابر است با:

(۱) صفر (۲) یک (۳) $\frac{1}{2}$ (۴) $\frac{1}{3}$

۸. جفت تصادفی (X, Y) دارای چگالی توأم $f(x, y) = (x+y+2)/8$ است، به طوری که $|x| < 1$ ، $P(X < Y)$ برابر است با:

(۱) $\frac{1}{4}$ (۲) $\frac{1}{2}$ (۳) $\frac{1}{3}$ (۴) $\frac{1}{6}$

۹. نمونه تصادفی $X_1, X_2, \dots, X_n$ را از توزیعی با میانگین ۲ و واریانس ۴ در نظر می‌گیریم. $\frac{1}{n} \sum_{i=1}^n X_i^2$

هنگامی که $n \rightarrow \infty$ به کدام در توزیع میل می‌کند؟

(۱) به یک متغیر نرمال (۲) به عدد ۸

(۳) به عدد ۲ (۴) به عدد ۴

۱۰. $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از توزیعی با واریانس متناهی σ^2 می‌باشد. $S^2 = \sum_{i=1}^n (x_i - \bar{x})^2 / n$

به عنوان برآوردکننده σ^2 به کار می‌بریم. با کدام n قدر مطلق اریبی S^2 برابر یک دهم σ^2 می‌شود؟

(۱) ۱۰۰ (۲) ۱۱ (۳) ۹ (۴) ۱۰

۱۱. دو متغیر تصادفی X_1 و X_2 ، به ترتیب با واریانس‌های ۱ و ۴، مستقل می‌باشند. $U = X_1 + X_2$ و $W = X_1 + cX_2$

برای چه مقدار c ناهمبسته می‌باشند؟

(۱) $\frac{1}{2}$ (۲) $\frac{1}{4}$ (۳) ۴ (۴) $-\frac{1}{4}$

۱۲. میان دو متغیر تصادفی X و Y ، با گشتاور دوم متناهی، رابطه $Y - bX = aX^2 + c$ برقرار است. ضریب همبستگی X و Y در چه صورت همواره برابر یک می‌شود؟

$$\begin{array}{ll} a = 0 & (۱) \\ b > 0, a = 0 & (۲) \\ b < 0, a = 0 & (۳) \\ b = 0, a = 0 & (۴) \end{array}$$

۱۳. x_1, x_2, x_3 یک نمونه تصادفی از توزیع برنولی با پارامتر P می‌باشد. $E(x_1 | x_1 + x_2 + x_3)$ برابر است با:

$$\begin{array}{llll} \bar{x} & (۱) & \frac{1}{3} \bar{x} & (۳) \\ 3\bar{x} & (۲) & \frac{1}{3} \bar{x} & (۴) \end{array}$$

۱۴. متغیر تصادفی X دارای چگالی $0 < x < 1$ و $f(x) = 2x$ می‌باشد. چارک اول کدام است؟

$$\begin{array}{llll} 1 & (۱) & 0.75 & (۴) \\ 0.5 & (۲) & 0.5 & (۳) \end{array}$$

۱۵. برای یک نمونه تصادفی ده‌تایی از توزیع $N(0, \sigma^2)$ ، $\sum_{i=1}^{10} X_i^2 = 1/6$ را مشاهده کرده‌ایم. MLE انحراف معیار کدام عدد است؟

$$\begin{array}{llll} 0.16 & (۱) & 0.4 & (۲) \\ 0.2 & (۳) & 0.4 & (۴) \end{array}$$

۱۶. برای پارامتر σ^2 (واریانس)، با استفاده از داده‌ها، یک فاصله اطمینان $(0.49, 0.25)$ با ضریب اطمینان 0.81 به‌دست آمده است. فاصله اطمینان و ضریب اطمینان با همین داده‌ها برای σ کدام است؟

$$\begin{array}{llll} (0.49, 0.25) \text{ با ضریب } 0.9 & (۱) & (0.7, 0.5) \text{ با ضریب } 0.9 & (۲) \\ (0.7, 0.5) \text{ با ضریب } 0.81 & (۳) & (0.7, 0.5) \text{ با ضریب } 0.81/2 & (۴) \end{array}$$

۱۷. یک نمونه تصادفی از توزیعی با پارامتر θ داریم، U و W و T به‌ترتیب آماره بسننده، برآوردگر گشتاوری و MLE برای θ می‌باشد. همواره داریم؟

$$\begin{array}{ll} U = T & (۲) \\ U = W & (۱) \\ U \text{ تابعی از } T & (۴) \\ W = T & (۳) \end{array}$$

۱۸. $x_1, x_2, \dots, x_n$ یک نمونه تصادفی از توزیع $N(\mu, \sigma^2)$ می‌باشد. فرض کنید $\bar{x}$ میانگین این نمونه و S^2 واریانس آن باشد. $P((\bar{x} - \mu)(S^2 - \sigma^2) > 0)$ برابر است با:

$$\begin{array}{llll} \frac{1}{2} & (۱) & \frac{1}{4} & (۲) \\ \frac{1}{8} & (۳) & 1 & (۴) \end{array}$$

۱۹. یک نمونه تصادفی $x_1, x_2, \dots, x_n$ از توزیع پواسن با میانگین λ در نظر می‌گیریم. فرض کنید S^2 واریانس ناریب این نمونه باشد. کران پایین برای واریانس S^2 برابر است با:

$$\begin{array}{llll} \lambda & (۱) & \frac{1}{\lambda} & (۲) \\ \frac{n}{\lambda} & (۳) & \frac{\lambda}{n} & (۴) \end{array}$$

۲۰. فرض کنید Z_1, Z_2, Z_3, Z_4 یک نمونه تصادفی از توزیع $N(0, 1)$ باشد. کدام یک از متغیرهای زیر توزیع F دارد؟

$$\begin{array}{ll} \left(\frac{Z_1 + Z_2}{Z_3 + Z_4}\right)^2 & (۱) \\ \left(\frac{Z_1 - Z_2}{Z_3 + Z_4}\right)^2 & (۳) \\ \left(\frac{Z_1 + Z_3}{Z_2 + Z_4}\right)^2 & (۲) \\ \left(\frac{Z_1 - Z_2}{Z_3 + Z_4}\right)^2 & (۴) \end{array}$$

۲۱. توان آزمون، برای آزمون فرض ساده H_0 در برابر فرض ساده H_1 چیست؟

(۱) احتمال رد H_0 هنگامی که H_0 نادرست باشد.

(۲) احتمال قبول H_0 هنگامی که H_1 نادرست باشد.

(۳) احتمال رد H_0 هنگامی که H_1 درست باشد.

(۴) احتمال قبول H_0 هنگامی که H_1 درست باشد.

۲۲. می‌خواهیم با مشاهده یک X ، فرض این که X دارای چگالی $4X^3$ می‌باشد، در برابر این که X دارای چگالی $3X^2$ می‌باشد را بیازماییم $0 < x < 1$ ، اگر خطای نوع اول $\frac{1}{16}$ باشد ناحیه بحرانی بر حسب X به چه صورت است؟

$$X < \frac{1}{2} \quad (1) \quad X < \frac{1}{4} \quad (2) \quad X > \frac{1}{2} \quad (3) \quad X > \frac{1}{4} \quad (4)$$

۲۳. نمونه تصادفی $x_1, x_2, \dots, x_n$ از توزیع $N(0, \sigma^2)$ را در نظر می‌گیریم. کدام آماره برای σ^2 بسنده است؟

$$\bar{x} \quad (1) \quad \bar{x}^2 \quad (2) \quad (\bar{x})^3 \quad (3) \quad e^{x^2} \quad (4)$$

۲۴. یک نمونه تصادفی از چگالی $x > 0$ و $f(x; \theta) = \frac{1}{\theta} e^{-\frac{x}{\theta}}$ در نظر می‌گیریم. برآورد کارا (efficient) برای θ کدام است؟

$$\bar{x} \quad (1) \quad \frac{1}{\bar{x}} \quad (2) \quad e^{\bar{x}} \quad (3) \quad \log \bar{x} \quad (4)$$

۲۵. یک نمونه تصادفی از توزیع $N(\mu, 1)$ در نظر می‌گیریم. کدام یک از آزمون‌های زیر UMP می‌باشد؟

$$\mu = 0 \text{ در برابر } \mu \neq 0 \quad (1) \quad \mu = 0 \text{ در برابر } \mu < \frac{1}{2} \quad (2) \quad \mu = 0 \text{ در برابر } \mu > 0 \quad (3) \quad \mu = 0 \text{ در برابر } \mu < -\frac{1}{2} \quad (4)$$

۲۶. متغیر تصادفی گسسته X دارای تابع مولد گشتاور $M(t) = \frac{1}{4}e^{-t} + \frac{1}{4}e^t + \frac{1}{4}$ می‌باشد. $\text{Var}(X)$ برابر است با:

$$\frac{1}{2} \quad (1) \quad \frac{1}{4} \quad (2) \quad 1 \quad (3) \quad \frac{1}{8} \quad (4)$$

۲۷. $x_1, x_2, \dots, x_n$ یک نمونه تصادفی از توزیع نرمال با میانگین μ می‌باشد. $P(\bar{x} < \mu, \bar{x} > x_1)$ برابر است با:

$$\frac{1}{2} \quad (1) \quad \frac{1}{4} \quad (2) \quad \frac{1}{8} \quad (3) \quad 1 \quad (4)$$

۲۸. x_1, x_2, x_3 به ترتیب دارای انحراف معیارهای σ_1 و σ_2 و σ_3 می‌باشد و دو به دو دارای ضریب همبستگی $\frac{1}{4}$ هستند.

$$\text{var}\left(\frac{x_1}{\sigma_1} + \frac{x_2}{\sigma_2} + \frac{x_3}{\sigma_3}\right) \text{ برابر است با:}$$

$$4 \quad (1) \quad 66 \quad (2) \quad 8 \quad (3) \quad 12 \quad (4)$$

روش‌های آماری

۲۹. اگر در رگرسیون خطی ساده Y روی X که تمام مفروضات آن برقرار است ضریب تعیین برابر یک باشد، کدام گزینه حتماً صحیح است؟

(۱) شیب خط رگرسیونی برابر یک است.

(۲) شیب خط رگرسیونی منفی است.

(۳) مجموع مربعات رگرسیونی با مجموع مربعات خطا برابر است.

(۴) مجموع مربعات خطا صفر است.

۳۰. در تحلیل رگرسیونی، تغییر واحد اندازه‌گیری متغیر پاسخ کدام یک از موارد زیر را تغییر نمی‌دهد؟

(۱) برآورد ثابت رگرسیونی (β_0)

(۲) ضریب تعیین (R^2)

(۳) برآورد واریانس شیب رگرسیونی ($\text{Var}(\hat{\beta}_1)$)

(۴) برآورد شیب رگرسیونی ($\hat{\beta}_1$)

۳۱. اگر احتمال بهبودی از یک بیماری P در نظر گرفته شود و در یک تحقیق فقط یک نفر از بین شش بیمار بهبودی پیدا کرده باشد، P-Value آزمون $H_0: p = 0.5$ در مقابل $H_1: p < 0.5$ برابر است با:

$$\begin{array}{llll} (۱) & \frac{6}{64} & (۲) & \frac{7}{64} \\ (۳) & \frac{12}{64} & (۴) & \frac{14}{64} \end{array}$$

۳۲. کدام یک از مدل‌های زیر قابل تبدیل به حالت خطی نیست؟

$$(۱) y_i = \frac{1}{1 + e^{\beta_0 + \beta_1 x_i + e_i}} \quad (۲) y_i = \beta_0 + \beta_1 \frac{1}{x_i^2 - 1}$$

$$(۳) y_i = \beta_0 + \beta_1 e^{\beta_2 x_i + e_i} \quad (۴) y_i = \beta_0 e^{\beta_1 x_i + e_i}$$

۳۳. برای آزمون فرضیه $\mu_1 = \mu_2 = \mu_3$ در مقابل $H_1: \mu_1 \neq \mu_2$ با توان ۹۰ درصد و خطای نوع اول ۵ درصد اگر اندازه اثر (effect size) را نصف انحراف معیار مشترک در نظر بگیریم، حداقل نمونه مورد نیاز در هر گروه برابر است

$$\text{با: } 10/5 = (Z_{0.975} + Z_{0.90})^2$$

$$(۱) 42 \quad (۲) 84 \quad (۳) 168 \quad (۴) 336$$

۳۴. از ویژگی‌های طرح بلوک تصادفی ناقص متعادل کدام است؟

(۱) تعداد بلوک‌ها و درمان‌ها (Treatments) برابر نمی‌باشد.

(۲) تعداد بلوک‌ها و درمان‌ها (Treatments) برابر می‌باشد.

(۳) کلیه درمان‌ها (Treatments) در تمام بلوک‌ها برابر و به دفعات یکسان ظاهر می‌شود.

(۴) تکرار هر درمان (Treatment) در کلیه بلوک‌ها برابر و کمتر از تعداد کل بلوک‌ها است.

۳۵. کدام یک از شاخص‌های زیر مربوط به ارزیابی برازش مدل رگرسیون چندگانه نیست؟

$$(۱) AIC \quad (۲) R_a^2 \quad (۳) PRESS \quad (۴) VIF$$

۳۶. اگر مدل کامل (Full model) $y_i = \beta_0 + \beta_1 x_i + \varepsilon_i$ و مدل کاهش یافته (reduced model) به صورت $y_i = \beta_0 + \varepsilon_i$ باشد، کدامیک از گزینه‌های زیر آماره آزمون برازش مدل کاهش یافته نسبت به مدل کامل است؟

$$F^* = \frac{SSE(F) - SSE(R)}{df_F - df_R} \div \frac{SSE(F)}{df_F} \quad (۱)$$

$$F^* = \frac{SSE(R) - SSE(F)}{df_R - df_F} \div \frac{SSE(F)}{df_F} \quad (۲)$$

$$F^* = \frac{SSE(F) - SSE(R)}{df_F - df_R} \div \frac{SSE(R)}{df_R} \quad (۳)$$

$$F^* = \frac{SSE(R) - SSE(F)}{df_R - df_F} \div \frac{SSE(R)}{df_R} \quad (۴)$$

۳۷. در یک طرح مربع لاتین 3×3 با ۲ تکرار درجه آزادی خطا چقدر است؟

- (۱) ۱۰ (۲) ۸ (۳) ۲ (۴) ۱

۳۸. اگر MSE میانگین مربعات خطا در رگرسیون باشد و توزیع خطاها نرمال نباشد، در این صورت:

(۱) $\sqrt{MSE}$ برآوردکننده نیرومند (Robust) σ (انحراف معیار خطا) نیست.

(۲) $\sqrt{MSE}$ برآوردکننده نیرومند (Robust) σ (انحراف معیار خطا) است.

(۳) $\frac{MSE}{n}$ برآوردکننده نیرومند (Robust) σ (انحراف معیار خطا) است.

(۴) MSE برآوردکننده نیرومند (Robust) σ (انحراف معیار خطا) است.

۳۹. اندازه نگرش و عملکرد ۶ فرد در مورد مصرف سیگار در مقیاس صفر تا ۱۰۰ نمره‌گذاری شده است. اگر در بین داده‌ها فقط سه زوج ناهماهنگ (Discordant) داشته باشیم و هیچ زوج هم رتبه‌ای نداشته باشیم ضریب همبستگی کندال برابر است با:

- (۱) ۰/۲ (۲) ۰/۵ (۳) ۰/۶ (۴) ۰/۸

۴۰. در مدل ANOVA با اندازه‌گیری‌های تکراری در r زمان متوالی، منظور از شرط کرویت (Sphericity) چیست؟

(۱) کوواریانس مشاهدات در دو زمان متوالی یکسان است.

(۲) واریانس مشاهدات در همه تکرارهای زمانی یکسان است.

(۳) همبستگی مشاهدات در همه تکرارهای زمانی یکسان است.

(۴) همبستگی مشاهدات در دو زمان متوالی یکسان است.

۴۱. در مدل ANOVA با اثرات تصادفی $y_{ij} = \mu_i + \varepsilon_{ij}$ که در آن μ_i ها و ε_{ij} ها به ترتیب متغیرهای تصادفی مستقل به صورت

$N(\mu, \sigma_\mu^2)$ و $N(0, \sigma_\varepsilon^2)$ و $i = 1, \dots, n, j = 1, \dots, r$ هستند. کدام گزینه سهم تغییرپذیری μ_i ها را آزمون می‌کند؟

$$\frac{\sigma_\mu^2}{\sigma_\mu^2 + \sigma_\varepsilon^2} \quad (۲) \quad \frac{\sigma_\mu^2}{\sigma_\varepsilon^2} \quad (۱)$$

$$\frac{\sigma_\varepsilon^2}{\sigma_\mu^2} \quad (۴) \quad \sigma_\mu^2 \quad (۳)$$

۴۲. در رگرسیون چندگانه X_1, X_2, X_3 و Y روی مقدار $SSR(X_3 | x_1, x_2)$ برابر با کدامیک از گزینه‌های زیر است؟

(۱) $SSR(x_1, x_2, x_3) - SSE(x_1, x_2)$

(۲) $SSR(x_1, x_2) - SSE(x_1, x_2, x_3)$

(۳) $SSR(x_1, x_2, x_3) - SSE(x_1, x_2)$

(۴) $SSR(x_1, x_2) - SSR(x_1, x_2, x_3)$

۴۳. اگر در یک مدل رگرسیون چندگانه خطی با p پارامتر، شاخص مالو (Mallow) را با C_p نشان دهیم، در کدامیک از

حالات زیر اعتبار (Validation) مدل بالاتر است؟

(۱) $C_p \approx p$ (۲) $C_p \leq p$ (۳) $C_p > p$ (۴) $C_p \geq p$

۴۴. کدامیک از گزینه‌های زیر، نتیجه قضیه گوس-مارکف در رگرسیون خطی ساده است؟

(۱) برآوردهای حداقل مربعات ضرایب مدل، نارایب، سازگار و در بین همه برآوردهای نارایب دارای کمترین واریانس هستند.

(۲) برآوردهای حداکثر درست‌نمایی ضرایب مدل، نارایب، سازگار و در بین همه برآوردهای نارایب دارای کمترین واریانس هستند.

(۳) برآوردهای حداقل مربعات ضرایب مدل، نارایب و در بین همه برآوردهای نارایب دارای کمترین واریانس هستند.

(۴) برآوردهای حداکثر درست‌نمایی ضرایب مدل، نارایب و در بین همه برآوردهای نارایب دارای کمترین واریانس هستند.

۴۵. در تبدیل باکس-کاکس (Box-Cox) برای مدل $y_i^\lambda = \beta_0 + \beta_1 x_i + \varepsilon_i$ جهت برآورد λ از کدام روش استفاده می‌شود؟

(۱) روش حداکثر درست‌نمایی (Maximum likelihood)

(۲) روش کمترین مربعات معمولی (Ordinary least square)

(۳) روش کمترین مربعات وزنی (Weighted least square)

(۴) روش ناپارامتری Lowess

۴۶. اگر در برازش مدل رگرسیون چندگانه Y روی X_1 و X_2 در یک نمونه تصادفی $n = 19$ و $\sum (y_i - \hat{y})^2 = 64$ باشد، درجه آزادی واریانس خطا برابر است با:

(۱) ۲ (۲) ۱۶ (۳) ۱۷ (۴) ۱۸

۴۷. برآورد حجم نمونه با استفاده از روش نمونه‌گیری تصادفی ساده، جهت برآورد شیوع یک بیماری در شهری بزرگ برابر

۱۰۰۰ نفر محاسبه گردیده است. اگر به‌جای نمونه‌گیری تصادفی ساده، نمونه‌گیری خوشه‌ای از خانوارها به‌عمل آوریم،

مشخص نماییم. در صورتی که متوسط بعد خانوار برابر ۴ و اثر طرح (Design effect) برابر ۳/۲ باشد، حدود چند خانوار

بایستی مورد مطالعه قرار گیرند؟

(۱) ۲۵۰ (۲) ۴۰۰ (۳) ۶۴۰ (۴) ۸۰۰

۴۸. در رگرسیون چندگانه روش حداقل مربعات وزنی (Weighted least squares) معمولاً در کدامیک از موارد زیر

استفاده می‌شود؟

(۱) در صورتی که هم خطی چندگانه وجود داشته باشد.

(۲) مقادیر پاسخ، پرت و دور افتاده باشد.

(۳) توزیع خطا نرمال نباشد.

(۴) واریانس خطا ثابت نباشد.

۴۹. فرض کنید تعداد تصادفات منجر به مرگ در طول هفته دارای توزیع پواسن باشد. اگر تعداد تصادفات در ۱۰ هفته پیاپی

برابر با ۳، ۱، ۲، ۴، ۳، ۰، ۱، ۲، ۱ باشد، احتمال مشاهده‌ی کدام مقدار در طول یک هفته بیش‌تر است؟

(۱) ۱ (۲) ۲ (۳) ۳ (۴) ۴

۵۰. در یک طرح آزمایشی کاملاً تصادفی، در صورتی که تعداد گروه‌ها ۴ و تعداد مشاهدات در هر گروه مساوی و برابر ۶ باشد، درجه‌ی آزادی خطا برابر است با:

- (۱) ۵ (۲) ۱۵ (۳) ۱۸ (۴) ۲۰

۵۱. در یک جدول توافقی با $r > 2$ سطر و $c > 2$ ستون، تعداد سطرها یا ستون‌ها را کاهش می‌دهیم. در این صورت پراکندگی توزیع با استفاده از ملاک کای-دو:

- (۱) افزایش می‌یابد. (۲) کاهش می‌یابد.

- (۳) تغییر نمی‌کند. (۴) برابر $(r-1)(c-1)$ می‌شود.

۵۲. اگر $SSE(F)$ خطای مدل رگرسیونی کامل (Full model) و $SSE(R)$ خطای مدل کاهش یافته (Reduced model) باشد، کدامیک از روابط زیر همواره برقرار است؟

- (۱) $SSE(F) < SSE(R)$ (۲) $SSE(F) \leq SSE(R)$

- (۳) $SSE(F) > SSE(R)$ (۴) $SSE(F) \geq SSE(R)$

۵۳. جهت بررسی وضعیت سلامت و بیماری در خانوارها در یک شهر بزرگ که لیست کامل خانوارها در دست نمی‌باشد، کدامیک از روش‌های نمونه‌گیری مناسب‌تر است؟

- (۱) تصادفی ساده (۲) منظم (سیستماتیک)

- (۳) طبقه‌بندی شده (۴) خوشه‌ای

۵۴. متغیر تصادفی x دارای میانگین نامعلوم θ و واریانس معلوم σ^2 می‌باشد. اگر $\hat{\theta}$ برآوردکننده‌ای اریب از θ باشد، مقدار اریبی چقدر است؟

- (۱) $\hat{\theta} - \theta$ (۲) $E(\hat{\theta}) - \theta$

- (۳) $\frac{E(\hat{\theta}) - \theta}{\sigma}$ (۴) $\frac{\hat{\theta} - \theta}{\sigma}$

۵۵. در مدل رگرسیون خطی Y روی X_1 و X_2 ، وقتی که X_2 متغیری کمی و X_1 متغیری کیفی در سه رده‌ی مختلف باشد، مشخص نمایید در ازاء مقادیر مختلف X_2 ، این مدل دارای چند شیب و چند عرض از مبدأ می‌باشد؟

- (۱) یک شیب و سه عرض از مبدأ (۲) یک شیب و دو عرض از مبدأ

- (۳) سه شیب و یک عرض از مبدأ (۴) دو شیب و دو عرض از مبدأ

۵۶. در مدل رگرسیون لجستیک $\text{Log} \frac{\pi}{1-\pi} = \beta_0 + \beta_1 X$ ، در ازای دو واحد افزایش در متغیر پیشگوی X ، نسبت بخت e برابر شده است. مقدار ضریب β_1 مدل چقدر است؟

- (۱) $\frac{1}{2}$ (۲) ۲ (۳) $\frac{e}{2}$ (۴) $2e$

۵۷. واریانس برآورد میانگین در نمونه‌گیری تصادفی ساده با جایگذاری در مقایسه با نمونه‌گیری تصادفی ساده بدون جایگذاری:

- (۱) کوچک‌تر یا مساوی است. (۲) مساوی است.

- (۳) کوچک‌تر است. (۴) بزرگ‌تر است.

۵۸. در مدل رگرسیون خطی که از مبدأ عبور می‌کند، اگر n و 2 و $i=1, \dots, n$ باشد، $e_i = \hat{y}_i - y_i$ ، برآورد ناریب واریانس خطای مدل کدام است؟

- (۱) $\sum_{i=1}^n e_i^2 / (n-2)$ (۲) $\sum_{i=1}^n e_i^2 / (n-1)$

- (۳) $\sum_{i=1}^n e_i^2 / (n-1)$ (۴) $\sum_{i=1}^n e_i^2 / (n-1)$

۵۹. در مدل رگرسیون $y_i = \beta_0 + \beta_1 x_{1i} + \beta_2 x_{2i} + \beta_3 x_{3i} + \epsilon_i, i = 1, 2, \dots, n$ برای انجام آزمون $H_0: \beta_1 = \beta_2 = \beta_3$ در مقابل $H_1: \beta_1 \neq \beta_2 \neq \beta_3$ ، آماره آزمون F دارای چه درجه‌ی آزادی می‌باشد؟

$$F(2, n-3) \quad (2)$$

$$F(1, n-3) \quad (1)$$

$$F(1, n-4) \quad (4)$$

$$F(2, n-4) \quad (3)$$

۶۰. در یک مدل تحلیل واریانس سه طرفه اگر عامل‌ها به‌ترتیب دارای ۴، ۳ و ۲ رده (Category) و حجم نمونه برای هر یک از حالات ترکیبی رده‌های این سه عامل که به‌طور تصادفی انتخاب شده‌اند ۱۱ نمونه باشد، درجه آزادی برای مجموع مجزورات باقی‌مانده برابر است با:

$$240 \quad (4)$$

$$120 \quad (3)$$

$$263 \quad (2)$$

$$60 \quad (1)$$

۶۱. در ارزیابی تأثیر یک روش گروه درمانی بر کیفیت زندگی سالمندان (y) از طرح قبل و بعد روی نمونه‌ای تصادفی از ۲۰ سالمند استفاده شده است. کدام مدل آماری زیر می‌تواند مناسب باشد؟

$$(i = 1, 2, j = 1, \dots, 20)$$

$$(2) \text{ طرح دو عاملی}$$

$$(1) \text{ طرح کاملاً تصادفی}$$

$$(4) \text{ طرح آشیانی (nested)}$$

$$(3) \text{ طرح بلوک‌های تصادفی}$$

۶۲. محقق می‌خواهد تأثیر نوعی بیماری (D) با ۴ رده (Category) و جنس (S) را بر کیفیت زندگی سالمندان بررسی نماید. اگر حجم نمونه برای هر یک از حالاتی که با ترکیب رده‌های ۲ متغیر فوق تشکیل می‌شود ۸ سالمند باشد، آنگاه امید ریاضی اثر متقابل بیماری و جنس عبارت است از:

$$\sigma_e^2 \text{ و } \sigma_{DS}^2 \text{ به ترتیب واریانس خطا و اثرات متقابل اند.}$$

$$\sigma_e^2 \quad (2)$$

$$\sigma_e^2 + 8\sigma_{DS}^2 \quad (1)$$

$$\sigma_e^2 + 16\sigma_{DS}^2 \quad (4)$$

$$\sigma_e^2 + 32\sigma_{DS}^2 \quad (3)$$

۶۳. در یک تحلیل کوواریانس مدل برآزش شده به‌صورت $y_{ij} = \alpha_i + \beta_i x + \epsilon_{ij}$ می‌باشد. در این صورت برآورد $E(Y | X)$ به‌صورت زیر خواهد بود:

$$(\bar{Y}_{..}, \bar{X}_{..}) \text{ به ترتیب میانگین کل } Y \text{ و } X \text{ و } \bar{Y}_i \text{ و } \bar{X}_i \text{ میانگین } Y \text{ و } X \text{ ها در گروه } i \text{ ام و } b_i \text{ برآورد } \beta_i \text{ است.}$$

$$\bar{Y}_i - b_i \bar{X}_i + b_i X \quad (2)$$

$$\bar{Y}_{..} - b_i \bar{X}_{..} + b_i X \quad (1)$$

$$\bar{Y}_{..} - \bar{Y}_i - bX \quad (4)$$

$$\bar{Y}_i - b\bar{X}_i + bX \quad (3)$$

۶۴. در یک تحلیل رگرسیون چندگانه Y بر حسب X_1 و X_2 ، آزمون فرض ضریب همبستگی چندگانه $H_0: \rho_{y,12}^2 = 0$ هم ارز کدام آزمون فرض است؟

$$H_0: \beta_1 = \beta_2 = 0 \quad (2)$$

$$H_0: \beta_1 = \beta_2 = 0 \quad (1)$$

$$H_0: \beta_1 = \beta_2 = 1 \quad (4)$$

$$H_0: \beta_1 = \beta_2 = 1 \quad (3)$$

۶۵. اگر سرعت جذب در بدن برای دو نوع داروی مسکن داری توزیع نرمال با واریانس‌های یکسان و معلوم باشد. آماره لازم برای آزمون فرض $H_0: \mu_1 = \mu_2$ کدام است؟

$$(\mu_1 \text{ و } \mu_2 \text{ میانگین سرعت جذب دو دارو است.})$$

$$t = \frac{n(\bar{x}_1 - 2\bar{x}_2)}{10S} \quad (2)$$

$$t = \frac{\sqrt{n}(\bar{x}_1 - 2\bar{x}_2)}{S\sqrt{5}} \quad (1)$$

$$Z = \frac{n(\bar{x}_1 - 2\bar{x}_2)}{5\sigma} \quad (4)$$

$$z = \frac{\sqrt{n}(\bar{x}_1 - 2\bar{x}_2)}{\sigma\sqrt{10}} \quad (3)$$

۶۶ در مطالعه اثر ۴ روش ماساژدرمانی بر فشار خون افراد سه عامل مخدوش وجود دارد و از مدل $y_{ijkm} = \mu + T_i + \beta_j + \gamma_k + w_m + \varepsilon_{ijkm}, i, j, k, m = 1, \dots, 4$ استفاده شده است. طرح مطالعه

کدام مورد زیر بوده است؟

(۱) طرح مربع یونانی-لاتین

(۲) تحلیل دو عاملی با دو اثر متقابل

(۳) تحلیل سه عاملی با یک اثر متقابل

(۴) تحلیل چهار عاملی با یک اثر متقابل

۶۷ در معادله رگرسیون $E(y | x) = \alpha + 2x_1 - 3x_2 + x_3$ ، میانگین نمونه‌ای متغیرهای X_i برای $i = 1, 2, 3$ در

رابطه $\bar{x}_i = 3(1+i)$ صدق می‌کنند. اگر $\bar{y} = \sum_{i=1}^3 \bar{x}_i$ باشد، برآورد α برابر است با:

(۱) -۳۰ (۲) ۲۴ (۳) -۲۴ (۴) ۳۰

۶۸ کدام یک از گزینه‌های زیر فرض معنی‌دار بودن ارتباط متغیر X_1 را با متغیر پاسخ در مدل زیر آزمون می‌کند؟

$$E(Y | X) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \beta_4 X_1 X_2$$

$$H_0: \beta_4 = 0 \quad (۲)$$

$$H_0: \beta_1 = \beta_4 \quad (۱)$$

$$H_0: \beta_1 = 0 \quad (۴)$$

$$H_0: \beta_1 = \beta_4 = 0 \quad (۳)$$

۶۹ در مقایسه میانگین نمرات کنکور چهار رشته دوره دکتری مدل کوواریانس $\varepsilon_{ij} + \gamma$ (بهره هوش) + γ (سن)

$y_{ij} = \alpha_i + \beta$ با فرض $\varepsilon_{ij} \sim N(0, \sigma_e^2)$ ، $i = 1, \dots, 4, j = 1, \dots, n_i$ برازش شده است. در این صورت تفاضل

میانگین گروه اول و سوم برابر است با:

$$\alpha_1 - \alpha_3 + \beta \quad (۲)$$

$$\alpha_1 + \alpha_3 + \gamma \quad (۱)$$

$$\alpha_1 - \alpha_3 \quad (۴)$$

$$\alpha_1 - \alpha_3 + \gamma \quad (۳)$$

آمار ریاضی و احتمالات

۱. گزینه (۳)

$$E(T) = E\left(\sum_{i=1}^n (X_i - \bar{X})^2\right) = E\left(\sum_{i=1}^n (X_i^2 + \bar{X}^2 - 2X_i\bar{X})\right) \quad \text{روش اول:}$$

$$= E\left(\sum_{i=1}^n X_i^2 + n\bar{X}^2 - 2\bar{X}\sum_{i=1}^n X_i\right) = E\left(\underbrace{\sum_{i=1}^n X_i^2}_{(1)} + \underbrace{E(n\bar{X}^2) - E(2\bar{X}\sum_{i=1}^n X_i)}_{(2)}\right)$$

$$(1) \left\{ \begin{array}{l} E\left(\sum_{i=1}^n X_i^2\right) = \sum_{i=1}^n E(X_i^2) = \sum_{i=1}^n (\text{Var}(X_i) + E(X_i)^2) = n\sigma^2 + n\bar{X}^2 \\ \text{Var}(X_i) = E(X_i^2) - E(X_i)^2 = \sigma^2 \\ E(X_i) = \bar{X} \end{array} \right.$$

$$(2) \left\{ \begin{array}{l} E(n\bar{X}^2) - E(2\bar{X}\sum_{i=1}^n X_i) = nE(\bar{X}^2) - 2nE(\bar{X})E(\bar{X}) = -nE(\bar{X}^2) \\ = -n(\text{Var}(\bar{X}) + E(\bar{X})^2) = -n\left(\frac{\sigma^2}{n} + \bar{X}^2\right) \\ = -\sigma^2 - n\bar{X}^2 \end{array} \right.$$

$$(1) \text{ و } (2) \Rightarrow E(T) = n\sigma^2 + n\bar{X}^2 - \sigma^2 - n\bar{X}^2 = n\sigma^2 - \sigma^2 = (n-1)\sigma^2$$

روش دوم:

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$$

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N (X_i - \mu)^2$$

$$E(T) = E\left(\sum_{i=1}^n (X_i - \bar{X})^2\right) = E((n-1)S^2) = (n-1)E(S^2) = (n-1)\sigma^2$$

نکات:

نکته ۱: نمونه تصادفی $X_1, X_2, \dots, X_n$ را از جامعه‌ای در نظر می‌گیریم و فرض می‌کنیم $g(X)$ تابعی باشد که برای آن $E(g(X))$ و $\text{Var}(g(X))$ وجود دارند. آن‌گاه:

$$E\left(\sum_{i=1}^n g(X_i)\right) = nE(g(X_1)) \quad , \quad \text{Var}\left(\sum_{i=1}^n g(X_i)\right) = n\text{Var}(g(X_1))$$

نکته ۲: فرض می‌کنیم $X_1, X_2, \dots, X_n$ ، نمونه‌ای تصادفی از جامعه‌ای با میانگین μ و واریانس متناهی σ^2 باشد، آن‌گاه داریم:

$$1) \quad E\bar{X} = \mu \quad , \quad 2) \quad \text{Var}\bar{X} = \frac{\sigma^2}{n} \quad , \quad 3) \quad ES^2 = \sigma^2$$

۲. گزینه (۴)

$$E(X_i) = E\left(\frac{\overbrace{X_i}^{\text{آماره فرعی}}}{\underbrace{\sum_{i=1}^n X_i}_{\text{آماره کامل}}}\right) \xrightarrow{\text{بنابر قضیه باسو}} E\left(\frac{X_i}{\sum_{i=1}^n X_i}\right) \times E\left(\sum_{i=1}^n X_i\right)$$

$$\Rightarrow E\left(\frac{X_i}{\sum_{i=1}^n X_i}\right) = \frac{E(X_i)}{E\left(\sum_{i=1}^n X_i\right)} = \frac{\mu}{n\mu} = \frac{1}{n} \Rightarrow E\left(\frac{X_1}{\bar{X}}\right) = \frac{E(X_1)}{E\left(\frac{1}{n} \sum_{i=1}^n X_i\right)} = \frac{\mu}{\frac{1}{n}(n\mu)} = 1$$

نکات: قضیه باسو: فرض کنید X دارای تابع چگالی احتمالی $f_\theta(X)$ ، $\theta \in \Theta$ باشد. اگر T آماره کامل کراندار و بسنده برای پارامتر θ باشد و U هر آماره‌ای باشد که توزیع آن بستگی به θ ندارد، آن‌گاه T و U از هم مستقل هستند. آماره کامل در خانواده نمایی: فرض کنید مشاهدات $X_1, X_2, \dots, X_n$ ، مستقل و هم توزیع از یک خانواده نمایی با تابع چگالی احتمال (pdf) یا تابع جرم احتمال (pmf) به فرم زیر هستند:

$$f(x|\theta) = h(x)c(\theta)\exp\left(\sum_{j=1}^K w(\theta_j)t_j(x)\right)$$

که $\theta = (\theta_1, \dots, \theta_K)$

آن گاه آماره $T(x) = \left(\sum_{i=1}^n t_1(X_i), \dots, \sum_{i=1}^n t_k(X_i) \right)$ ، یک آماره کامل در فضای پارامتر Θ است.

آماره فرعی خانواده پارامترهای مقیاس: فرض کنید $X_1, \dots, X_n$ مشاهدات مستقل و هم توزیع از یک خانواده پارامتر

مقیاس با تابع توزیع تجمعی $F\left(\frac{x}{\sigma}\right)$ ، $\sigma > 0$ باشند. آن گاه هر آماره‌ای که تنها از طریق $n-1$ مقدار $\frac{X_1}{X_n}, \dots, \frac{X_{n-1}}{X_n}$

به نمونه وابسته باشد، یک آماره فرعی است.

برای مثال $1 + \frac{X_1}{X_n} + \dots + \frac{X_{n-1}}{X_n} = \frac{X_1 + \dots + X_n}{X_n}$ ، یک آماره فرعی است. برای نشان دادن این موضوع داریم:

$$\begin{aligned} X_i = \sigma Z_i &\Rightarrow F(y_1, \dots, y_{n-1} | \sigma) = P\left(\frac{X_1}{X_n} \leq y_1, \dots, \frac{X_{n-1}}{X_n} \leq y_{n-1}\right) \\ &= P\left(\frac{\sigma Z_1}{\sigma Z_n} \leq y_1, \dots, \frac{\sigma Z_{n-1}}{\sigma Z_n} \leq y_{n-1}\right) = P\left(\frac{Z_1}{Z_n} \leq y_1, \dots, \frac{Z_{n-1}}{Z_n} \leq y_{n-1}\right) \end{aligned}$$

۳. گزینه (۳)

$$P \text{ (از ۴ مهره انتخابی، ۳ تا زوج باشند)} = \frac{\binom{3}{3} \binom{3}{1}}{\binom{6}{4}} = \frac{1 \times 3}{15} = \frac{1}{5} = \frac{2}{10}$$

نکته: فرض کنید $S = \{s_1, \dots, s_N\}$ یک فضای نمونه متناهی باشد و $P(\{S_i\})$ ، آن گاه:

$$P(A \text{ رخداد پیشامد}) = \sum_{S_i \in A} P(\{S_i\}) = \sum_{S_i \in A} \frac{1}{N} = \frac{\text{تعداد عناصر موجود در } A}{\text{تعداد عناصر موجود در } S}$$

حال برای محاسبه تعداد عناصر موجود در A و S با استفاده از روش‌های شمارش داریم:

تعداد جایگشت‌های ممکن به اندازه γ از n شیء		
با جایگذاری	بدون جایگذاری	
n^γ	$\frac{n!}{(n-\gamma)!}$	دارای اهمیت ترتیب
$\binom{n+\gamma-1}{\gamma}$	$\binom{n}{\gamma}$	بدون اهمیت ترتیب

با توجه به این که در صورت سوال، نمونه‌ها بدون جایگذاری انتخاب می‌شوند و ترتیب مهم نمی‌باشد، هم برای صورت کسر و

هم برای مخرج کسر از رابطه $\binom{n}{\gamma}$ استفاده می‌کنیم.

۴. گزینه (۱)

$$F_X(x) = P(X \leq x) = \int_{-\infty}^x f_X(t) dt = \int_{-\infty}^x \frac{1}{\pi(1+t^2)} dt = \frac{1}{\pi} (\arctan t) \Big|_{-\infty}^x \\ = \frac{1}{\pi} \left(\arctan x + \frac{\pi}{2} \right) = \frac{1}{\pi} \arctan x + \frac{1}{2}$$

۵. گزینه (۴)

$$P(X+Y=0) = P(X=0, Y=0) \rightarrow \text{چون دامنه توزیع پواسن شامل اعداد غیرمنفی است} \\ = P(X=0)P(Y=0) \rightarrow Y, X \text{ با توجه به استقلال بین} \\ = \frac{e^{-\lambda_1} \lambda_1^0}{0!} \times \frac{e^{-\lambda_2} \lambda_2^0}{0!} = e^{-\lambda_1} e^{-\lambda_2} = e^{-(\lambda_1+\lambda_2)} = e^{-(2+3)} = e^{-5}$$

نکته: متغیر تصادفی X دارای توزیع پواسن با پارامتر λ ($X \sim P(\lambda)$) است، اگر دارای تابع احتمال زیر باشد:

$$P(X=x|\lambda) = f_\lambda(x) = \frac{e^{-\lambda} \lambda^x}{x!}, x=0,1,\dots; \lambda>0$$

$$E_\lambda(x) = \lambda$$

$$\text{Var}_\lambda(x) = \lambda$$

$$M_X(t) = \exp[\lambda(e^t - 1)], t \in \mathbb{R}$$

در توزیع پواسن با پارامتر λ داریم:

۶. گزینه (۳)

$$P(|X-\mu| \geq K\sigma) \leq \frac{1}{K^2} \text{ حالت خاصی از نامساوی چبیف}$$

$$P(|X-\mu| \geq 2) \leq \frac{1}{4} = \frac{1}{16} \\ \downarrow \\ K\sigma=2 \Rightarrow K\left(\frac{1}{2}\right)=2 \Rightarrow K=4$$

نکات: نامساوی‌های از نوع چبیف:

نامساوی مارکف: اگر X یک متغیر تصادفی و $h(0)$ تابعی غیرمنفی باشد که $E[h(X)] < \infty$ ، آن‌گاه برای هر $a > 0$:

$$P(h(X) \geq a) \leq \frac{E(h(X))}{a}$$

نامساوی چبیف: با انتخاب $h(x) = |x|$ ، در نامساوی مارکف و با توجه به: $P(|X| \geq a) = P(X^2 \geq a^2)$ داریم:

$$P(|X| \geq a) \leq \frac{E(X^2)}{a^2}$$

$$P(|Z| \geq t) \leq \sqrt{\frac{2}{\pi}} \frac{e^{-\frac{t^2}{2}}}{t}, t > 0$$

نکته: اگر Z دارای توزیع نرمال استاندارد باشد، آن‌گاه:

نکته: فرض کنید X یک متغیر تصادفی با تابع مولد گشتاور $-h < t < h$ ، $M_X(t)$ باشد. آن‌گاه:

$$P(X \geq a) \leq e^{-at} M_X(t), 0 < t < h$$

$$P(X \leq a) \leq e^{-at} M_X(t), -h < t < 0$$

۷. گزینه (۱)

$$X \sim U(-1, 1) \Rightarrow \text{Var}(x) = \frac{(1 - (-1))^2}{12} = \frac{4}{12} = \frac{1}{3}$$

$$\Rightarrow P(|X| > \sqrt{r \text{Var}(x)}) = P(|X| > \sqrt{r \times \frac{1}{3}}) = P(|X| > 1) = P(X > 1) + P(X < -1)$$

با توجه به این که مقادیر X برای $X > 1$ و $X < -1$ تعریف نشده است لذا داریم:

$$P(|X| > \sqrt{r \text{Var}(x)}) = 0$$

متغیر تصادفی X دارای توزیع یکنواخت در فاصله (θ_1, θ_2) است $(X \sim U(\theta_1, \theta_2))$ ، اگر دارای تابع چگالی احتمال زیر باشد:

$$f(x | \theta_1, \theta_2) = \begin{cases} \frac{1}{\theta_2 - \theta_1}, & \theta_1 \leq x \leq \theta_2, -\infty < \theta_1 < \theta_2 < \infty \\ 0, & \text{سایر نقاط} \end{cases}$$

اگر $X \sim U(\theta_1, \theta_2)$ آن گاه:

$$E_{\theta_1, \theta_2}(x) = \frac{\theta_1 + \theta_2}{2}$$

$$\text{Var}_{\theta_1, \theta_2}(x) = \frac{(\theta_2 - \theta_1)^2}{12}$$

$$M_X(t) = \begin{cases} \frac{e^{\theta_2 t} - e^{\theta_1 t}}{\theta_2 - \theta_1} \times t, & t \neq 0 \\ 1, & t = 0 \end{cases}$$

۸. گزینه (۲)

تابع $f(x, y)$ تابع چگالی توأم بردار تصادفی دو متغیره پیوسته (X, Y) می باشد، اگر برای هر $A \subset R^2$ داشته باشیم:

$$P((X, Y) \in A) = \int_A \int f(x, y) dx dy$$

$$\begin{cases} |X| < 1 \Rightarrow -1 < X < 1 \\ |Y| < 1 \Rightarrow -1 < Y < 1 \end{cases}$$

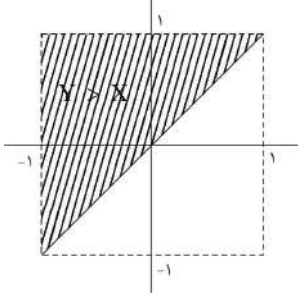
$$P(x < y) = P(-1 < x < y < 1) = \int_{-1}^1 \int_{-1}^y \frac{x+y+2}{8} dx dy$$

$$= \frac{1}{8} \int_{-1}^1 \left(\frac{1}{2} x^2 + xy + 2x \Big|_{-1}^y \right) dy = \frac{1}{8} \int_{-1}^1 \frac{y^2}{2} + 2y + \frac{y}{2}$$

$$= \frac{1}{8} \left(\frac{1}{2} y^2 + 2y + \frac{y}{2} \Big|_{-1}^1 \right) = \frac{4}{8} = \frac{1}{2}$$

روش دیگر حل مسئله به صورت زیر می باشد:

با توجه به این که دامنه تغییرات X و Y بین -1 و 1 می باشد، بنابراین مساحت کل برابر با یک است، لذا داریم:



$$P(X < Y) = \frac{1}{2}$$

نکات:

اگر $f(x, y)$ تابع چگالی احتمال توأم بردار تصادفی دو متغیره پیوسته (X, Y) باشد، داریم:

$$E(g(x, y)) = \iint g(x, y) f(x, y) dx dy$$

$$f_X(x) = \int f(x, y) dy$$

$$f_Y(y) = \int f(x, y) dx$$

۹. گزینه (۲)

با استفاده از قضیه ارگودیک که تعمیمی از قانون اعداد بزرگ می باشد، وقتی $n \rightarrow \infty$ داریم:

$$\frac{1}{n} \sum_{i=1}^n h(X_i) = E(h(x))$$

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n X_i^2 = E(X^2) = \text{Var}(X) + E(X)^2 = 4 + 2^2 = 8$$

بنابراین:

نکات: مفاهیم همگرایی

همگرایی در احتمال: دنباله متغیرهای تصادفی $X_1, \dots, X_n$ در احتمال به متغیر تصادفی X همگراست، اگر برای

$$\lim_{n \rightarrow \infty} P(|X_n - X| \geq \varepsilon) = 0 \quad \text{یا} \quad \lim_{n \rightarrow \infty} P(|X_n - X| < \varepsilon) = 1 \quad \text{هر } \varepsilon > 0$$

کوچک باشد و با افزایش n بسیار کم تر شود. به عبارت دیگر، احتمال این که X_n از X را نشان می دهد. اگر $\{X_n\}$ به X همگرا شود، $P(|X_n - X| > \varepsilon)$ باید کوچک باشد و با افزایش n بسیار کم تر شود. به عبارت دیگر، احتمال این که X_n از X فاصله داشته باشد با افزایش n باید

$$\lim_{n \rightarrow \infty} P(|X_n - X| > \varepsilon) = 0 \quad \text{به صفر میل نماید. یعنی داریم:}$$

قانون ضعیف اعداد بزرگ: فرض کنید $X_1, X_2, \dots$ متغیرهای تصادفی مستقل و هم توزیع با $EX_i = \mu$

$$\text{و } \text{Var } X_i = \sigma^2 < \infty \text{ باشند. میانگین دنباله را به صورت } \bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i \text{ تعریف می کنیم. آن گاه برای هر } \varepsilon > 0:$$

$$\lim_{n \rightarrow \infty} P(|\bar{X}_n - \mu| < \varepsilon) = 1$$

یعنی $\bar{X}_n$ در احتمال به μ همگراست.

نکته: فرض کنید $X_1, X_2, \dots$ در احتمال به متغیر تصادفی X همگرا باشد و h را یک تابع پیوسته در نظر بگیریم. آن گاه $h(x_1), h(x_2), \dots$ در احتمال به $h(x)$ همگرا است.

همگرایی قریب به یقین: نوعی از همگرایی که نسبت به همگرایی در احتمال قوی تر است. دنباله متغیرهای تصادفی $X_1, X_2, \dots$ همگرایی قریب به یقین به متغیر تصادفی X دارند، اگر برای هر $\varepsilon > 0$:

$$P\left(\lim_{n \rightarrow \infty} |X_n - X| < \varepsilon\right) = 1$$

اگر وقوع پیشامدی قریب به یقین باشد، در این صورت همواره رخ می دهد، اما پیشامدهای دیگری هم از لحاظ نظری در فضای احتمال مورد بررسی امکان وقوع دارند. اگر چه با افزایش فضای نمونه، احتمال رخ دادن هر پیشامد دیگری به طور مجانبی به صفر می گراید. به طور مثال، فضای نمونه S با بازه بسته $[0, 1]$ با توزیع احتمال یکنواخت را در نظر می گیریم.

متغیر تصادفی $X(S) = S, X_n(S) = S + S^n$ را تعریف می کنیم. برای هر $S \in [0, 1]$ ، اگر $n \rightarrow \infty$ آن گاه $S^n \rightarrow 0$. بنابراین $X_n(S) \rightarrow S = X(S)$. با این حال $X_n(1) = 2$ (برای هر n)، بنابراین $X_n(1)$ به عدد $X(1) = 1$ همگرا نمی شود. اما از آن جایی که همگرایی در بازه $[0, 1]$ اتفاق می افتد و $P([0, 1]) = 1$ ، X_n تقریباً همه جا (قریب به یقین) به عدد X همگرا است.

قانون قوی اعداد بزرگ: فرض کنید $X_1, X_2, \dots$ متغیرهای تصادفی مستقل و هم توزیع

با $EX_i = \mu, \text{Var } X_i = \sigma^2 < \infty$ باشند و میانگین دنباله به صورت $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$ تعریف شود. آن گاه برای

$$P\left(\lim_{n \rightarrow \infty} |\bar{X}_n - \mu| < \varepsilon\right) = 1 \quad \text{هر } \varepsilon > 0:$$

یعنی $\bar{X}_n$ تقریباً همه جا (قریب به یقین) به μ همگراست.

همگرایی در توزیع: دنباله متغیرهای تصادفی $X_1, X_2, \dots$ در توزیع به متغیر تصادفی X همگرا است

اگر $\lim_{n \rightarrow \infty} F_{X_n}(x) = F_X(x)$ (در همه نقاط x که $F_X(x)$ پیوسته است)

نکته: همگرایی در احتمال، همگرایی در توزیع را نتیجه می دهد.

قضیه حد مرکزی: فرض کنید $X_1, X_2, \dots$ دنباله ای از متغیرهای تصادفی مستقل و هم توزیع باشند که تابع مولد گشتاور

آن ها در همسایگی صفر وجود دارد. هم چنین فرض کنید $EX_i = \mu$ و $\text{Var } X_i = \sigma^2 < \infty$ ، $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$

و $G_n(x)$ بیانگر تابع توزیع تجمعی $\frac{\sqrt{n}(\bar{X}_n - \mu)}{\sigma}$ باشد. آن گاه برای هر $X, -\infty < X < \infty$ داریم:

$$\lim_{n \rightarrow \infty} G_n(x) = \int_{-\infty}^x \frac{1}{\sqrt{2\pi}} e^{-y^2/2} dy;$$

یعنی $\frac{\sqrt{n}(\bar{X}_n - \mu)}{\sigma}$ ، به طور تقریبی دارای توزیع نرمال استاندارد است. به عبارت دیگر، قضیه حد مرکزی بیان می کند که

با فرض شرایط خاص، میانگین تعدادی متغیر تصادفی مستقل، به طور تقریبی دارای توزیع نرمال خواهد بود.

قضیه اسلاتسکی: اگر X_n در توزیع به X و Y_n در احتمال به عدد ثابت a همگرا باشند، آن گاه:

الف- $Y_n X_n$ در توزیع به aX همگرا است، ب- $X_n + Y_n$ در توزیع به $X + a$ همگرا است.

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

$$\left| E(S)^2 - \sigma^2 \right| = \frac{1}{n} \sigma^2 \Rightarrow \left| E \left(\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n} \right) - \sigma^2 \right| = \frac{1}{n} \sigma^2 *$$

$$\begin{aligned} E(S^2) &= E \left(\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \right) = \sigma^2 \Rightarrow E \left(\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \right) \\ &= E \left(\frac{n-1}{n} \times \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \right) = E \left(\frac{n-1}{n} S^2 \right) = \frac{n-1}{n} \sigma^2 \end{aligned}$$

$$\begin{aligned} * \Rightarrow \left| E \left(\frac{n-1}{n} S^2 \right) - \sigma^2 \right| &= \frac{1}{n} \sigma^2 \Rightarrow \left| \frac{n-1}{n} \sigma^2 - \sigma^2 \right| = \frac{1}{n} \sigma^2 \\ \Rightarrow \left| \frac{\sigma^2}{n} \right| &= \frac{\sigma^2}{n} \Rightarrow n = 1 \end{aligned}$$

نکات: میانگین مربعات خطای (MSE) برآوردگر ω از پارامتر θ ، تابعی از θ است که به صورت $E_{\theta}(\omega - \theta)^2$ تعریف می‌شود. به عبارت دیگر:

$$\text{MSE} = E_{\theta}(\omega - \theta)^2 = \text{Var}_{\theta} \omega + (E_{\theta} \omega - \theta)^2 = \text{Var}_{\theta} \omega + (\text{Bias}_{\theta} \omega)^2$$

میزان اریبی برآوردگر ω از پارامتر θ به صورت $E_{\theta} \omega - \theta$ تعریف می‌شود. برآوردگری که میزان اریبی آن برابر صفر شود،

یک برآوردگر نارایب نامیده می‌شود، یعنی $E_{\theta} \omega = \theta$

بنابراین شاخص MSE شامل دو مؤلفه می‌باشد که یکی میزان تغییرپذیری برآوردگر (دقت) و دیگری میزان اریبی آن (صحت) را می‌سنجد.

$$\text{Cov}(U, \omega) = \text{Cov}(X_1 + X_2, X_1 + CX_2)$$

$$= \text{Cov}(X_1, X_1) + C \text{Cov}(X_1, X_2) + \text{Cov}(X_1, X_2) + C \text{Cov}(X_2, X_2)$$

$$= \text{Var}(X_1) + (C+1) \text{Cov}(X_1, X_2) + C \text{Var}(X_2)$$

برای این که U و ω ناهمبسته باید $\text{Cov}(U, \omega) = 0$ باشد. هم‌چنین با توجه به صورت سوال، X_1 و X_2 مستقل هستند

$$\text{Var}(X_1) + C \text{Var}(X_2) = 0 \Rightarrow 1 + 4C = 0 \Rightarrow C = -\frac{1}{4} \quad \text{پس } \text{Cov}(X_1, X_2) = 0 \text{ است. بنابراین:}$$

۱۲. گزینه (۲)

با استفاده از تعریف ضریب همبستگی بین دو متغیر X و Y که به صورت $\rho_{XY} = \frac{\text{Cov}(X, Y)}{\sigma_X \sigma_Y}$ می باشد داریم:

$$\begin{aligned}\rho_{xy} &= \frac{\text{Cov}(X, aX^r + bX + C)}{\sqrt{\text{Var}(X)} \sqrt{\text{Var}(aX^r + bX + C)}} \\ &= \frac{\text{Cov}(X, aX^r) + \text{Cov}(X, bX) + \text{Cov}(X, C)}{\sqrt{\text{Var } X} \sqrt{\text{Var}(aX^r) + \text{Var}(bX) + r \text{Cov}(aX^r, bX)}} \\ &= \frac{a \text{Cov}(X, X^r) + b \text{Var}(X)}{\sqrt{\text{Var } X} \sqrt{a^r \text{var}(X^r) + b^r \text{Var}(X) + rab \text{Cov}(X^r, X)}}\end{aligned}$$

تنها گزینه ای که به جواب منتهی می شود، گزینه ۲ است.

۱۳. گزینه (۱)

فرض کنید $S_n = X_1 + X_2 + \dots + X_n$ داریم:

$$\begin{aligned}E(X_i | S_n) &= E(X_1 | S_n) = \dots = E(X_n | S_n) \\ &\Rightarrow E(X_1 | S_n) + \dots + E(X_n | S_n) = E(S_n | S_n) = S_n \\ &\Rightarrow E(X_i | S_n) = \frac{S_n}{n}, i = 1, \dots, n \\ &\Rightarrow E(X_1 | S_r) = \frac{S_r}{r} = \frac{X_1 + X_2 + X_3}{3} = \bar{X}\end{aligned}$$

۱۴. گزینه (۳)

با استفاده از تعریف تابع توزیع که به صورت $F_X(x) = P(X \leq x) = \int_{-\infty}^x f_X(t) dt$ می باشد، میانه توزیع برابر با

$$P(X \leq m) = F_X(m) = \int_{-\infty}^m f_X(t) dt = \frac{1}{4} \quad \text{عدد است که در آن مقدار تابع توزیع برابر با } \frac{1}{4} \text{ باشد یعنی:}$$

با توجه به این که میانه برابر با چارک دوم است، برای به دست آوردن چارک اول کافیست به جای عدد $\frac{1}{4}$ ، عدد $\frac{1}{4}$ را قرار

دهیم. یعنی جایی از توزیع که مقدار توزیع تجمعی تا آن نقطه برابر با $\frac{1}{4}$ شده است:

$$F_X(Q_1) = \int_{-\infty}^{Q_1} f_X(t) dt = \frac{1}{4} \xrightarrow{\text{با توجه به دامنه } X} \int_0^{Q_1} t dt = \frac{1}{4}$$

$$\Rightarrow t^2 \Big|_0^{Q_1} = \frac{1}{4} \Rightarrow Q_1^2 = \frac{1}{4} \Rightarrow Q_1 = \pm \frac{1}{2}$$

با توجه به دامنه X ، $Q_1 = \frac{1}{2}$ درست است.

$$\begin{aligned}
 X_i &\sim N(0, \sigma^2) \Rightarrow f(x | \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2\sigma^2}(X_i^2)} \\
 \Rightarrow L(\sigma^2 | X_1, \dots, X_n) &= (2\pi\sigma^2)^{-\frac{n}{2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=1}^n X_i^2\right) \\
 \Rightarrow \text{Log } L(\sigma^2 | \underline{X}) &= -n \ln \sigma - \frac{1}{2\sigma^2} \sum_{i=1}^n X_i^2 \\
 \Rightarrow \frac{\partial}{\partial \sigma^2} \text{Log } L &= \frac{-n}{\sigma} + \frac{1}{\sigma^3} \sum_{i=1}^n X_i^2 = 0 \\
 \Rightarrow \sigma^2 &= \sum_{i=1}^n X_i^2 / n = \frac{1/6}{1} = 0.16 \Rightarrow \sigma = 0.4
 \end{aligned}$$

نکات:

روش ML، روشی برای به دست آوردن برآوردگر ماکسیمم درست‌نمایی می‌باشد که به اختصار MLE نامیده می‌شود که براساس تابع درست‌نمایی است.

فرض کنیم $X = (X_1, \dots, X_n)$ ، بردار n متغیر تصادفی با تابع چگالی احتمال توأم $f_{\underline{\theta}}(x)$ که $\underline{\theta} \in \Theta \subseteq \mathbb{R}^K$ باشد. برای هر مقدار داده شده‌ی $X = x$ ، تابع درست‌نمایی X را تابع چگالی احتمال توأم X یعنی $f_{\underline{\theta}}(x)$ تعریف می‌کنیم که به صورت تابعی از $\underline{\theta}$ در نظر گرفته می‌شود و آن را با نماد $L(\underline{\theta})$ نشان می‌دهیم:

$$L(\underline{\theta}) = f_{\underline{\theta}}(x)$$

اگر $X_1, \dots, X_n$ ، نمونه تصادفی n تایی از توزیعی از خانواده‌ی چگالی‌های $\{f_{\underline{\theta}}(x) : \underline{\theta} \in \Theta\}$ باشد، آن گاه:

$$L(\underline{\theta} | \underline{X}) = L(\underline{\theta}_1, \dots, \underline{\theta}_K | X_1, \dots, X_n) = \prod_{i=1}^n f(X_i | \underline{\theta}_1, \dots, \underline{\theta}_K)$$

MLE پارامتر $\underline{\theta}$ با کمک حل معادله $\frac{\partial}{\partial \underline{\theta}} L(\underline{\theta}) = 0$ (یا معادل آن $\frac{\partial}{\partial \underline{\theta}} \ln L(\underline{\theta}) = 0$) به دست می‌آید. بنابراین اگر $\delta(x)$ برآوردگری برای $\underline{\theta}$ باشد به طوری که:

$$i) P_{\underline{\theta}}(\delta(x) \in \Theta) = 1, \forall \underline{\theta} \in \Theta$$

$$ii) L(\delta(X)) \geq L(\underline{\theta}), \forall \underline{\theta} \in \Theta$$

آن گاه $\delta(X)$ برآوردگر ماکسیمم درست‌نمایی $\underline{\theta}$ در نظر گرفته می‌شود.

توزیع نرمال: متغیر تصادفی X دارای توزیع نرمال با پارامترهای μ, σ^2 ($X \sim N(\mu, \sigma^2)$) است، اگر دارای تابع چگالی احتمال زیر باشد:

$$f_{\mu, \sigma^2}(x) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left\{-\frac{1}{2\sigma^2}(x - \mu)^2\right\}, x \in \mathbb{R}, \mu \in \mathbb{R}, \sigma > 0$$

۱۶. گزینه (۳)

با توجه به خاصیت پایایی برآوردگر ماکسیمم درست‌نمایی (اگر $\hat{\theta}$ برآوردگر ماکسیمم درست‌نمایی θ باشد، آن گاه برای هر تابعی از θ مثل $\tau(\theta)$ ، برآوردگر ماکسیمم درست‌نمایی $\tau(\hat{\theta})$ برابر با $\tau(\theta)$ است)، برای به‌دست آوردن فاصله اطمینان برای σ کفایت از کران‌ها جذر گرفته شود.
ذکر این نکته ضروری است که سطح اطمینان تغییری نخواهد کرد.

۱۷. گزینه (۴)

اگر T یک آماره بسنده برای پارامتر θ باشد و MLE منحصر به فرد θ ، یعنی $\hat{\theta}$ وجود داشته باشد، آن گاه $\hat{\theta}$ باید تابعی از T باشد و اگر MLE منحصر به فرد پارامتر θ وجود نداشته باشد، آن گاه یک MLE پارامتر θ را می‌توان انتخاب کرد که تابعی از T باشد.

۱۸. گزینه (۱)

$(\bar{X}, S^2)$ یک آماره بسنده توأم کامل برای (μ, σ^2) است. از طرفی برای هر σ^2 ثابتی، $\bar{X}$ یک آماره بسنده کامل برای μ است و S^2 دارای توزیع کای دو با $n-1$ درجه آزادی است که توزیع آن به μ بستگی ندارد. بنابراین، بنابر قضیه باسو، برای هر μ و σ^2 ثابت، $\bar{X}$ و S^2 مستقل از هم هستند، لذا هر ترکیب خطی از آن‌ها نیز از هم مستقل هستند.
$$P((\bar{X} - \mu)(S^2 - \sigma^2) > 0) = P(\bar{X} - \mu)P(S^2 - \sigma^2) > 0$$

این رابطه زمانی برقرار است که هر دو احتمال مثبت و یا هر دو احتمال منفی باشند، پس داریم:

$$\begin{aligned} ۲۹ \quad P((\bar{X} - \mu)(S^2 - \sigma^2) > 0) &= (P(\bar{X} - \mu) > 0)(P(S^2 - \sigma^2) > 0) + (P(\bar{X} - \mu) < 0)(P(S^2 - \sigma^2) < 0) \\ &= P(\bar{X} > \mu)P(S^2 > \sigma^2) + P(\bar{X} < \mu)P(S^2 < \sigma^2) \\ &= \frac{1}{2} \times \frac{1}{2} + \frac{1}{2} \times \frac{1}{2} = \frac{1}{2} \end{aligned}$$

نکات:

خانواده توزیع‌های تولید شده توسط آماره T کامل است، اگر هیچ برآوردگر ناریب صفر غیر از برآوردگر صفر وجود نداشته باشد. یعنی تنها برآوردگر ناریب صفر در این خانواده خود صفر است. به عبارت دیگر:

$$\begin{aligned} E_{\theta}(g(T)) &= 0, \quad \forall \theta \in \Theta \\ \Rightarrow P_{\theta}(g(T) = 0) &= 1, \quad \forall \theta \in \Theta \end{aligned}$$

باید توجه نمود که کامل بودن یک ویژگی از خانواده توزیع‌های احتمالی است نه یک ویژگی از یک توزیع خاص. به طور مثال اگر $X \sim N(0, 1)$ ، با تعریف $g(x) = x$ داریم:

$$E g(X) = 0 \Rightarrow P(g(X) = 0) = 0$$

که مشاهده می‌شود X یک آماره کامل نیست. با این وجود این حالتی از یک توزیع خاص می‌باشد نه یک خانواده از توزیع‌ها. اما اگر $X \sim N(\theta, 1)$ آن گاه:

$$E g(X) = 0 \Rightarrow P(g(X) = 0) = 1, \quad \forall \theta \in \Theta$$

بنابراین خانواده توزیع‌های $N(\theta, 1)$ کامل می‌باشد.

$$f(x|\lambda) = \frac{e^{-\lambda} \lambda^x}{x!} \Rightarrow \text{Log } f(x|\lambda) = -\lambda + x \ln \lambda - \ln(x!)$$

$$\Rightarrow \frac{\partial}{\partial \lambda} \ln f(x|\lambda) = -1 + \frac{x}{\lambda} \Rightarrow \left(\frac{\partial}{\partial \lambda} \ln f(x|\lambda) \right)^2 = 1 + \frac{x^2}{\lambda^2} - \frac{2x}{\lambda}$$

$$E \left[\left(\frac{\partial}{\partial \lambda} \ln f(x|\lambda) \right)^2 \right] = 1 + \frac{E(x^2)}{\lambda^2} - \frac{2E(x)}{\lambda} = 1 + \frac{\lambda + \lambda^2}{\lambda^2} - \frac{2\lambda}{\lambda} = \frac{1}{\lambda} \quad (1)$$

(در توزیع پواسن مقدار میانگین برابر با واریانس است یعنی $E(X) = \text{Var}(X) = \lambda$)

$$[\omega'(\theta) + b'(\theta)]^2 = [(\lambda)' + 0]^2 = 1 \quad (2)$$

$$(1), (2) \Rightarrow \text{Var}(T) \geq \frac{1}{n \times \frac{1}{\lambda}} = \frac{\lambda}{n}$$

نکات:

نامساوی کرامر راثو، کران پایین برای واریانس برآوردگرهای انتخابی ارائه می‌دهد.

فرض کنید $X_1, \dots, X_n$ دارای توزیع توأم با تابع چگالی احتمال توأم $P_\theta(X_1, \dots, X_n)$ باشند. اگر T برآوردگر پارامتر $\omega(\theta)$ باشد، آن‌گاه نامساوی کرامر-راثو براساس یک نمونه تصادفی n تایی برای هر $\theta \in \Theta$ به یکی از دو روش زیر خواهد بود:

$$(1) \quad \text{Var}_\theta(T) \geq \frac{[\omega'(\theta) + b'(\theta)]^2}{E_\theta \left[\frac{\partial}{\partial \theta} \ln P_\theta(X_1, \dots, X_n) \right]^2}$$

$$(2) \quad \text{Var}_\theta(T) \geq \frac{(\omega'(\theta) + b'(\theta))^2}{n \times E_\theta \left[\left(\frac{\partial}{\partial \theta} \ln f_\theta(x) \right)^2 \right]}$$

به‌طوری‌که: $b(\theta) = E_\theta(T) - \omega(\theta)$ بیانگر اریبی برآوردگر T در برآورد پارامتر $\omega(\theta)$ است.

$$1) (Z_1 + Z_2) \sim N(0, 2) \Rightarrow \frac{Z_1 + Z_2}{\sqrt{2}} \sim N(0, 1) \Rightarrow \left(\frac{Z_1 + Z_2}{\sqrt{2}} \right)^2 \sim X_{(1)}^2$$

می‌دانیم:

$$2) (Z_1 - Z_2) \sim N(0, 2) \Rightarrow \frac{Z_1 - Z_2}{\sqrt{2}} \sim N(0, 1) \Rightarrow \left(\frac{Z_1 - Z_2}{\sqrt{2}} \right)^2 \sim X_{(1)}^2$$

$$(1), (2) \Rightarrow \frac{(Z_1 + Z_2)^2 / 2}{(Z_1 - Z_2)^2 / 2} = \frac{(Z_1 + Z_2)^2}{(Z_1 - Z_2)^2} \sim \frac{X_{(1)}^2}{X_{(1)}^2} \sim F_{(1,1)}$$

۲۱. گزینه (۳)

اگر $\emptyset$ یک تابع آزمون با ناحیه بحرانی C و آماره آزمون $T(X)$ باشد:

$$\alpha = P(\text{خطای نوع اول}) = P(H_0 \text{ درست است} | T(X) \in C) = E_{H_0} \{ \emptyset(X) \}$$

$$\beta = P(\text{خطای نوع دوم}) = P(H_1 \text{ را قبول کنیم وقتی که } H_0 \text{ درست است}) = P_{H_1}(T(X) \notin C) = E_{H_1}(1 - \emptyset(X))$$

$$\beta^* = \emptyset = 1 - \beta = 1 - E_{H_1}(1 - \emptyset(X)) = E_{H_1}(\emptyset(X)) = P_{H_1}(T(X) \in C) = P(H_1 \text{ را رد کنیم هنگامی که } H_1 \text{ درست است})$$

۲۲. گزینه (۲)

با کمک لیم نیمن-پیرسون مسئله را حل می‌کنیم:

$$\begin{cases} H_0 : f_0(x) = 4x^3 \\ H_1 : f_1(x) = 3x^2 \end{cases}$$

$$\frac{f_1(x)}{f_0(x)} > K \Rightarrow \frac{3x^2}{4x^3} > K \Rightarrow x < \frac{3}{4K} \Rightarrow x < C$$

۳۱

$$\emptyset^*(x) = \begin{cases} 1 & X < C \\ 0 & \text{در غیر این صورت} \end{cases}$$

$$\alpha = E_{H_0}(\emptyset^*(X)) = \int_0^C 4x^3 dx = x^4 \Big|_0^C = C^4$$

$$\alpha = \frac{1}{16} \Rightarrow C^4 = \frac{1}{16} \Rightarrow C = \frac{1}{2}$$

بنابراین ناحیه بحرانی برابر با $X < \frac{1}{2}$ است.

نکات:

لیم نیمن-پیرسون: اگر X یک متغیر تصادفی با تابع چگالی احتمال $f(x|\theta)$ باشد و علاقه‌مند به آزمون $H_0: \theta = \theta_0$ در مقابل $H_1: \theta = \theta_1$ باشیم.

اگر $\emptyset^*$ را برای هر $0 \leq K \leq \infty$ به صورت زیر تعریف کنیم:

$$\emptyset^*(X) = \begin{cases} 1 & (x \in C) \quad f(x|\theta_1) > K f(x|\theta_0) \\ 0 & (x \in C^c) \quad f(x|\theta_1) < K f(x|\theta_0) \end{cases}$$

(C ناحیه بحرانی و C^c ناحیه پذیرش)

به طوری که $\alpha = E_{\theta_0}(\emptyset^*(X))$ باشد، آن‌گاه $\emptyset^*(X)$ پرتوان‌ترین آزمون در سطح α برای آزمون H_0 در مقابل H_1 است.

۲۳. گزینه (۴)

با استفاده از قضیه دسته‌بندی فیشِر-نیمن داریم:

$$f(\underline{x}|\sigma^2) = \underbrace{\left(\frac{1}{2\pi\sigma^2}\right)^{-\frac{n}{2}} e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n X_i^2}}_{g(T(\underline{x})|\sigma^2)} \times \underbrace{1}_{h(\underline{x})}$$

در نتیجه $\sum_{i=1}^n X_i^2$ یک آماره بسنده برای پارامتر σ^2 است.

نکات:

قضیه دسته‌بندی فیشِر-نیمن: آماره $T(X)$ برای پارامتر θ بسنده است اگر و تنها اگر برای هر $\theta \in \Theta$ توابع غیرمنفی g و h وجود داشته باشند که:

$$f(\underline{x}|\theta) = g(T(\underline{x})|\theta)h(\underline{x})$$

به‌طوری‌که $g(T(X)|\theta)$ فقط از طریق $T(X)$ به x بستگی دارد و $h(\underline{x})$ به θ بستگی ندارد. با توجه به این‌که هر

تابع یک به یک از آماره بسنده نیز بسنده می‌باشد، بنابراین $e^{\frac{T(X)}{\sigma^2}}$ نیز یک آماره بسنده است.

۲۴. گزینه (۱)

اگر واریانس T_n (برآوردگر نارایب $\omega(\theta)$) با کران پایین کرامر-رائو برابر باشد، T_n را یک برآوردگر θ' را می‌نامیم. کران

پایین کرامر-رائو در این حالت به‌صورت $\frac{\omega'(\theta)^2}{I(\theta)}$ (ماتریس اطلاع) است. لذا کارایی برآوردگر نارایب T_n به‌صورت

زیر است:

$$e_n(T_n) = \frac{(\omega'(\theta))^2 / I(\theta)}{\text{Var}_{\theta}(T_n)}$$

$e_n(T_n)$ عددی بین صفر و یک است و برای یک برآوردگر کارا، کارایی یک است.

از طرفی اگر T یک برآوردگر UMVU پارامتر $\omega(\theta)$ باشد و T^* نیز یک برآوردگر نارایب $\omega(\theta)$ باشد، آنگاه کارایی T^* عبارتست از:

$$e_n^*(T^*) = \frac{V_{\theta}(T)}{V_{\theta}(T^*)}$$

بنابراین $1 \leq e_n^*(T^*) < \infty$ و کلیه برآوردگرهای UMVU با این تعریف برآوردگر کارا هستند. لذا برای به‌دست آوردن برآورد کارا برای پارامتر θ کفایت برآوردگر UMVU پارامتر θ را به‌دست آوریم، لذا داریم:

$$f(x|\theta) = \frac{1}{\theta} e^{-\frac{x}{\theta}} \Rightarrow L(\underline{x}|\theta) = \theta^{-n} e^{-\frac{1}{\theta} \sum_{i=1}^n X_i}$$

با توجه به فرم نمایی توزیع، $T = \sum_{i=1}^n X_i$ ، آماره بسنده کامل برای پارامتر θ است. حال تابعی از آماره بسنده کامل را به

دست می‌آوریم که برای پارامتر θ نارایب باشد، یعنی $E(g(T)) = \theta$.

$$T = \sum_{i=1}^n X_i \sim \exp(n\theta) \Rightarrow E(T) = n\theta \Rightarrow E\left(\frac{T}{n}\right) = \theta$$

لذا $\bar{X} = \frac{\sum_{i=1}^n X_i}{n} = \bar{T}$ ، $\frac{T}{n}$ پارامتر θ UMVUE است.

نکات: قضیه رائو-بلکول-لهمن-شفه: فرض کنید X یک متغیر تصادفی با توزیعی از خانواده توزیع‌های $\{P_\theta : \theta \in \Theta\}$ باشد. اگر $T(X)$ یک آماره بسنده کامل برای θ باشد، آن‌گاه برای هر پارامتر برآوردپذیر $\omega(\theta)$ ، برآوردگر ناریب یکتا با کم‌ترین واریانس (UMVUE) بر پایه $T(X)$ وجود دارد.

آماره $T(X)$ یک برآوردگر ناریب با کم‌ترین واریانس برای $\omega(\theta)$ می‌باشد اگر:

$$E_\theta(T(X)) = \omega(\theta), \forall \theta \in \Theta$$

و برای هر برآوردگر ناریب $h(X)$ پارامتر $\omega(\theta)$ داشته باشیم:

$$\text{Var}_\theta(T(X)) \leq \text{Var}_\theta(h(X)), \forall \theta \in \Theta$$

۲۵. گزینه (۴)

پرتوان‌ترین آزمون یکنواخت برای هر آزمون فرضی وجود ندارد و فقط برای حالات خاصی با شرایط معین وجود دارد. حتی در موارد بسیار ساده که علاقه‌مند به آزمون فرض ساده در مقابل فرض مرکب با دو عضو هستیم، ممکن است پرتوان‌ترین آزمون یکنواخت وجود نداشته باشد. در حالت کلی، برای آزمون فرض‌های دو طرفه، پرتوان‌ترین آزمون یکنواخت وجود ندارد. اما به‌طور کلی، در صورتی که خانواده توابع توزیع دارای خاصیت MLR باشند، آن‌گاه آزمون UMP برای آزمون فرض‌های یک طرفه وجود دارد.

نکته: خانواده $\{f_\theta : \theta \in \Theta \subset \mathbb{R}\}$ دارای خاصیت نسبت درست‌نمایی یکنوا (MLR) در آماره $T(X)$ می‌باشد اگر برای

هر $\theta_1 > \theta_0$ ، نسبت $\frac{f_{\theta_1}(x)}{f_{\theta_0}(x)}$ تابعی غیر نزولی از $T(x)$ باشد.

۲۶. گزینه (۱)

اگر X دارای تابع مولد گشتاور $M_X(t)$ باشد، آن‌گاه:

$$E X^n = M_X^{(n)}(0) = \left. \frac{d^n}{dt^n} M_X(t) \right|_{t=0}$$

به عبارت دیگر گشتاور n ام برابر با مشتق n ام تابع مولد گشتاور در نقطه $t=0$ است. از طرفی داریم:

$$\text{Var}(X) = E(X^2) - [E(X)]^2$$

لذا برای به دست آوردن واریانس کفایت گشتاور اول و دوم متغیر X را محاسبه کنیم:

$$E(X) = \left. \frac{d}{dt} M_X(t) \right|_{t=0} = \left(\left. \frac{-1}{4} e^{-t} + \frac{1}{4} e^t \right|_{t=0} \right) = 0$$

$$E(X^2) = \left. \frac{d^2}{dt^2} M_X(t) \right|_{t=0} = \left(\left. \frac{1}{4} e^{-t} + \frac{1}{4} e^t \right|_{t=0} \right) = \frac{1}{2}$$

$$\Rightarrow \text{Var}(X) = E(X^2) - E(X)^2 = \frac{1}{2}$$

۲۷. گزینه (۲)

$$P(\bar{X} > x_i) = P(\bar{X} < x_i) = \frac{1}{2}$$

در توزیع نرمال داریم:

$$P(\bar{X} > \mu) = P(\bar{X} < \mu) = \frac{1}{2}$$

$$P(\bar{X} < \mu, \bar{X} > x_i) = P(\bar{X} < \mu) P(\bar{X} > x_i) = \frac{1}{2} \times \frac{1}{2} = \frac{1}{4}$$

در نتیجه:

۲۸. با توجه به نادرست بودن صورت سوال حذف شده است.

$$\text{با فرض } \text{Var} \left[\frac{x_1}{\sigma_1} + \frac{x_2}{\sigma_2} + \frac{x_3}{\sigma_3} \right] \text{ داریم:}$$

$$\begin{aligned} \text{Var} \left[\frac{x_1}{\sigma_1} + \frac{x_2}{\sigma_2} + \frac{x_3}{\sigma_3} \right] &= \frac{1}{\sigma_1^2} \text{Var}(X_1) + \frac{1}{\sigma_2^2} \text{Var}(X_2) + \frac{1}{\sigma_3^2} \text{Var}(X_3) \\ &\quad + \frac{2}{\sigma_1 \sigma_2} \text{Cov}(X_1, X_2) + \frac{2}{\sigma_1 \sigma_3} \text{Cov}(X_1, X_3) + \frac{2}{\sigma_2 \sigma_3} \text{Cov}(X_2, X_3) \\ &= \frac{1}{\sigma_1^2} \times \sigma_1^2 + \frac{1}{\sigma_2^2} \times \sigma_2^2 + \frac{1}{\sigma_3^2} \times \sigma_3^2 + 2\rho_{12} + 2\rho_{13} + 2\rho_{23} = 1 + 1 + 1 + 1 + 1 + 1 = 6 \end{aligned}$$

بنابراین با فرض درست بودن سوال، جواب گزینه ۲ است.

نکات:

نکته ۱: اگر X و Y دو متغیر تصادفی باشند و a و b ضرایب ثابت، آن گاه:

$$\text{Var}(ax + by) = a^2 \text{Var}(x) + b^2 \text{Var}(y) + 2ab \text{Cov}(x, y)$$

اگر X و Y مستقل از یکدیگر باشند داریم:

$$\text{Var}(ax + by) = a^2 \text{Var}(x) + b^2 \text{Var}(y)$$

نکته ۲: ضریب همبستگی بین دو متغیر تصادفی X و Y به صورت زیر می باشد:

$$\rho_{XY} = \frac{\text{Cov}(X, Y)}{\sigma_X \sigma_Y}$$

روش های آماری

۲۹. گزینه (۴)

مدل رگرسیون خطی ساده y بر روی X به صورت زیر می باشد:

$$y_i = \beta_0 + \beta_1 X_i + \varepsilon_i, i = 1, \dots, n$$

که در آن ε_i ها دارای توزیع $N(0, \sigma^2)$ می باشند و ε_i ها دو به دو ناهمبسته هستند یعنی $\text{cov}(\varepsilon_i, \varepsilon_j) = 0$. این مدل

رگرسیونی یک مدل ساده، خطی در پارامترها و خطی در متغیر پیش بین می باشد. ساده است به این دلیل که تنها یک متغیر پیش بین وجود دارد، خطی در پارامترها است چرا که هیچ پارامتری به صورت نمایی یا ضرب و تقسیم پارامتری دیگری بیان نشده است و همچنین خطی در متغیر پیش بین است چرا که این متغیر با توان یک ظاهر شده است. به این مدل خطی در پارامترها و متغیر پیش بین، مدل مرتبه اول می گویند.

برای تحلیل مدل رگرسیونی از رویکرد آنالیز واریانس که بر مبنای تقسیم مجموع مربعات و درجه آزادی متناظر با متغیر پاسخ y می باشد استفاده می شود.

$$SSTO = SSR + SSE$$

مجموع مربعات کل (SSTO)، میزان تغییر پذیری مشاهدات پاسخ (y_i) را اندازه می گیرد که هر چه تغییر پذیری بین

$$SSTO = \sum_{i=1}^n (y_i - \bar{y})^2$$

مشاهدات پاسخ بیش تر باشد SSTO بزرگ تر می شود که برابر است با

مجموع مربعات رگرسیونی (SSR)، معیاری است که آن بخش از تغییرات متغیر پاسخ که توسط خط رگرسیون بیان می شود را نشان می دهد. به عبارت دیگر، هر چه مقدار SSR نسبت به SSTO بیش تر شود، اثر متغیر X در محاسبه تغییرات کلی مشاهدات پاسخ بیش تر است. مقدار SSR به صورت زیر تعریف می شود:

$$SSR = \sum (\hat{y}_i - \bar{y})^2$$

مجموع مربعات خطا (SSE)، میزانی از تغییر پذیری متغیر پاسخ است که توسط خط رگرسیون قابل بیان نمی باشد و به صورت زیر تعریف می شود:

$$SSE = \sum (y_i - \hat{y}_i)^2$$

ضریب تعیین (R^2)، معیاری است که با استفاده از آن می توان اثر پیش گوی X را در کاهش تغییر پذیری متغیر پاسخ (کاهش عدم قطعیت در پیش بینی y) محاسبه نمود و برابر است با:

$$R^2 = \frac{SSR}{SSTO} = 1 - \frac{SSE}{SSTO}$$

۳۵

باتوجه به این که $0 \leq SSE \leq SSTO$ لذا $0 \leq R^2 \leq 1$.

بنابراین ضریب تعیین به عنوان میزان کاهش در تغییر پذیری متغیر پاسخ هنگام استفاده از پیش گوی X می باشد. یعنی هر چه R^2 بزرگ تر باشد، استفاده از متغیر X کاهش بیش تری در مجموع تغییرات کلی متغیر پاسخ ایجاد می کند (اثر متغیر X در محاسبه تغییرات کلی مشاهدات پاسخ بیش تر است)

اگر همه مشاهدات بر روی خط رگرسیون برازش شده قرار بگیرند آن گاه $SSE = 0$ و $R^2 = 1$. به عبارت دیگر تمام تغییرات متغیر پاسخ y_i توسط پیش گوی X قابل بیان است.

هنگامی که خط رگرسیون برازش داده شده به صورت افقی باشد آن گاه $b_1 = 0$ و $\hat{y}_i = \bar{y}$. بنابراین $SSE = SSTO$ و $R^2 = 0$. به عبارت دیگر رابطه خطی بین X و y وجود ندارد و پیش گوی x هیچ گونه کمکی در کاهش تغییر پذیری مشاهدات y_i با استفاده از مدل رگرسیونی نمی کند.

۳.۰. گزینه (۲)

تغییر مقیاس متغیر منجر به تغییرات متناظر در مقیاس ضرایب و خطاهای استاندارد خواهد شد، اما هیچ تغییری در معنی داری و تفسیر ضرایب ایجاد نمی کنند.

اگر متغیر X_j در عدد c ضرب شود آن گاه ضریب متناظر با آن بر عدد c تقسیم می شود، اما سایر ضرایب مدل ثابت می مانند.

اگر متغیر پاسخ در عدد c ضرب شود آن گاه تمام ضرایب رگرسیونی در عدد c ضرب می شوند. هیچ یک از آماره های آزمون t و F تحت تأثیر تغییر واحدهای اندازه گیری متغیرها قرار نمی گیرند.

Dependent Independent	y	cy	Y
X_1	$\beta_1 (se_1)$	$c\beta_1 (c \times se_1)$	—
cX_1	—	—	$\beta_1 / c (se_1 / c)$
X_p	$\beta_p (se_p)$	$c\beta_p (c \times se_p)$	$\beta_p (se_p)$
Intencept	$\beta_0 (se_0)$	$c\beta_0 (c \times se_0)$	$\beta_0 (se_0)$
SSR	SSR	$c^2 \times SSR$	SSR
R-squared	R^2	R^2	R^2

۳۱. گزینه (۲)

$$X \sim \text{Bin}\left(n, p = \frac{1}{2}\right) \sim \text{احتمال بهبودی}$$

$$p\text{-value} = p(RH_0 | H_0) = p(X \leq 1) = p(X = 0) + p(X = 1) = \binom{6}{0} \left(\frac{1}{2}\right)^6 + \binom{6}{1} \left(\frac{1}{2}\right)^6 = \frac{1}{64} + \frac{6}{64} = \frac{7}{64}$$

۳۲. گزینه (۳)

فرم کلی مدل رگرسیون خطی با $p-1$ متغیر پیشگو به صورت زیر می باشد:

$$y_i = \beta_0 + \beta_1 X_{i1} + \beta_2 X_{i2} + \dots + \beta_{p-1} X_{i,p-1} + \varepsilon_i, i = 1, \dots, n$$

که در آن پارامترهای متناظر با متغیرهای $X_1, \dots, X_{p-1}$ و β_0 عرض از مبدا مدل می باشد. مانده ها دارای توزیع نرمال $\varepsilon_i \sim N(0, \sigma^2)$ می باشند.

خطی بودن در این مدل اشاره به خطی بودن برحسب پارامترها (β_i) و نه خطی بودن سطوح پاسخ است. بنابراین گوئیم رگرسیون برحسب پارامترها خطی می باشد اگر بتوان آن را به فرم زیر نوشت:

$$y_i = c_{i0}\beta_0 + c_{i1}\beta_1 + c_{i2}\beta_2 + \dots + c_{i,p-1}\beta_{p-1} + \varepsilon_i$$

که در آن ضرایب $c_{i0}, \dots, c_{i,p-1}$ ها هستند که همان پیشگوهای مدل می باشند. برای مثال مدل رگرسیونی $y_i = \beta_0 + \beta_1 X_{i1} + \beta_2 X_{i2} + \varepsilon_i$ برحسب ضرایب $c_{i0} = 1$ و $c_{i1} = X_{i1}$ و $c_{i2} = X_{i2}$ برحسب پارامترهای خطی می باشد.

اما مدل رگرسیونی $y_i = \beta_0 \exp(\beta_1 X_{i1}) + \varepsilon_i$ یک مدل رگرسیون غیرخطی می باشد.

همچنین مدل $y_i = \beta_0 + \beta_1 X_{i1} + \beta_2 X_{i1}^2 + \beta_3 X_{i2} + \beta_4 X_{i1} X_{i2} + \varepsilon_i$ نیز یک مدل رگرسیون خطی برحسب پارامترها می باشد چرا که پارامترها به صورت خطی وارد مدل شده اند. حال به بررسی خطی بودن مدل های رگرسیونی مطرح شده می پردازیم:

$$1) y_i = \frac{1}{1 + e^{\beta_0 + \beta_1 X_{i1} + c_i}} \Rightarrow 1 + e^{\beta_0 + \beta_1 X_{i1} + c_i} = \frac{1}{y_i}$$

$$\Rightarrow \ln\left(\frac{1}{y_i} - 1\right) = \beta_0 + \beta_1 X_{i1} + c_i$$

لذا با انجام تبدیل لگاریتمی مشاهده می شود که مدل تبدیل یافته برحسب پارامترهای خطی می باشد.

$$2) y_i = \beta_0 + \beta_1 \frac{1}{X_{i1}^2 - 1}$$

با در نظر گرفتن $c_1 = \frac{1}{X_i^2 - 1}$ داریم: $y_i = \beta_1 + c_1 \beta_1$ که به فرم خطی تبدیل می‌شود.

$$3) y_i = \beta_0 + \beta_1 e^{\beta_2 X_i + e_i} \Rightarrow \frac{\ln y_i}{\ln \beta_0} = \ln \beta_1 + \beta_2 X_i + e_i$$

$$\Rightarrow \ln y_i = (\ln \beta_0) (\ln \beta_1) + (\ln \beta_0) \beta_2 X_i + (\ln \beta_0) e_i$$

پس از انجام تبدیل ملاحظه می‌گردد که مدل تبدیل یافته بر حسب پارامترهای غیر خطی است.

$$4) y_i = \beta_0 e^{\beta_1 X_i + e_i} \Rightarrow \ln y_i = \ln \beta_0 + \beta_1 X_i + \varepsilon_i$$

لذا با تبدیل لگاریتمی، مدل تبدیل یافته بر حسب پارامترهای خطی می‌باشد.

۳۳. گزینه (۲)

برای به دست آوردن حجم نمونه مورد نیاز در طرح دو نمونه‌ای در صورتی که حجم مورد نیاز در یک گروه k برابر حجم گروه دیگر باشد داریم:

$$n_2 = \frac{\left(\frac{Z_{\alpha}}{2} + Z_{\beta} \right)^2 \sigma^2 (1 + 1/k)}{\varepsilon^2}, n_1 = kn_2$$

که در آن σ انحراف معیار مشترک و ε اندازه اثر نامیده می‌شوند. در صورتی که حجم نمونه در دو گروه برابر در نظر گرفته

$$n_1 = n_2 = \frac{\left(\frac{Z_{\alpha}}{2} + Z_{\beta} \right)^2 \sigma^2}{\varepsilon^2}$$

شود داریم:

$$n_1 = n_2 = \frac{(2(1/5) \times \sigma^2)}{\left(\frac{1}{2} \sigma \right)^2} = \frac{2 \times 10/5 \times \sigma^2}{\frac{1}{4} \sigma^2} = 8 \times 10/5 = 16$$

لذا باتوجه به اطلاعات سؤال داریم:

۳۴. گزینه (۴)

در طرح‌های بلوکی، واحدهای آزمایشی ناهمگن به بلوک‌های همگن تقسیم می‌شوند و سپس تصادفی‌سازی از واحدهای آزمایشی درون هر بلوک انجام می‌شود.

به‌طور مثال فرض کنید در دو کارخانه A و B محصول نان در ۴ درجه حرارت مختلف تولید می‌شود. از آنجایی که ممکن است حجم نان تولید شده تحت تأثیر فرایندهای مختلف و مواد خام مورد استفاده در این دو کارخانه قرار بگیرد، و باتوجه به این که نان‌های تولید شده در هر کارخانه تقریباً مشابه به هم هستند، بنابراین استفاده از طرح بلوکی با اندازه ۴ در هر بلوک توصیه می‌شود.

طرح بلوکی کامل تصادفی

کارخانه (بلوک)	واحدهای آزمایش (درجه حرارت)			
	۱	۲	۳	۴
A	بالا	پایین	خیلی بالا	متوسط
B	متوسط	بالا	خیلی بالا	پایین

لذا مدل آماری خطی باید منعکس کننده اثر مداخله (درجه حرارت) و اثر بلوک (کارخانه تولید) باشد:

$$y = [\text{overall constant}] + [\text{Treatment Effect}] + [\text{Block Effect}] + [\text{Experimental Error}]$$

پاسخنامه تشریحی سال ۸۸-۸۷

برای مثال فوق اثر درجه حرارت با استفاده از سه ایندیکاتور X_1 و X_2 و X_3 و اثر بلوک با ایندیکاتور X_4 وارد مدل می‌شوند. یعنی:

$$Y_{ij} = \beta_0 + \beta_1 X_{ij1} + \beta_2 X_{ij2} + \beta_3 X_{ij3} + \beta_4 X_{ij4} + \varepsilon_{ij}$$

گاهی مواقع باتوجه به شرایط قادر به انجام بررسی همه‌ی مداخلات (Treatment) در یک بلوک نیستیم و لذا اندازه هر بلوک کم‌تر از اندازه ملاحظات می‌شود.

به‌طور مثال فرض کنید که در یک کارخانه برای بررسی کیفیت یک محصول لینی از ۵ فرمول‌بندی مختلف استفاده شده است. افراد مختلفی جهت بررسی کیفیت این محصول انتخاب می‌شوند. اما از آنجایی که فرمول‌بندی این محصولات خیلی مشابه یکدیگر می‌باشند، برای این که افراد تشخیص‌های درست‌تری داشته باشند تنها هر فرد ۳ محصول را بررسی خواهد کرد. باتوجه به این که افراد مختلف دارای سلاقی متفاوتی می‌باشند بنابراین هر فرد به‌عنوان یک بلوک در نظر گرفته می‌شود. باتوجه به این محدودیت، مشاهده می‌شود که هر فرد نمی‌تواند هر پنج فرمول محصول را در یک زمان ارزیابی کند و تنها سه مورد از پنج مورد محصول را می‌توان مورد ارزیابی قرار دهد. لذا هر فرد بیانگر یک بلوک ناقص می‌باشد.

فرمول محصول					افراد ارزیاب (بلوک)
۵	۴	۳	۲	۱	
		x	x	x	۱
	x		x	x	۲
x			x	x	۳
	x	x		x	۴
x		x		x	۵
x	x			x	۶
	x	x	x		۷
x		x	x		۸
x	x		x		۹
x	x	x			۱۰

در طرح‌های بلوکی تصادفی ناقص متعادل، تکرار هر مداخله (درمان) در کلیه بلوک‌ها برابر و کم‌تر از تعداد کل بلوک‌ها است. در مثال فوق هر مداخله دقیقاً سه بار با مداخله دیگر تکرار شده است. به‌طور مثال، فرمول‌بندی ۱ و ۲ باهم در بلوک‌های ۱، ۲ و ۳ تکرار شده‌اند.

۳۵. گزینه (۴)

از هر مجموعه‌ای با $p-1$ متغیر پیشگو، 2^{p-1} مدل‌های جایگزین می‌توانند ساخته شوند. این محاسبه براساس این واقعیت است که هر پیشگو را می‌توان از مدل خارج یا وارد مدل نمود. به‌عنوان مثال $2^4 = 16$ مدل مختلف از مجموع ۴ متغیر پیشگو می‌توان تشکیل داد. به این صورت که در ابتدا می‌توان مدل بدون متغیرهای X را به داده‌ها برازش داد یعنی $y_i = \beta_0 + \varepsilon_i$. سپس مدل‌های با تنها یک متغیر پیشگو (X_1, X_2, X_3, X_4)، مدل‌های با دو متغیر پیشگو (X_1, X_2, X_3 و X_1, X_2, X_4 و X_1, X_3, X_4 و X_2, X_3, X_4) و غیره را به داده‌ها برازش داد.

با افزایش تعداد متغیرهای پیشگو، بررسی دقیق همه مدل‌های رگرسیونی احتمالی امکان‌پذیر نمی‌باشد. لذا روش‌های انتخاب مدل برای شناسایی گروه کوچکی از مدل‌های رگرسیونی که براساس معیارهای مشخصی «مدل‌های مناسب‌تری» هستند توسعه یافته است. معیارهای زیادی برای مقایسه مدل‌های رگرسیونی وجود دارند که مهم‌ترین آن‌ها عبارتند از:

$$PRESS_p, SBC_p, AIC_p, C_p, R_{a,p}^2, R_p^2$$

معیار R_p^2 یا SSE_p : R_p^2 که به عنوان ضریب تعیین چندگانه تعریف می شود، معیاری است که با استفاده از آن می توان اثر متغیرهای توضیحی $X_1, \dots, X_p$ را در کاهش تغییرپذیری متغیر پاسخ محاسبه نمود که عبارتست از:

$$R_p^2 = 1 - \frac{SSE_p}{SSTO}$$

دلیل استفاده از معیار R_p^2 این نیست که مدلی را انتخاب کنیم که دارای بیشترین مقدار R_p^2 باشد. چرا که مقدار ضریب تعیین با اضافه کردن متغیرهای توضیحی به مدل کاهش پیدا نمی کند و بنابراین مدلی که شامل همه متغیرهای توضیحی باشد ($p-1$ متغیر) دارای بزرگترین مقدار R_p^2 است.

بنابراین هدف استفاده از معیار R_p^2 در انتخاب زیر مدل رگرسیونی مناسب از تعداد $p-1$ متغیر توضیحی اینست که دنبال نقطه ای باشیم که بعد از آن اضافه کردن متغیرهای توضیحی تغییر معنی داری در مقدار ضریب تعیین ایجاد نکند.

معیار $R_{a,p}^2$ یا MSE_p :

به دلیل محدودیت های معیار R_p^2 ، معیار $R_{a,p}^2$ (ضریب تعیین تعدیل شده) به عنوان معیاری جایگزین پیشنهاد می شود:

$$R_{a,p}^2 = 1 - \left(\frac{n-1}{n-p} \right) \times \frac{SSE_p}{SSTO} = 1 - \frac{MSE_p}{\frac{SSTO}{n-1}}$$

این معیار، تعداد پارامترهای مدل رگرسیونی را از طریق درجه آزادی در نظر می گیرد.

با توجه به فرمول مقدار $R_{a,p}^2$ افزایش می یابد اگر و تنها اگر مقدار MSE_p کاهش پیدا کند چرا که مقدار $\frac{SSTO}{n-1}$ ثابت است.

برخلاف معیار R_p^2 ، با افزایش تعداد پارامترها مقدار $R_{a,p}^2$ ممکن است کاهش پیدا کند، بدین صورت که مقدار افزایش در معیار R_p^2 به حدی کم است که برای جبران از دست دادن درجه آزادی اضافی کافی نیست.

معیار C_p (آماره مالوس):

این معیار با مجموع میانگین مربع خطای n مقدار برازش شده از هر زیر مدل رگرسیون مرتبط است.

$$\Gamma_p = \frac{1}{\sigma^2} \left[\sum_{i=1}^n (E(\hat{y}_i) - \mu_i)^2 + \sum_{i=1}^n \sigma^2 (\hat{y}_i)^2 \right]$$

یک برآوردگر مناسب Γ_p ، C_p می باشد:

$$C_p = \frac{SSE_p}{MSE(X_1, \dots, X_{p-1})} - (n - 2p)$$

اگر مدل برازش داده شده رگرسیونی با $p-1$ متغیر توضیحی دارای اریبی نباشد آن گاه $\mu_i \equiv E\{\hat{y}_i\}$ و داریم:

$$E\{C_p\} \approx p$$

بنابراین هنگامی که مقادیر C_p برای همه مدل های رگرسیونی احتمالی در مقابل p ترسیم شوند، مدل های با اریبی کم تر در نزدیکی خط $C_p = p$ قرار دارند. مدل های با اریبی قابل توجه بالای این خط و مدل های بدون اریبی پایین این خط قرار می گیرند.

در استفاده از معیار C_p به دنبال شناسایی زیرمجموعه‌هایی از متغیرهای $X_1, \dots, X_p$ هستیم که برای آن‌ها (۱) مقدار C_p کوچک است و (۲) مقدار C_p در نزدیکی p است.

مقدار کوچک C_p بیانگر کوچک بودن مجموع میانگین مربعات خطا است و نزدیکی مقدار C_p به p نشان‌دهنده کم بودن اربابی مدل رگرسیونی است.

معیارهای AIC_p و SBC_p :

علاوه بر معیارهای $R^2_{a,p}$ و C_p ، معیارهای اطلاعاتی آکائیکه (AIC_p) و بیزی شوارتز (SBC_p) نیز برای اضافه کردن هر متغیر پیشگو در مدل جرمه‌ای را در نظر می‌گیرند که با در نظر گرفتن این جرمه نتیجه‌گیری می‌شود که آیا اضافه کردن پیشگوی اضافه به مدل باعث به‌دست آوردن برآوردهای بهتری می‌شود یا خیر. این معیارها عبارتند از:

$$AIC_p = n \ln SSE_p - n \ln n + 2p$$

$$SBC_p = n \ln SSE_p - n \ln n + [\ln n]p$$

معیار $RRESS_p$ (مجموع پیش‌بینی مربعات):

معیاری است که با استفاده از آن می‌توان بررسی نمود که مقادیر برازش شده یک زیر مدل رگرسیونی تا چه حد می‌توانند پاسخ‌های مشاهده شده y_i را به‌خوبی پیش‌بینی نمایند.

مجموع مربعات خطا $SSE = \sum (y_i - \hat{y}_i)^2$ نیز دارای چنین ویژگی است. اما تفاوتی که معیار $RRESS$ با آن دارد در اینست که برای به‌دست آوردن مقدار برازش شده $\hat{y}_i$ ، فرد i ام از مطالعه حذف شده و برآورد تابع رگرسیونی برای زیر مدل مفروض با استفاده از $n-1$ مشاهده دیگر صورت می‌گیرد و با استفاده از آن مقدار پیش‌بینی شده $\hat{y}_{i(i)}$ برای فرد i ام به‌دست می‌آید.

لذا مقدار خطای پیش‌بینی آماره $PRESS$ برای فرد i ام عبارتست از: $y_i - \hat{y}_{i(i)}$ و مقدار $PRESS_p$ عبارتست از مجموع مربعات خطای پیش‌بینی برای همه n فرد، یعنی:

$$PRESS_p = \sum_{i=1}^n (y_i - \hat{y}_{i(i)})^2$$

مدل‌های با مقادیر کم‌تر $PRESS_p$ به‌عنوان کاندیدای انتخاب مدل بهینه در نظر گرفته می‌شوند. عامل تورم واریانس (VIF) شاخصی است که برای بررسی وجود هم‌خطی چندگانه در مدل رگرسیونی می‌پردازد.

۳۶. گزینه (۲)

روش آزمون خطی عمومی شامل سه مرحله اصلی است:

(۱) مدل کامل

ابتدا با مدلی شروع می‌کنیم که به‌نظر مناسب‌ترین مدل برای داده‌ها باشد که به مدل کامل یا نامحدود معروف است. برای مدل رگرسیونی خطی ساده با خطاهای نرمال، مدل کامل عبارتست از:

$$y_i = \beta_0 + \beta_1 X_i + \varepsilon_i$$

این مدل را با استفاده از روش کم‌ترین مربعات یا روش درست‌نمایی ماکزیمم به داده‌ها برازش داده و مجموع مربعات خطای آن را به دست می‌آوریم.

$$SSE(F) = \sum [y_i - (b_0 + b_1 X_i)]^2 = \sum (y_i - \hat{y}_i)^2 = SSE$$

۲) مدل کاهش یافته

مدل کاهش یافته شامل مدلی است که در آن متغیرهایی را که مایل به بررسی اثرات آن‌ها هستیم از مدل خارج می‌نماییم. در مورد مدل رگرسیون خطی ساده با خطاهای نرمال داریم:

$$H_0: \beta_1 = 0$$

$$H_a: \beta_1 \neq 0$$

لذا مدل کاهش یافته آن عبارتست از:

$$y_i = \beta_0 + \varepsilon_i$$

مشابه حالت قبل با استفاده از روش کم‌ترین مربعات یا روش درست‌نمایی ماکزیمم این مدل را به داده‌ها برازش داده و مجموع مربعات خطای آن را به دست می‌آوریم.

$$SSE(R) = \sum (y_i - b_0)^2 = \sum (y_i - \bar{y})^2 = SSTO$$

۳) آماره آزمون

حال به بررسی مجموع مربعات خطای به دست آمده از برازش دو مدل کامل و کاهش یافته می‌پردازیم یعنی $SSE(F)$ و $SSE(R)$.

همواره داریم $SSE(F) \leq SSE(R)$. دلیل این امر اینست که پارامترهای بیش‌تری در مدل کامل هستند، بنابراین برازش بهتری بر روی داده‌ها ایجاد می‌کنند که دارای انحرافات کم‌تری در اطراف تابع رگرسیون هستند. هنگامی که مقدار $SSE(F)$ خیلی کم‌تر از مقدار $SSE(R)$ نباشد، آن‌گاه استفاده از مدل کامل در مقایسه با مدل کاهش یافته تأثیر چندانی در کاهش تغییرپذیری مشاهدات y_i در اطراف خط برازش داده شده نخواهد داشت، لذا استفاده از مدل کاهش یافته به جای مدل کامل توصیه می‌شود.

بنابراین تفاوت کم $SSE(R) - SSE(F)$ بیانگر برقراری فرض H_0 و تفاوت زیاد $SSE(R) - SSE(F)$ بیانگر برقراری فرض H_a می‌باشد.

آماره آزمون مناسب برای بررسی فرضیه H_0 در مقابل H_a عبارتست از:

$$F^* = \frac{SSE(R) - SSE(F)}{df_R - df_F} \div \frac{SSE(F)}{df_F}$$

که تحت فرض H_0 دارای توزیع $F(1 - \alpha; df_R - df_F, df_F)$ می‌باشد. لذا قاعده تصمیم‌گیری براساس آن عبارتست از:

$$\begin{cases} \text{If } F^* \leq F(1 - \alpha; df_R - df_F, df_F), \text{ Conclude } H_0. \\ \text{If } F^* > F(1 - \alpha; df_R - df_F, df_F), \text{ Conclude } H_a \end{cases}$$

۳۷. گزینه (۱)

گاهی مواقع در طرح‌های بلوکی به جای استفاده از یک متغیر بلوک‌بندی شده، بیش‌تر از یک متغیر بلوک‌بندی شده مورد استفاده قرار می‌گیرد تا از خطاهای آزمایشی، تغییرات مربوط به هریک از متغیرهای بلوک‌بندی شده حذف شوند. به عنوان مثال متغیرهای بلوک‌بندی ممکن است شامل سن و درآمد افراد باشد که هر بلوک خاص شامل افرادی با گروه سنی مشخص و درآمد معلومی می‌باشند.

با این حال، استفاده از طرح‌های بلوکی با بیش از یک متغیر بلوک‌بندی، اغلب نیازمند تعداد زیادی واحدهای آزمایشی است. به طور مثال اگر سن و درآمد هر کدام شامل ۶ کلاس باشند آن‌گاه ۳۶ بلوک تشکیل می‌شود، حال اگر تعداد درمان‌هایی که قرار است مورد مطالعه قرار بگیرند ۶ درمان باشد، آن‌گاه مطالعه نیازمند ۲۱۶ فرد برای بررسی است.

پاسخنامه تشریحی سال ۸۸-۸۷

با در نظر گرفتن ملاحظات مربوط به هزینه ممکن است مجاز به استفاده از ۲۱۶ واحد آزمایشی نباشیم، لذا باید دنبال طرح جایگزینی بود که با استفاده همزمان این دو متغیر بلوکی و تعداد کم‌تری از واحدهای آزمایشی، واریانس خطای آزمایشی را به اندازه کافی کاهش دهد. طرح مربع لاتین، طرح مفیدی برای چنین وضعیتی می‌باشد. فرض کنید آزمایش شامل سه درمان A و B و C باشد. همچنین فرض کنید که روزهای هفته (دوشنبه، سه‌شنبه و چهارشنبه) و اپراتور (شماره ۱ و شماره ۲ و شماره ۳) به عنوان متغیرهای بلوک‌بندی استفاده می‌شوند. یک طرح مربع لاتین به صورت زیر است:

روز	اپراتور		
	۳	۲	۱
دوشنبه	C	A	B
سه‌شنبه	B	C	A
چهارشنبه	A	B	C

توجه داشته باشید که هر اپراتور هر درمان را اجرا می‌کند و تمام درمان‌ها در هر روز اجرا می‌شوند. طرح مربع لاتین دارای خصوصیات زیر می‌باشد:

(۱) دارای r درمان (Treatment) می‌باشد.

(۲) دارای دو متغیر بلوک‌بندی هریک با r کلاس می‌باشد.

(۳) هر سطر و هر ستون شامل تمام درمان‌ها می‌باشد.

مزایای این طرح عبارتست از:

(۱) استفاده همزمان از دو متغیر بلوک‌بندی به جای یک متغیر بلوک‌بندی باعث کاهش بیش‌تری در تغییرات خطاهای آزمایشی می‌شود.

(۲) در طرح‌های اندازه‌های تکراری (repeated measure) با استفاده از طرح مربع لاتین قادر به بررسی ترتیب مکانی اثرات درمان‌ها هستیم.

معایب این طرح عبارتند از:

(۱) تعداد کلاس‌های هر متغیر بلوک‌بندی باید برابر با تعداد درمان‌ها باشد که این محدودیت در عمل سخت است.

(۲) فرضیه‌های مدل محدود می‌باشند (به عنوان مثال هیچ اثر متقابل بین هریک از متغیرهای بلوکی با درمان و همچنین بین دو متغیر بلوکی وجود ندارد).

(۳) استفاده از طرح مربع لاتین زمانی که تعداد کمی درمان مورد مطالعه قرار می‌گیرند، به تعداد کمی از درجه آزادی برای خطای آزمایشی منجر خواهد شد. از طرف دیگر، هنگامی که تعداد درمان‌های مورد مطالعه زیر باشند، درجه آزادی خطای آزمایشی ممکن است بزرگ‌تر از ضرورت باشد.

باتوجه به این محدودیت هنگامی که بیش از ۸ درمان مورد بررسی قرار بگیرند، طرح مربع لاتین به ندرت مورد استفاده قرار می‌گیرد. اگر تعداد درمان‌ها کم (کم‌تر از ۴) باشد آن‌گاه نیازمند تکرارهای اضافه درمان‌ها هستیم.

مدل طرح مربع لاتین با درمان‌های ثابت و اثرات بلوکی به صورت زیر است:

$$y_{ijk} = \mu_{...} + p_i + k_j + \tau_k + \varepsilon_{ijk}$$

که:

$\mu_{...}$ is constant

p_i, k_j, τ_k are constants subject to the restrictions $\sum p_i = \sum k_j = \sum \tau_k = 0$

ε_{ijk} are independent $N(0, \sigma^2)$

$i = 1, \dots, r$; $j = 1, \dots, r$; $k = 1, \dots, r$

p_i اثر متغیر بلوکی سطری، k_j اثر متغیر بلوک ستونی و τ_k اثر اصلی درمان می‌باشد. جدول زیر، جدول آنالیز واریانس طرح مربع لاتین را با اثرات ثابت نشان می‌دهد:

Source of variation	ss	df	MS	$E\{MS\}$
Row blocking variable	SSROW	$r-1$	$MSROW = \frac{SSROW}{r-1}$	$\sigma^2 + r \frac{\sum p_i^2}{r-1}$
Column blocking variable	SSCOL	$r-1$	$MSCOL = \frac{SSCOL}{r-1}$	$\sigma^2 + r \frac{\sum k_j^2}{r-1}$
Treatments	SSTR	$r-1$	$MSTR = \frac{SSTR}{r-1}$	$\sigma^2 + r \frac{\sum \tau_k^2}{r-1}$
Error	SSRem	$(r-1)(r-2)$	$MSRem = \frac{SSRem}{(r-1)(r-2)}$	σ^2
Total	SSTO	r^2-1		

در حالتی که n تکرار در هر سلول داشته باشیم مدل به صورت زیر می‌شود:

$$y_{ijkm} = \mu_{...} + p_i + k_j + \tau_k + \varepsilon_{ijkm}$$

$$i = 1, \dots, r; j = 1, \dots, r; k = 1, \dots, r; m = 1, \dots, n$$

جدول زیر، جدول آنالیز واریانس طرح مربع لاتین با تعداد n تکرار در هر سلول را نشان می‌دهد:

Source of variation	ss	df	MS	MS
Row blocking variable	SSROW	$r-1$		MSROW
Column blocking variable	SSCOL	$r-1$		MSCOL
Treatments	SSTR	$r-1$		MSTR
Error	SSRem	$nr^2 - 3r + 2$		MSRem
Total	SSTO	$nr^2 - 1$		

$$r = 3 \Rightarrow df(\text{Error}) = nr^2 - 3r + 2 = 2(9) - (3 \times 3) + 2 = 11$$

باتوجه به اطلاعات مسئله داریم:

۳۸. گزینه (۱)

واریانس (σ^2) عبارت خطا (ε_i) در مدل رگرسیون $y_i = \beta_0 + \beta_1 X_i + \varepsilon_i$ با خطاهای نرمال برای به‌دست آوردن استنباط‌های مختلف در مورد تابع رگرسیونی و پیش‌بینی مقادیر y باید برآورد شود. یک برآوردگر نارایب برای σ^2 ، MSE می‌باشد یعنی $E(MSE) = \sigma^2$.

$$s^2 = MSE = \frac{SSE}{n-p-1} = \frac{\sum (y_i - \hat{y}_i)^2}{n-p-1}$$

که در آن SSE دارای درجه آزادی $n-p-1$ متناظر با p متغیر توضیحی وارد شده در مدل می‌باشد. اگر توزیع خطاها در مدل رگرسیونی انحراف زیادی از توزیع نرمال داشته باشد آن‌گاه MSE برآوردگر نیرومند σ^2 نخواهد بود. برآوردگر نیرومند، برآوردگری است که در صورت انحراف از توزیع نرمال، $E|MSE| = \sigma^2$. در صورت عدم برخورداری خطاها از توزیع نرمال، برآوردگر انحراف مطلق میانه (MAD) برآوردگری نیرومند برای σ خواهد بود.

$$MAD = \frac{1}{n} \text{median}\{ |e_i - \text{median}\{e_i\}| \}$$

۳۹. گزینه (۳)

ضریب همبستگی رتبه‌ای کندال (τ) یک آماره ناپارامتری است که برای سنجش همبستگی آماری میان دو متغیر تصادفی رتبه‌ای به کار می‌رود.

فرض کنید $(X_1, Y_1), (X_2, Y_2), \dots, (X_n, Y_n)$ مقادیری از دو متغیر با توزیع توأم X و Y باشند. دو زوج (X_i, Y_i) و (X_j, Y_j) را سازگار (concordant) می‌نامیم اگر رتبه‌های دو مشاهده باهم برابر باشند یعنی هرگاه $X_i > X_j$ آن‌گاه $Y_i > Y_j$ یا اگر $X_i < X_j$ آن‌گاه $Y_i < Y_j$ برقرار باشند.

دو زوج ناسازگار (discordant) نامیده می‌شوند هرگاه $X_i > X_j$ آن‌گاه $Y_i < Y_j$ یا اگر $X_i < X_j$ آن‌گاه $Y_i > Y_j$ باشد. اگر $X_i = X_j$ یا $Y_i = Y_j$ ، زوج نه سازگار و نه ناسازگار نامیده می‌شوند. آن‌گاه ضریب همبستگی کندال به صورت زیر محاسبه می‌شود:

$$\tau = \frac{\text{تعداد زوج‌های ناسازگار} - \text{تعداد زوج‌های سازگار}}{\frac{1}{2}n(n-1)}$$

مقدار این ضریب بین $+1$ و -1 که این دو مقدار متناظر با تطابق و عدم تطابق کامل است. در صورت سؤال $n = 6$ فرد به لحاظ متغیرهای نگرش (X) و عملکرد (Y) مورد بررسی قرار گرفته‌اند. تعداد کل زوج‌های مورد بررسی عبارتست از

$$15 = \binom{6}{2} \text{ زوج. با توجه به این که ۳ زوج ناسازگار بوده و هیچ زوج هم رتبه‌ای نیز نداشته‌ایم داریم.}$$

$$\tau = \frac{12-3}{\frac{1}{2} \times 6 \times (5)} = \frac{9}{15} = \frac{3}{5} = 0.6$$

۴۰. گزینه (۳) ۴۴

طرح اندازه‌های تکراری افراد یکسانی (شخص، فروشگاه، کارخانه و ...) برای هریک از روش‌های مورد مطالعه مورد استفاده قرار می‌دهد. این کار معمولاً به دو صورت رخ می‌دهد:

(۱) زمانی که افراد چندین بار در طول زمان اندازه‌گیری می‌شوند تا تغییرات روش‌های مورد مطالعه را مشاهده نمایند.
(۲) هنگامی که افراد تحت وضعیت‌ها یا آزمایش‌های مختلف قرار می‌گیرند و پاسخ به هریک از این وضعیت‌ها با هم مقایسه می‌شوند.

برای مثال با استفاده از یک طرح اندازه‌های تکراری می‌توان بررسی نمود که آیا مصرف سیگار در بین افراد با مصرف بالای سیگار بعد از یک برنامه هیپنوتراپی (به عنوان مثال با سه نقطه زمانی: بلافاصله قبل، یک ماه بعد و ۶ ماه بعد از برنامه هیپنوتیزم) تغییر کرده است.

به عنوان مثال دیگر، برای بررسی این که آیا رنگ شیشه جلوی ماشین بر میزان سرعت راننده تأثیر می‌گذارد یا خیر می‌توان با استفاده از طرح اندازه‌های تکراری میزان سرعت ماشین را در چهار وضعیت بدون رنگ، رنگ ضخیم، رنگ متوسط و رنگ تیره شیشه جلوی ماشین بررسی نمود.

مزیت اصلی طرح‌های اندازه‌گیری تکراری این است که دقت خوبی برای مقایسه مداخلات ارائه می‌دهند چرا که همه منابع تغییرپذیری بین افراد از خطای آزمایش حذف می‌شوند و فقط تغییرات درون فردی وارد خطاهای آزمایش می‌شوند. هنگامی که قصد استفاده از طرح اندازه‌های تکراری را برای تجزیه و تحلیل داده‌ها داریم، ابتدا باید اطمینان حاصل نمود که این داده‌ها با استفاده از طرح اندازه‌های تکراری قابل تحلیل هستند. بنابراین تنها زمانی استفاده از طرح اندازه‌های تکراری مناسب است که فرضیه‌های مورد نیاز این طرح در مورد داده‌ها برقرار باشد. که مهم‌ترین آن‌ها عبارتند از:

- (۱) عدم وجود داده‌های پرت معنی‌دار در مشاهدات
- (۲) داده‌ها در گروه‌های مورد مطالعه به صورت تقریبی دارای توزیع نرمال باشند.
- (۳) برقراری شرط کرویت: یعنی همبستگی بین مشاهدات در همه تکرارهای زمانی یکسان است.

۴۱. گزینه (۲)

در طرح آنالیز واریانس و برخی روش‌های دیگر دو نوع از عوامل وجود دارند: عوامل با اثرات ثابت و عوامل با اثرات تصادفی. مناسب بودن هر کدام از این دو نوع بستگی به زمینه مشکل، سؤالات مورد علاقه و نحوه جمع‌آوری داده‌ها دارد. عامل اثرات ثابت: داده‌ها برای تمام سطوح عامل مورد نظر جمع‌آوری شده‌اند. به‌طور مثال: در یک آزمایش هدف مقایسه اثرات سه دوز مختلف یک دارو بر روی پاسخ بیماران می‌باشد. در این‌جا، "مقدار مصرف دوز" عامل، و سه سطح مختلف دوز، سطوح عامل مورد مطالعه می‌باشند و هیچ مطلبی در مورد دوزهای دیگر گفته نشده است. عامل اثرات تصادفی: عامل مورد نظر دارای سطوح زیادی است که علاقه‌مند به بررسی تمام سطوح عامل می‌باشیم اما تنها یک نمونه تصادفی از سطوح در داده‌ها وجود دارد. به‌طور مثال یک تولیدکننده بزرگ محصولات لوازم خانگی، علاقه‌مند به مطالعه اثر اپراتور ماشین بر کیفیت نهایی محصول است. محقق یک نمونه تصادفی از اپراتورها را از تعداد زیادی از اپراتورها برای بررسی کیفیت نهایی محصول انتخاب می‌نماید. تجزیه و تحلیل داده‌ها بسته به این‌که عامل مورد نظر ثابت یا تصادفی می‌باشد متفاوت است. بنابراین اگر عامل به‌صورت نادرست طبقه‌بندی شود ممکن است منجر به نتیجه‌گیری نادرست شود. مدل Anova با اثرات تصادفی با حجم نمونه برابر در هریک از سطوح عامل عبارتست از:

$$y_{ij} = \mu_i + \varepsilon_{ij}$$

Where:

$$\mu_i \text{ are independent } N(\mu, \sigma_\mu^2)$$

$$\varepsilon_{ij} \text{ are independent } N(0, \sigma^2)$$

$$\mu_i \text{ and } \varepsilon_{ij} \text{ are independent random variables}$$

$$i = 1, \dots, r ; j = 1, \dots, n$$

۴۵

در مدل Anova با اثرات تصادفی، استنباط‌ها معمولاً به‌طور مجزا در مورد هر سطح عامل انجام نمی‌شود و علاقه‌مند به استنباط در مورد کل سطوح عامل هستیم. بنابراین میزان کل تغییرپذیری متغیر پاسخ را که به‌وسیله تغییرپذیری سطوح عامل قابل محاسبه است به‌دست می‌آوریم که عبارتست از:

$$\frac{\sigma_\mu^2}{\sigma_y^2} = \frac{\sigma_\mu^2}{\sigma_\mu^2 + \sigma^2}$$

که مقداری بین صفر و یک را شامل می‌شود. هرچه درصد تغییرات بیان شده متغیر پاسخ توسط سطوح عامل مورد نظر بیشتر باشد یعنی σ_μ^2 نسبت به σ^2 بزرگ‌تر باشد، این مقدار به‌سمت یک میل می‌نماید.

لذا برای بررسی این‌که آیا میانگین μ_i ها در سطوح مختلف برابر هستند یا خیر آزمون زیر را انجام می‌دهیم:

$$H_0 : \sigma_\mu^2 = 0$$

$$H_a : \sigma_\mu^2 > 0$$

فرضیه H_0 بیانگر اینست که همه μ_i ها برابر هستند یعنی $\mu_i \equiv \mu$. اما فرضیه H_a بیانگر این است که μ_i ها متفاوت هستند. تحلیل واریانس Anova با اثرات تصادفی مشابه با Anova با اثرات ثابت می‌باشد.

$$F^* = \frac{MSTR}{MSE} = \frac{\sigma^2}{\sigma^2 + n\sigma_\mu^2}$$

اگر $F^* > F[1-\alpha; r-1, r(n-1)]$ ، آنگاه فرضیه صفر رد می‌شود.

۴۲. گزینه (۲) این سوال حذف شده است.

مجموع مربعات اضافی (Extra sum of squares)، میزان کاهش در مجموع مربعات خطا را هنگامی که یک یا چند متغیر پیشگوی جدید وارد مدل رگرسیونی می‌شوند (با توجه به این که متغیرهای توضیحی دیگر در مدل هستند) را اندازه‌گیری می‌کند.

به‌طور معادل می‌توان مجموع مربعات اضافی را به‌عنوان میزان افزایش در مجموع مربعات رگرسیونی هنگامی که یک یا چند متغیر توضیحی وارد مدل می‌شوند تعریف کرد.

به‌طور مثال $SSR(X_p | X_1)$ مجموع مربعات اضافی مربوط به متغیر X_p وقتی که متغیر X_1 در مدل وجود دارد براساس دو تعریف بالا به‌صورت زیر است:

$$\begin{cases} SSR(X_p | X_1) = SSE(X_1) - SSE(X_1, X_p) \\ SSR(X_p | X_1) = SSR(X_1, X_p) - SSR(X_1) \end{cases}$$

همچنین داریم:

$$SSTO = SSR(X_1, X_p) + SSE(X_1, X_p) = SSR(X_1) + SSR(X_p | X_1) + SSE(X_1, X_p)$$

باتوجه به صورت سؤال متغیرهای X_1 و X_p در مدل وجود دارند، بنابراین مجموع مربعات اضافی مربوط به متغیر X_p هنگامی که دو متغیر X_1 و X_p در مدل هستند برابر است با:

$$SSR(X_p | X_1, X_p) = SSR(X_1, X_p, X_p) - SSR(X_1, X_p)$$

یا

$$SSR(X_p | X_1, X_p) = SSE(X_1, X_p) - SSE(X_1, X_p, X_p)$$

جواب در بین هیچ‌کدام از گزینه‌ها نمی‌باشد.

۴۳. گزینه (۱)

به سؤال ۳۵ رجوع شود.

۴۴. گزینه (۳)

مدل رگرسیون خطی ساده با یک متغیر پیشگو به‌صورت زیر می‌باشد:

$$y_i = \beta_0 + \beta_1 X_i + \varepsilon_i$$

که در آن ε_i ها خطاهای تصادفی با $E(\varepsilon_i) = 0$ و $\sigma^2(\varepsilon_i) = \sigma^2$ می‌باشند و دو به دو ناهمبسته هستند یعنی $\sigma(\varepsilon_i, \varepsilon_j) = 0$ برای همه $i \neq j$.

برای برآورد پارامترهای این مدل از روش حداقل مربعات استفاده می‌شود که این روش میزان انحراف مشاهدات را از مقادیر موردانتظار در نظر می‌گیرد یعنی:

$$Q = \sum_{i=1}^n (y_i - (\beta_0 + \beta_1 X_i))^2$$

باحل این معادله نسبت به پارامترهای β_0 و β_1 برآوردهای کم‌ترین مربعات پارامترها به‌صورت زیر به‌دست می‌آیند:

$$b_0 = \bar{y} - b_1 \bar{X}$$

$$b_1 = \frac{\sum (X_i - \bar{X})(y_i - \bar{y})}{\sum (X_i - \bar{X})^2}$$

قضیه گاوس - مارکف بیان می کند که تحت شرایط مدل رگرسیونی، برآوردهای حداقل مربعات b_0 و b_1 ، برآوردهایی نارایب و در بین همه برآوردهای نارایب خطی دارای کمترین واریانس هستند.

$$E(b_1) = \beta_1 \text{ و } E(b_0) = \beta_0.$$

دارای کمترین واریانس در بین همه برآوردهای نارایب خطی هستند یعنی این برآوردها نسبت به هر برآوردهای دیگری متعلق به کلاس برآوردهای نارایب که یک ترکیب خطی از مشاهدات پاسخ هستند تغییرپذیری کمتری دارند.

۴۵. گزینه (۱)

هنگامی که یک مدل رگرسیونی برای برازش به داده ها مورد استفاده قرار می گیرد، معمولاً نمی توان مطمئن بود که مدل مناسب داده ها می باشد. یک یا چند ویژگی مدل مانند خطی بودن تابع رگرسیونی یا نرمال بودن مانده های مدل، ممکن است برای داده های موجود برقرار نباشد. از این رو مهم است تا بررسی مناسب بودن مدل برای داده ها قبل از انجام استنباط براساس آن مدل صورت پذیرد. اگر مدل رگرسیون خطی برای برازش بر روی داده ها مناسب نباشد دو انتخاب اساسی وجود دارد:

(۱) رها کردن مدل رگرسیونی و توسعه و استفاده از یک مدل مناسب تر

(۲) استفاده از تبدیلات مناسب تا مدل رگرسیون برای داده های تبدیل شده مناسب باشد.

عدم نرمال بودن خطاها و نابرابری واریانس خطاها معمولاً باهم رخ می دهند. لذا برای اصلاح این انحرافها در مدل رگرسیون خطی ساده باید تبدیل مناسب بر روی متغیر پاسخ y صورت پذیرد. روش باکس - کاکس به طور خودکار یک تبدیل مناسب را از یک خانواده از تبدیلات توانی بر روی متغیر y شناسایی می کند.

خانواده تبدیلات توانی به فرم $y' = y^\lambda$ می باشد که پارامتر λ با استفاده از داده ها مشخص می شود. لذا مدل رگرسیونی به فرم زیر در می آید:

$$y_i^\lambda = \beta_0 + \beta_1 X_{i1} + \varepsilon_i$$

مدل رگرسیونی شامل یک پارامتر اضافی λ علاوه بر سایر پارامترهای مدل می باشد که باید برآورد شود. روش باکس - کاکس با استفاده از روش درست نمایی ماکزیمم پارامتر λ را مشابه برآورد سایر پارامترها برآورد می نماید.

۴۶. گزینه (۲)

فرم کلی مدل رگرسیون خطی با $p-1$ متغیر پیشگو به صورت زیر می باشد:

$$y_i = \beta_0 + \beta_1 X_{i1} + \beta_2 X_{i2} + \dots + \beta_{p-1} X_{i,p-1} + \varepsilon_i$$

Where:

$\beta_0, \beta_1, \dots, \beta_{p-1}$ are parameters

$X_{i1}, \dots, X_{i,p-1}$ are known constants

ε_i are independent $N(0, \sigma^2)$

$i = 1, \dots, n$

فرم ماتریسی مدل فوق به صورت زیر می باشد:

$$y_{n \times 1} = X_{n \times p} \beta_{p \times 1} + \varepsilon_{n \times 1}$$

که:

$$y_{n \times 1} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix}, \quad X_{n \times p} = \begin{bmatrix} 1 & X_{11} & X_{12} & \dots & X_{1,p-1} \\ 1 & X_{21} & X_{22} & \dots & X_{2,p-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & X_{n1} & X_{n2} & \dots & X_{n,p-1} \end{bmatrix}$$

$$\beta_{p \times 1} = \begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_{p-1} \end{bmatrix}, \quad \varepsilon_{n \times 1} = \begin{bmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \vdots \\ \varepsilon_n \end{bmatrix}$$

جدول Anova برای مدل رگرسیون خطی به صورت زیر می باشد:

Source of variation	ss	df	MS
Regression	$SSR = b'x'y - \frac{1}{n}x'Jy$	$p-1$	$MSR = \frac{SSR}{p-1}$
Error	$SSE = y'y - b'x'y$	$n-p$	$MSE = \frac{SSE}{n-p}$
Total	$SSTO = y'y - \frac{1}{n}y'Jy$	$n-1$	

که در آن J یک ماتریس $n \times n$ با مقادیر یک می باشد و b برآورد حداقل مربعات ضرایب رگرسیونی مدل می باشد یعنی $b = (x'x)^{-1}x'y$.

در مدل رگرسیون خطی، MSE یک برآوردگر ناریب σ^2 می باشد یعنی $E(MSE) = \sigma^2$. لذا درجه آزادی واریانس خطا برابر است با $n-p$.

باتوجه به این که تعداد مشاهدات برابر $n = 19$ و تعداد پارامترهای مدل برابر با $p = 3$ می باشد، بنابراین درجه آزادی واریانس خطا برابر است با $n-p = 16$.

۴۷. گزینه (۴)

۴۸

حجم نمونه روش نمونه گیری تصادفی ساده $\times$ اثر طرح = حجم نمونه مورد نیاز در نمونه گیری خوشه ای

$$N_{\text{cluster}} = 3 / 2 \times 1000 = 3200 \text{ نفر}$$

باتوجه به این که متوسط تعداد افراد هر خانوار ۴ نفر می باشد لذا $800 = \frac{3200}{4}$ خانوار برای برآورد شیوع بایستی مورد مطالعه قرار بگیرند.

۴۸. گزینه (۴)

برای کاهش و یا حذف نابرابری واریانس خطا یک روش استفاده از تبدیلات مناسب بر روی متغیر پاسخ y می باشد. مشکلی که با تبدیل متغیر y ممکن است رخ دهد اینست که منجر به ایجاد یک رابطه رگرسیونی نامناسب شود. زمانی که به وسیله یک مدل رگرسیونی مناسب تشخیص داده شد که واریانس خطاها ثابت نیستند، روش حداقل مربعات وزنی جایگزین مناسبی برای تبدیلات متغیر پاسخ می باشد که روشی براساس تعمیم مدل رگرسیون چندگانه است.

$$y_i = \beta_0 + \beta_1 X_{i1} + \beta_2 X_{i2} + \dots + \beta_{p-1} X_{i,p-1} + \varepsilon_i$$

Where:

$\beta_0, \beta_1, \dots, \beta_{p-1}$ are parameters

$X_{i1}, \dots, X_{i,p-1}$ are known constants

ε_i are independent $N(0, \sigma_i^2)$

$i = 1, \dots, n$

لذا در این مدل خطاها دارای واریانس ثابت σ^2 نمی باشند.

برای برآورد ضرایب رگرسیونی مدل تعمیم یافته می توان از برآوردگرهای مدل رگرسیون معمولی با واریانس ثابت σ^2 استفاده نمود. این برآوردگرها ناریب می باشند اما دیگر دارای واریانس مینیمم نیستند. برای به دست آوردن برآوردگرهای ناریب با کمترین واریانس، باید میزان اطلاعاتی را که هر کدام از مشاهدات در مورد تابع رگرسیون فراهم می کنند به صورت وزنی وارد مدل نمائیم به این صورت که مشاهدات با واریانس کمتر اطلاعات دقیق تری را در مورد تابع رگرسیون به دست می دهند.

روش های مختلفی برای وزن دهی مشاهدات وجود دارد که در فرم ساده آن اگر واریانس خطاها معلوم باشد $w_i = \frac{1}{\sigma_i^2}$

در نظر گرفته می شود لذا فرم تابع درست نمایی مشاهدات به صورت زیر می شود:

$$L(\beta) = \left[\prod_{i=1}^n \left(\frac{w_i}{\sqrt{\pi}} \right)^{\frac{1}{2}} \right] \exp \left[-\frac{1}{2} \sum_{i=1}^n w_i (y_i - \beta_0 - \beta_1 X_{i1} - \dots - \beta_{p-1} X_{i,p-1})^2 \right]$$

۴۹. گزینه (۲)

اگر متغیر تصادفی x دارای توزیع پواسن با پارامتر λ باشد آن گاه دارای تابع احتمال زیر است:

$$f_{\lambda}(x) = \frac{e^{-\lambda} \lambda^x}{x!}, \quad x = 0, 1, 2, \dots, \lambda > 0$$

اگر $E(x) = \lambda$ آن گاه $x \sim p(\lambda)$

لذا احتمال مشاهده متوسط تعداد تصادفات در طول یک هفته بیش تر است یعنی:

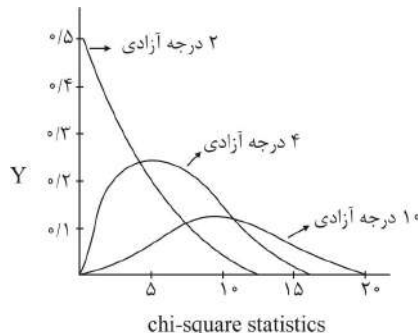
۴۹ $E(X) = \frac{3+1+\dots+2+1}{10} = \frac{20}{10} = 2$

۵۰. گزینه (۴)

به سؤال ۴۱ رجوع شود.

۵۱. گزینه (۱)

اگر متغیر تصادفی X دارای توزیع کای- دو با n درجه آزادی باشد یعنی $X \sim \chi^2_{(n)}$ ، آن گاه میانگین و واریانس آن برابر است با:



یکی از ویژگی های توزیع کای- دو اینست که با افزایش درجه آزادی شکل توزیع شبیه به توزیع نرمال می شود. بنابراین زمانی که تعداد سطرها یا ستون ها کاهش پیدا کند، درجه آزادی نیز کاهش می یابد که این امر باعث می شود شکل منحنی از توزیع نرمال فاصله داشته باشد به این معنی که پراکندگی توزیع افزایش می یابد.

شماره تماس با نمایندگی‌های فعال و رسمی آناهید گستر خلیلی

نماینده	شماره تماس	نماینده	شماره تماس
اراک (آقای حسینی)	۰۹۹۱۳۷۱۳۷۵۰	سیزوار (آقای شریفان)	۰۹۱۰۱۷۱۱۸۷۲
اردبیل (خانم عاصمی‌زاده)	۰۹۱۹۹۱۰۱۲۴۸	شاهرود (آقای سلمانی)	۰۹۹۱۳۷۱۳۷۴۱
ارومیه (آقای آرمیون)	۰۹۱۹۹۱۰۱۲۴۱	شهرکرد (خانم تقی‌پور)	۰۹۱۹۹۱۰۱۲۴۹
اسفراین (خانم اسماعیل‌زاده)	۰۹۱۹۶۳۲۱۸۵۲	شیراز (آقای فروردین/ خانم هوشمندی)	۰۹۱۹۵۷۳۰۱۵۲
اصفهان (آقای کیانی)	۰۹۱۹۵۷۳۰۱۵۰	فسا (خانم شعبانی)	۰۹۱۹۳۷۱۳۷۴۲
اهواز (آقای دکتر رضازاده)	۰۹۱۹۵۷۳۰۱۵۴	قزوین (خانم چگینی)	۰۹۱۹۵۷۳۰۱۴۹
ایلام (آقای دکتر بوچانی)	۰۹۱۹۰۹۸۰۳۱۸	کرج (آقای دکتر علیرضایی)	۰۹۱۹۷۷۸۱۹۴۷
بروجرد (آقای احمدی‌صدر)	۰۹۱۹۶۸۵۳۱۱۶	کرمان (خانم ضمانتی‌یار)	۰۹۱۳۱۹۸۱۲۷۰
بم (خانم سرحدی‌نژاد)	۰۹۹۱۳۷۱۳۷۴۹	کرمانشاه (آقای ابراهیمی)	۰۹۱۹۵۷۳۰۱۴۸
بیرجند (آقای دکتر بهروان)	۰۹۱۹۵۹۰۷۲۰۳	کوهدشت (خانم یاریان)	۰۹۱۹۵۳۷۱۸۹۰
تبریز (خانم عاصمی‌زاده)	۰۹۱۹۵۷۳۰۱۴۷	گرگان (آقای شاه‌علی)	۰۹۱۹۹۱۰۱۲۴۷
جهرم (آقای یاعلی جهرمی)	۰۹۰۱۳۷۳۷۸۹۶	مراغه (آقای صمدی)	۰۹۱۹۵۷۳۳۱۷۹
جیرفت (خانم دکتر محمدی)	۰۹۱۹۹۱۰۱۲۴۰	مشهد (خانم دکتر شجاعی)	۰۹۹۱۳۷۱۳۷۴۴
خرم‌آباد (خانم میر)	۰۹۱۹۲۷۰۵۸۷۸	میاندوآب (آقای صمدی)	۰۹۱۹۵۷۳۳۱۷۹
رشت (خانم تنها)	۰۹۱۹۵۷۳۰۱۵۳	نیشابور (آقای موسوی)	۰۹۱۹۶۳۵۰۷۶۸
رفسنجان (آقای یوسفی)	۰۹۱۹۶۸۲۹۲۸۰	هرمزگان (آقای زارعی)	۰۹۱۹۵۷۳۳۲۸۴
زاهدان (آقای مهمان‌دوست)	۰۹۱۹۹۱۰۱۲۴۵	همدان (آقای سوری)	۰۹۱۹۵۷۳۰۱۵۵
زنجان (خانم دکتر هوشیار)	۰۹۱۹۲۷۰۵۸۷۱	یاسوج (آقای مرتضوی)	۰۹۱۹۵۹۰۷۲۰۴
ساری (آقای دکتر اکبری)	۰۹۱۹۷۷۸۱۹۴۴	یزد (خانم آزاد)	۰۹۱۹۹۱۰۱۲۴۳

بانک کتاب ناهید



«هر کتابی، از هر انتشاراتی را از ما بخواهید»

✓ جامع‌ترین بانک کتاب

✓ تحویل روزانه

✓ ارسال به تمامی نقاط کشور

✓ سفارش کتاب به‌صورت تلفنی و آنلاین

www.NIBS.ir



کتاب دانشگاهی، فنی و مهندسی، علوم پزشکی، علوم انسانی، عمومی،

ادبی، مذهبی، کمک آموزشی، کودک و نوجوان و کتاب نفیس

فروشگاه: تهران - خیابان انقلاب - روبه‌روی درب اصلی دانشگاه تهران

پاساژ فروزنده - طبقه همکف - پلاک ۳۳۱

تلفن: ۰۲۱ - ۶۶۴۸۹۳۷۵ - ۰۲۱ - ۶۶۴۸۹۳۴۹