

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

دعای پیش از مطالعه

اَللّٰهُمَّ اَخْرِجْنِيْ مِنْ ظُلُمَاتِ الْوَهْمِ وَاكْرِمْنِيْ بِنُورِ الْفَهْمِ
اَللّٰهُمَّ افْتَحْ عَلَيْنَا اَبْوَابَ رَحْمَتِكَ وَاَنْشُرْ عَلَيْنَا خَزَائِنَ عُلُوْمِكَ
بِرَحْمَتِكَ يَا اَرْحَمَ الرَّاحِمِيْنَ

پروردگارا، خارج کن مرا از تاریکی های فکر و گرامی بدار به نور فهم
پروردگارا، بگشای بر ما در های 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.



درسنامه روش‌های آماری

(رگرسیون، روش‌های ناپارامتری، نمونه‌گیری و طرح آزمایشات)

ویژه کارشناسی ارشد رشته‌ی آمار زیستی

تألیف و گردآوری:

میثم اسلامی

سرشناسه

: اسلامی، میثم، ۱۳۶۹-

عنوان و نام پدیدآور

: AGK درسنامه روش‌های آماری (رگرسیون، روش‌های ناپارامتری، نمونه‌گیری و طرح آزمایشات): ویژه

مشخصات نشر

: تهران: آناهید گستر خلیلی، ۱۴۰۳.

مشخصات ظاهری

: ۱۴۴ص؛ مصور، جدول؛ ۲۲ × ۲۹ س.م.

شابک

: 978-622-4907-42-4

وضعیت فهرست‌نویسی

: فیپا

موضوع

: زیست‌سنجی -- راهنمای آموزشی -- آزمون‌ها و تمرین‌ها (عالی)

Biometry -- Study and teaching -- Examinations, questions, ect (Higher)

دانشگاه‌ها و مدارس عالی -- آزمون‌ها

Universities and colleges -- Examinations

رده‌بندی کنگره

: QH۲۲۳/۵

رده‌بندی دیویی

: ۵۷۴/۱۵۱۹۵

شماره کتابشناسی ملی

: ۹۸۶۰۹۰۶



انتشارات آناهید گستر خلیلی

نام: AGK درسنامه روش‌های آماری

(رگرسیون، روش‌های ناپارامتری، نمونه‌گیری و طرح آزمایشات)

(ویژه کارشناسی ارشد رشته‌ی آمار زیستی)

تألیف و گردآوری: میثم اسلامی

ناشر: آناهید گستر خلیلی

نوبت و سال چاپ: اول. ۱۴۰۳

شمارگان: ۱۰۰۰

چاپ و صحافی: گیلان

مدیر تولید: امیرحسین خلیلی

مدیر فنی و هنری: مریم آرده

تایپ و صفحه‌آرایی: آرزو راضی

بهاء: ۲۵۰۰۰۰ تومان

آموزشگاه/فروشگاه: ۰۲۱۶۶۵۶۸۶۲۱

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

مدیر فروش: ۰۹۱۹۱۶۸۵۹۷۰ (ارتباط از طریق پیامک)

 agkpublisher@gmail.com

 www.agkins.com

طلیحه سخن مدیر تولید:

سپردن کار به اهلش نشانه‌ی خردمندی است.

برای تولید بسته‌های آموزشی تمام تلاش خود را به کار گرفته‌ایم تا محصولی شایسته را تقدیم شما عزیزان نماییم.

مجموعه‌ای که پیش روی شماست براساس معتبرترین و به‌روزترین منابع کنکور و به قلم مجرب‌ترین اساتید تألیف شده است و مطالعه آن موجب صرفه‌جویی در زمان و هزینه شما خواهد شد. ارائه بالاترین آمار قبولی در کنکور کارشناسی ارشد و دکتری در دهه اخیر بیان‌گر توانایی این مجموعه در راستای اهداف آموزشی خود می‌باشد. اینک با اتکا به تجربه‌ای ۱۷ ساله می‌توانیم ادعا کنیم که مجموعه حاضر، بی‌شک از جمله کامل‌ترین و به‌روزترین بسته‌های آموزشی موجود در عرصه کنکور می‌باشد.

❧ برخی از ویژگی‌های بسته‌های آموزشی:

- منطبق با معتبرترین و به‌روزترین منابع کنکور
- گردآوری و تألیف توسط مجرب‌ترین اساتید کنکور
- تدریس مطالب از سطح پایه تا پیشرفته به بیان ساده و روان
- انتقال بهتر مفاهیم با استفاده از شکل، جدول و نمودار
- تألیف مبتنی بر نحوه‌ی طراحی سوالات کنکور
- مطالعه یک منبع منسجم به جای مطالعه چندین منبع مختلف
- منبعی خودآموز و بدون نیاز به مدرس

اینک علاوه بر مطالعه بسته‌های آموزشی استفاده از خدمات و محصولات زیر را به عنوان ضمانتی برای قبولی به شما داوطلبان عزیز توصیه می‌کنم؛

❧ آزمون‌های آزمایشی: خوش بود گر محک تجربه آید به میان

داوطلب شجاع کیست؟ داوطلبی است که خود را مورد سنجش (خود را با خود) و ارزیابی (خود را با دیگران) قرار دهد تا پیش از کنکور به نقاط ضعف خود پی ببرد و آن‌ها را تقویت نماید. آزمون‌های آزمایشی دارای یک برنامه‌ریزی منسجم است و به شما اولتیماتوم می‌دهد. یعنی شما باید در یک زمان معین مطالب معینی را مطالعه نمایید.

آزمون‌های آزمایشی مجموعه دارای سوالات استاندارد، پاسخنامه تشریحی و کارنامه تحلیلی می‌باشد. صادقانه شرکت در آزمون‌های آزمایشی را به شما عزیزان توصیه می‌کنم.

❧ مشاوره تخصصی و برنامه‌ریزی: زیر نظر حرفه‌ای‌ها درس بخوانید ...

در صورتی که نمی‌توانید با توجه به زمان آزادتان برنامه‌ی مطالعاتی استاندارد را برای خود تدوین کنید حتماً از افراد متخصص کمک بخواهید. افرادی که تجربیات خود را صادقانه در اختیار شما قرار می‌دهند و برنامه‌ریزی مناسب با شرایط شما را انجام می‌دهند.

📖 کتاب تست:

کوچک‌ترین واحد یادگیری در کنکور، تست است. شما می‌توانید برای این‌که سطح یادگیری خود را پس از مطالعه بسنجید از کتاب‌های تست مجموعه (تست کنکور) و ماطراحان (تست‌های تألیفی) استفاده کنید. این کتاب‌ها به‌صورت موضوعی طبقه‌بندی شده‌اند و می‌توانند نحوه‌ی طرح سوالات کنکور و سطح آن‌ها را به‌صورت واضح به شما نشان دهند.

📖 درسنامه‌های نکات داغ و کتاب‌های الگوریتم مناسب‌ترین منابع برای دوران مرور و جمع‌بندی:

بیش‌تر داوطلبان حین مطالعه نگران خلاصه‌نویسی و مرور آموخته‌های خود هستند و برای خلاصه‌نویسی زمان زیادی را صرف می‌کنند. اینک به شما مژده می‌دهیم که این‌بار زمان به نفع شما در حال گذر است و دیگر با استفاده از درسنامه‌های نکات داغ و کتاب‌های الگوریتم نیاز به خلاصه‌نویسی ندارید.

شما می‌توانید نظرات، انتقادات و پیشنهادهای سازنده خود را در مورد بسته‌های آموزشی به شماره‌تلفن همراه ۰۹۱۹۱۶۸۵۹۷۰ فقط و فقط پیامک نمایید.

مشاور و مدیر تولید

امیرحسین خلیلی

طلیعه سخن مؤلف:

مدت کوتاهی است که خداوند نعمت زندگی را به بنده عطا فرموده، ولی همین مدت کم ولی طولانی! که همچون تُنگ کوچک ولی عمیقی بود، دریافتیم که خوشبختی یعنی واقف بودن به این که هرچه داریم از رحمت خداست و هر آن چه نداریم از حکمت او، احساس خوشبختی در رسیدن به خواسته‌ها نبود، بلکه لذت بردن از داشته‌ها بود و حیف که همین چند صباح را در غفلت گذراندم.

همه این‌ها را گفتم تا بدانید همین که به این مرحله و این جایگاه رسیده‌اید یعنی قدم بزرگی در زندگی برداشته‌اید. از همین داشته‌ها لذت ببرید و با پشتوانه این حس خوشبختی و آرامش، به مسیر خود ادامه دهید. در این مجموعه شاهد چکیده‌ای خواهید بود از تمام کتب منبع، کتب مرتبط و همچنین تألیفات اینجانب که شامل درسنامه جامع روش‌های آماری، آموزش EasyFit و ریاضی برای علوم پزشکی است.

با بنده در ارتباط باشید و برای مطالعه درس روش‌های آماری، برنامه‌ریزی صحیح و علمی به کار گیرید. اگر بیش از چهار ماه برای مطالعه این درسنامه فرصت دارید، اصلاً شتاب‌زده عمل نکنید، ابتدا مباحث هر جلسه (از پیش تعیین شده) را مرور سریع کنید و پس از مطالعه، نگاهی به منابع اصلی داشته باشید و در انتها برای تسلط کافی به سراغ تست‌های کنکور یا مثال‌های درسنامه بروید.

میثم اسلامی

Meysam.islami@gmail.com

بخش اول: رگرسیون

فصل اول: رگرسیون خطی ساده	۱۰
فصل دوم: رگرسیون چندگانه	۲۷
فصل سوم: بررسی مانده‌ها	۳۴
فصل چهارم: مدل‌های پیچیده	۴۰

بخش دوم: روش‌های ناپارامتری

فصل اول: آماره‌های ترتیبی	۵۰
فصل دوم: آماره‌های رتبه‌ای	۵۹
فصل سوم: آزمون‌های معروف براساس رتبه	۶۱
فصل چهارم: همبستگی	۶۶
فصل پنجم: آزمون دوها	۷۰
فصل ششم: آزمون‌های نیکویی برازش	۷۲
فصل هفتم: توزیع تجربی و سه آزمون معروف	۷۶

بخش سوم: نمونه‌گیری

فصل اول: آشنایی با برخی مفاهیم پایه	۸۴
فصل دوم: نمونه‌گیری تصادفی ساده	۸۶
فصل سوم: نمونه‌گیری تصادفی با طبقه‌بندی	۹۳
فصل چهارم: نمونه‌گیری با احتمال متغیر	۱۰۰
فصل پنجم: نمونه‌گیری سیستماتیک	۱۰۴
فصل ششم: نمونه‌گیری خوشه‌ای	۱۰۸
فصل هفتم: برآورد نسبتی و رگرسیونی	۱۱۳

بخش چهارم: طرح آزمایشات

فصل اول: آزمایش‌های تک عاملی	۱۲۱
فصل دوم: بلوک‌های تصادفی شده	۱۲۸
فصل سوم: طرح عاملی با دو عامل و عام	۱۳۷
فصل چهارم: برخی طرح‌های آماری	۱۴۱

بخش اول

رگرسیون

راهنمای مطالعه این کتاب

در ابتدای کار کاملاً صریح اعلام می‌کنیم این درسنامه حاصل تجربه و به زبان شخص مؤلف نگارش شده است ولی بدیهی است که سؤالات چند گزینه‌ای، فرمول‌ها، قضایا و مفاهیم پایه از کتاب‌های منبع گردآوری شده است و تنها هنر مؤلف بیان ساده‌تر و طبقه‌بندی شده‌ی مطالب است. در هریک از بخش‌ها ابتدا درسنامه‌ای جامع و کافی بیان گردیده و انتظار می‌رود پس از مطالعه و پاسخ به سؤالات هر بخش، توانایی کسب نمره لازم را داشته باشید.

رگرسیون خطی ساده

در کتاب بانک سوالات درسنامه جامع آمار زیستی و کتاب آمار ناپارامتری و نمونه‌گیری، تألیفات اینجانب از انتشارات مؤسسه دکتر خلیلی به‌طور مفصل به بیان ضریب زیبایی همبستگی پرداختیم و باتوجه به جنس هریک از متغیرها به نحوه کاربست ضرایب همبستگی پیرسون، اسپیرمن و کندال اشاره شد. ضریب همبستگی بین دو متغیر همانند قد و وزن افراد تنها به وجود و عدم وجود رابطه می‌پردازد و نمی‌توان هریک را از روی دیگری پیش‌بینی کرد.

در این درس با استفاده از فنون ریاضی یک رابطه خطی بین دو متغیر را به یک مدل ریاضی تبدیل می‌کنیم و با استفاده از این مدل که سعی داریم بهترین مدل باشد، هریک از متغیرها را با استفاده از دیگری تقریب خواهیم زد.

Y متغیر وابسته $\longrightarrow$ متغیری است که قصد داریم برای آن پیش‌بینی به‌دست بیاوریم و اثر متغیر مستقل را بر روی آن می‌سنجیم.

x متغیر مستقل $\longrightarrow$ متغیری که اثر آن را بر روی متغیر وابسته می‌سنجیم و با استفاده از آن متغیر وابسته را پیش‌بینی می‌کنیم.

مثال ۱: برای پیش‌بینی میزان فروش فروشندگان براساس میزان سابقه فروش، ۱۵ فروشنده را انتخاب می‌کنیم.

$(x_i, Y_i) = (\text{میزان فروش فرد } i\text{-ام (میلیون در ماه)}, \text{میزان سابقه فرد } i\text{-ام (ماه)}) \longrightarrow \text{فرد } i\text{-ام}$

$(x_1, Y_1) \longrightarrow (15, 2/1) \longrightarrow \text{فرد اول}$

$(x_2, Y_2) \longrightarrow (15, 1/8) \longrightarrow \text{فرد دوم}$

$\vdots \quad \quad \quad \vdots$

$(x_{15}, Y_{15}) \longrightarrow (5, 5/5) \longrightarrow \text{فرد پانزدهم}$

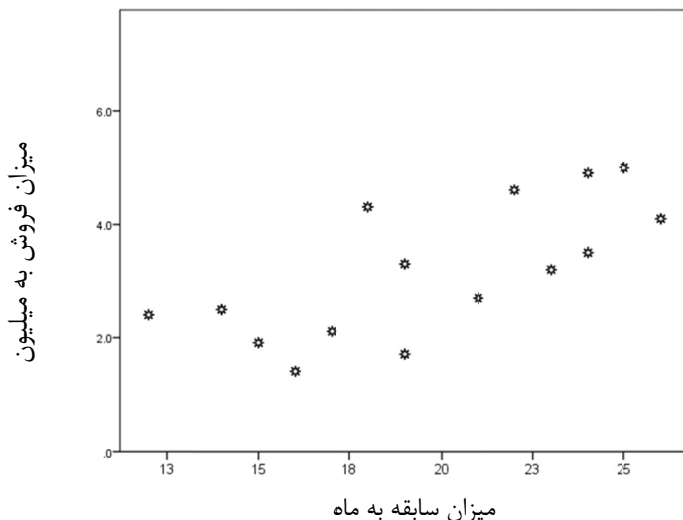
تمامی داده‌ها در جدول زیر قابل مشاهده است:

جدول ۱-۱

Sales	1.4	1.7	1.9	2.1	2.4	2.5	2.7	3.2	3.3	3.5	4.1	4.3	4.6	4.9	5.0
Record	16	19	15	17	12	14	21	23	19	24	26	18	22	24	25

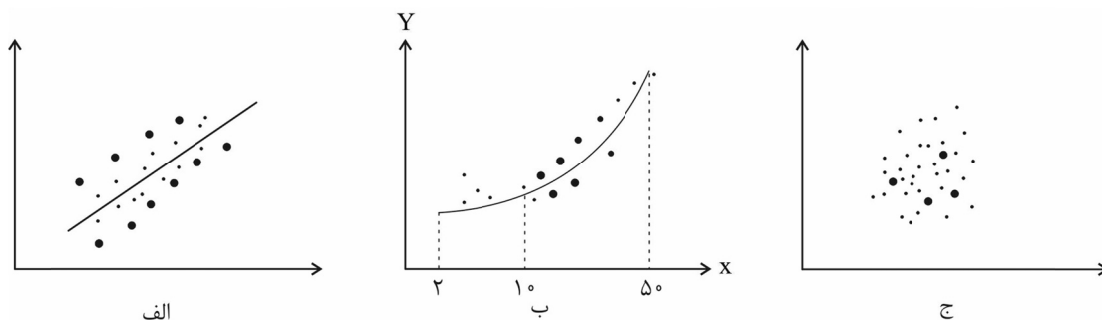
فصل اول: رگرسیون خطی ساده

نمودار پراکنش بین میزان سابقه (Record) و میزان فروش (Sales) افراد به صورت زیر است:



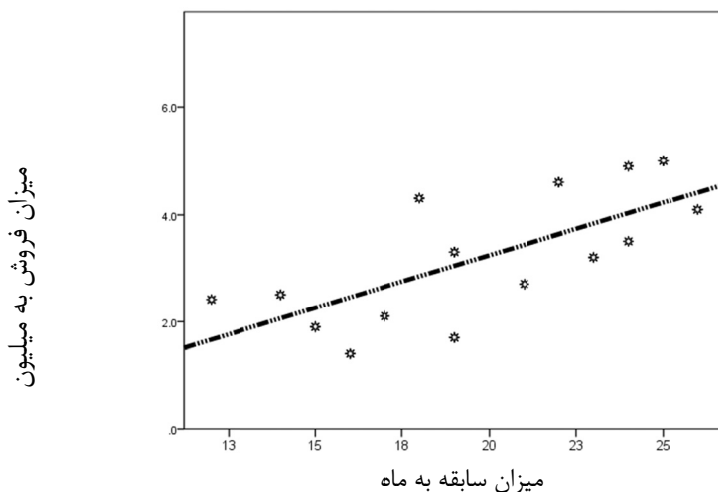
نگاره ۱-۱

در حالت کلی روابط زیادی بین متغیرها می‌تواند برقرار باشد، به سه نمودار زیر دقت کنید.



نگاره ۲-۱

برای انجام تجزیه و تحلیل رگرسیونی بایستی رابطه بین دو متغیر همچون نمودار «الف» در نگاره ۲-۱، خطی باشد. البته در نمودار «ب» می‌توان برای $2 \leq X \leq 10$ ، از فنون رگرسیونی استفاده شود زیرا رابطه‌ی خط مستقیم به دست آمده از مشاهدات وابسته به این فاصله ممکن است یک تابع کاملاً مناسب برای این فاصله باشد چون رابطه دو متغیر در این فاصله خطی است اما برای $2 \leq X \leq 50$ و سایر جاها ممکن است تابع نباشد. با توجه به نگاره ۱-۱ چون رابطه دو متغیر خطی است می‌توان یک خط از بین نقاط عبور داد و معادله آن را به صورت $Y = f(X)$ در نظر گرفت.



نگاره ۳-۱

رگرسیون چندگانه

ابتدا به مفاهیم زیر دقت کنید:

۱. رگرسیون خطی ساده: اگر یک متغیر مستقل و یک متغیر وابسته داشته باشیم.
۲. رگرسیون خطی چندگانه: اگر چند متغیر مستقل و یک متغیر وابسته داشته باشیم.
۳. رگرسیون خطی چند متغیره: اگر چند متغیر مستقل و چند متغیر وابسته داشته باشیم.

رگرسیون چندگانه^۱

پس در صورتی که با یک متغیر وابسته و چند متغیر مستقل سروکار داشته باشیم آن گاه برای برآزش مدل، بایستی مدل رگرسیون چندگانه را به کار گیریم.

اگر k متغیر مستقل داشته باشیم آن گاه در رگرسیون چندگانه مدل ما به شرح زیر است:

$$Y = \beta_1 + \beta_2 X_1 + \beta_3 X_2 + \dots + \beta_p X_k + \varepsilon$$

در این حالت که k تعداد متغیر مستقل و p پارامتر داریم که ضرایب رگرسیونی هستند و رابطه $p = k + 1$ بین آنها برقرار است. برای پیشگویی می توان همچون حالتی که با یک متغیر مستقل سروکار داریم عمل نمود، یعنی:

$$E(\hat{Y} | X_1 = x_1, \dots, X_k = x_k) = b_1 \times 1 + b_2 \times x_1 + b_3 \times x_2 + \dots + b_p \times x_k$$

که در آن $\hat{Y}$ پیش بینی مقدار y است. اگر مشاهده اول برابر واحد فرض شود، مقدار b_1 همان مقدار ثابت (در حالت خطی برابر عرض از مبدأ) است. اگر از هر متغیر مستقل n مشاهده داشته باشیم، معادله رگرسیون را می توان به صورت زیر نوشت:

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

ضرایب رگرسیونی به صورت زیر برآورد می شوند:

$$b = \hat{\beta} = (X^T X)^{-1} (X^T Y)$$

که در آن X^T ترانزاده ماتریس X است و $(X^T X)^{-1}$ معکوس ماتریس $(X^T X)$ است.

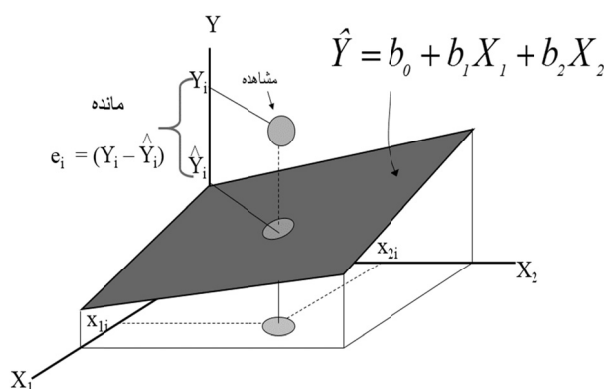
بخش اول: رگرسیون

جدول تحلیل واریانس در این مدل به صورت زیر است:

جدول تحلیل واریانس

منبع تغییرات	مجموع مربعات	درجه آزادی	میانگین مربعات	آماره F
رگرسیون (انحراف تبیین نشده)	SSR	$p-1$	MSR	$\frac{MSR}{MSE}$
مانده‌ها (انحراف تبیین شده)	SSE	$n-p$	MSE	
جمع کل	SST	$n-1$		

که در آن مجموع تغییرات با همان مفهوم قبل هستند، با این تفاوت که به جای یک متغیر، با چند متغیر مستقل سروکار داریم.



$SST =$ کل تغییرات در Y یا مجموع مجذورات کل

$$SST = \underline{Y}^T \left(\underline{1} - \frac{1}{n} \underline{1} \underline{1}^T \right) \underline{Y}$$

$SSE =$ منبع تغییرات ناشی از خطا

$$SSE = \underline{Y}^T \left(\underline{1} - \underline{X} (\underline{X} \underline{X}^T)^{-1} \underline{X}^T \right) \underline{Y}$$

$SSR =$ تغییرات تبیین شده

$$SSR = SST - SSE$$

نحوه‌ی تصمیم‌گیری همچون حالت قبلی است. اگر $F > F_{1-\alpha, p-1, n-p}$ باشد آن‌گاه فرض صفر (یعنی عدم رابطه خطی بین تمامی متغیرها) در سطح معناداری α پذیرفته نمی‌گردد.

خطای معیار برآورد (خطای استاندارد برآورد)

همان‌طور که به یاد دارید واریانس بیانگر میزان پراکندگی در داده‌ها می‌باشد. پس اگر واریانس خطای پیش‌بینی کم‌تر باشد آن‌گاه می‌توان نتیجه گرفت که مدل ما دقت بالایی در پیش‌بینی دارد. واریانس خطای پیش‌بینی به صورت:

$$MSE = \frac{SSE}{n-p} = \frac{\underline{Y}^T \left(\underline{1} - \underline{X} (\underline{X} \underline{X}^T)^{-1} \underline{X}^T \right) \underline{Y}}{n-p}$$

فاصله اطمینان $(1-\alpha) \times 100\%$ برای پیش‌بینی Y در یک نقطه جدید X_h

$$\left(\hat{Y} - t_{1-\frac{\alpha}{2}, n-p} \sqrt{MSE \times \left(X_h^T (X^T X)^{-1} X_h + 1 \right)}, \hat{Y} + t_{1-\frac{\alpha}{2}, n-p} \sqrt{MSE \times \left(X_h^T (X^T X)^{-1} X_h + 1 \right)} \right)$$

بررسی مانده‌ها

همان‌طور که به یاد دارید برای انجام آزمون‌های پارامتری همچون تی مستقل بایستی برخی فروض از جمله کمی بودن داده‌ها و نرمال بودن توزیع آن‌ها برقرار باشد و در غیر این صورت آزمون‌های مذکور و در حالت کلی استنباط‌های آماری از درجه اعتبار ساقط هستند.

اعتبار مدل و میزان شایستگی آن را به وسیله مانده‌ها مورد بررسی قرار می‌دهیم. در صورتی که n زوج مشاهده به صورت (x_i, Y_i) داشته باشیم آن گاه n مانده به صورت زیر داریم:

$$e_i = Y_i - \hat{Y}_i, \quad i = 1, 2, \dots, n$$

Y_i : مقدار مشاهده شده

$\hat{Y}_i$: مقدار پیش بینی شده برای Y_i براساس مدل

یک مدل رگرسیون معتبر است هرگاه:

(۱) خطاها مستقل از یکدیگر باشند.

(۲) خطاها دارای میانگین صفر و واریانس ثابت σ^2 باشند.

(۳) خطاها دارای توزیع نرمال باشند.

برای بررسی مانده‌ها روش‌های متعددی وجود دارد که چند مورد را بررسی خواهیم کرد:

(A) با استفاده از چند نکته تحلیلی

نکته ۳-۱: فرض سوم یا نرمال بودن خطاها برای آزمون‌های F و β ها موردنیاز است.

نکته ۳-۲: اگر مدل دارای پارامتر β باشد آن گاه $\sum_{i=1}^n e_i = 0$ است زیرا:

$$\sum \hat{Y}_i = \sum Y_i \Rightarrow \sum e_i = \sum (Y_i - \hat{Y}_i) = \sum Y_i - \sum \hat{Y}_i = 0$$

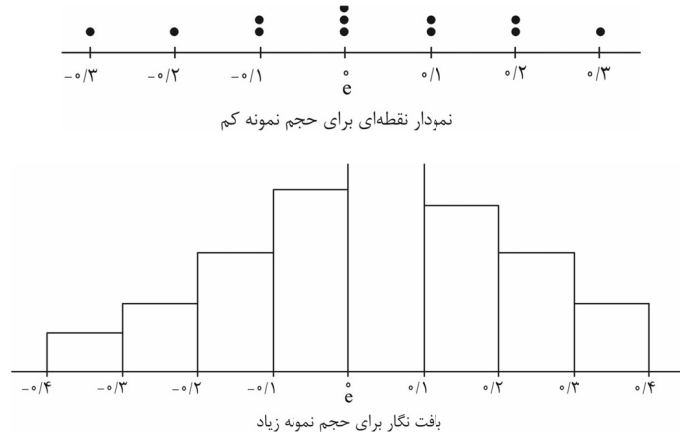
نکته ۳-۳: در رگرسیون چندگانه بردارهای e و $\hat{Y}$ ناهمبسته هستند یعنی برای هر i و j ($i \neq j$) داریم:

$$\text{Cov}(e_i, \hat{Y}_j) = 0$$

(B) نمودار نقطه‌ای

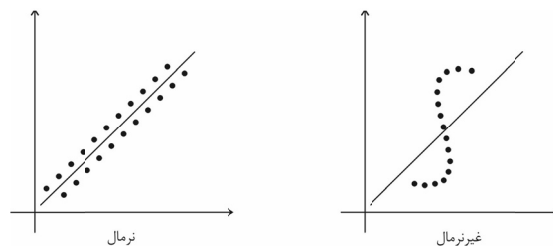
یکی دیگر از راه‌های بررسی مانده‌ها، نمودار نقطه‌ای و در حالتی که تعداد مشاهدات یا به عبارتی تعداد زوج داده‌ها زیاد باشد، نمودار هیستوگرام یا بافت نگار است.

اگر مدل برازش داده شده یک مدل رگرسیونی مناسب باشد آن‌گاه نمودار نقطه‌ای مانده‌ها دارای توزیع نرمال با میانگین صفر است که نباید در سمت چپ و راست نمودار نقاط زیادی یا به عبارتی داده دور افتاده مشاهده شود. نمودار نقطه‌ای زیر، بیانگر مناسب بودن مدل مرتبط با آن است:



(C) کاغذ احتمال

می‌توان روی کاغذهای احتمال معمولی، نمودار نرمال یا نمودار نیم نرمال از مانده‌ها را رسم کرد. در صورتی که نقاط، بیش‌تر حول خط $Y = x$ باشند آن‌گاه مانده‌ها دارای توزیع نرمال هستند. به نگاره‌های زیر دقت کنید.



(D) آزمون فرض

یکی از فروض اساسی در رگرسیون آن است که $\varepsilon_i \sim N(0, \sigma^2)$ باشد و طبق روش‌های آمار ریاضی می‌توان نشان داد که $\frac{\varepsilon_i}{\sigma} \sim N(0, 1)$ و اگر مدل رگرسیونی مناسب و صحیح باشد آن‌گاه میانگین توان دوم خطاها یا به عبارتی MSE برای σ^2 نااریب است:

$$E(MSE) = E \left(\frac{\sum_{i=1}^n (e_i - \bar{e})^2}{n-p} \right) = \frac{\sum e_i = 0}{n-p} E \left(\frac{\sum_{i=1}^n e_i^2}{n-p} \right) = \sigma^2$$

اما در صورتی این برآوردگر نااریب است که $\frac{\varepsilon_i}{\sigma} \sim N(0, 1)$ باشد. پس باید آزمون فرض روبه‌رو انجام شود.

$$\begin{cases} H_0: \frac{\varepsilon_i}{\sigma} \sim N(0, 1) \\ H_1: \frac{\varepsilon_i}{\sigma} \not\sim N(0, 1) \end{cases}$$

برای انجام این آزمون باید مقادیر $\frac{e_i}{\sqrt{MSE}}$ را برای $i = 1, 2, \dots, n$ محاسبه کرد و برای مقایسه حاصل، نمودار نقطه‌ای ترسیم کرد. در صورتی که ۹۵٪ از نقاط در فاصله $(-1/96, +1/96)$ قرار بگیرد فرض صفر پذیرفته می‌شود.

توجه: معمولاً مانده‌ها مستقل نیستند و دارای واریانس مشترک نیستند و به ماتریس X وابسته هستند و در نتیجه استفاده از این روش ممکن است گمراه کننده بوده و روش کاملاً دقیقی نیست.

مدل‌های پیچیده

مدل‌های به‌ظاهر غیر خطی

در ادامه برخی مدل‌های مهم که به‌ظاهر غیرخطی هستند ولی ذاتاً خطی هستند را بیان می‌کنیم که همیشه مورد توجه طراح محترم سؤالات بوده و هست و خواهد بود و آن‌ها را با یک تبدیل مناسب می‌توان به ذات خود برگرداند.

الف) مدل ضربی $\xleftarrow{\text{تبدیل مناسب}}$ لگاریتم

$$\ln Y = \ln \alpha + \beta \ln x_1 + \gamma \ln x_2 + \delta \ln x_3 + \ln \varepsilon \quad \xleftarrow{\ln} \quad Y = \alpha x_1^\beta x_2^\gamma x_3^\delta \varepsilon$$

$$W = \beta_0 + \beta_1 Z_1 + \beta_2 Z_2 + \beta_3 Z_3 + \varepsilon^*$$

ε : یک متغیر تصادفی

$\alpha, \beta, \gamma, \delta$: پارامترهای نامعلوم

$$\varepsilon^* = \ln \varepsilon \sim N(0, \sigma^2)$$

ب) مدل نمایی $\xleftarrow{\text{تبدیل مناسب}}$ لگاریتم

$$\ln Y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \ln \varepsilon \quad \xleftarrow{\ln} \quad Y = e^{\beta_0 + \beta_1 x_1 + \beta_2 x_2} \varepsilon$$

$$\downarrow$$

$$W = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \varepsilon^*$$

ج) مدل معکوس $\xleftarrow{\text{تبدیل مناسب}}$ معکوس طرفین

$$\frac{1}{Y} = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \varepsilon \quad \xleftarrow{\quad} \quad Y = \frac{1}{\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \varepsilon}$$

$$\downarrow$$

$$W = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \varepsilon^*$$

(د) مدل نمایی و پیچیده‌تر $\xleftarrow{\text{تبدیل مناسب}}$ معکوس طرفین و سپس لگاریتم

$$\ln\left(\frac{1}{Y} - 1\right) = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \varepsilon \quad \leftarrow \quad Y = 1 + e^{\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \varepsilon} \quad \leftarrow \quad Y = \frac{1}{1 + e^{\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \varepsilon}}$$

$$W = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \varepsilon^*$$

نکته: توجه شود که برای بررسی مانده‌ها بایستی متغیر تبدیل یافته را بررسی کنیم و عبارت خطا یعنی ε^* دارای توزیع $N(0, \sigma^2 I_n)$ خواهد بود و ضرایب از روش کم‌ترین توان‌های دوم خطا برآورده می‌شود.

متغیرهای ظاهری در رگرسیون چندگانه

در مدل‌هایی که تا به حال بیان شده، متغیرهای مستقل مقادیر مختلفی را می‌تواند به خود بگیرد و به عبارتی فضای نمونه‌ای آن‌ها پیوسته است. اما گاهی متغیر مستقل به صورت گسسته است و تنها چند عضو محدود را در فضای نمونه‌ای خود دارد. تصور کنید در مثال وزن و قد افراد، بخواهیم وضعیت ابتلا به فشارخون را به عنوان یک متغیر مستقل وارد مدل کنیم. بدیهی است که این متغیر تنها دو مقدار به خود می‌گیرد.

$$Z = \begin{cases} 1 & \text{مبتلا به فشار خون} \\ 0 & \text{عدم ابتلا به فشار خون} \end{cases}$$

و مدل رگرسیونی به صورت $Y = \beta_0 + x\beta_1 + z\beta_2 + \varepsilon$ خواهد بود.

Z : وضعیت فشارخون

x : مثلاً وزن

Y : مثلاً قد

اما گاهی متغیر ظاهری بیش از دو انتخاب را در بردارد. به عنوان مثال بخواهیم رنگ چشم افراد را هم به صورت متغیر ظاهری به مدل اضافه کنیم و بدانیم همه افراد دارای چشم‌هایی با رنگ مشکی، قهوه‌ای و سبز هستند در این صورت سه متغیر ظاهری به صورت زیر خواهیم داشت:

$$Z_1 = \begin{cases} 1 & \text{رنگ چشم فرد مشکی است} \\ 0 & \text{رنگ چشم فرد مشکی نیست} \end{cases}$$

$$Z_2 = \begin{cases} 1 & \text{رنگ چشم فرد قهوه‌ای است} \\ 0 & \text{رنگ چشم فرد قهوه‌ای نیست} \end{cases}$$

$$Z_3 = \begin{cases} 1 & \text{رنگ چشم فرد آبی است} \\ 0 & \text{رنگ چشم فرد آبی نیست} \end{cases}$$

اما نیاز نیست هر سه متغیر را وارد مدل کنیم زیرا هرگاه $Z_1 = 0$ و $Z_2 = 0$ آن‌گاه بدون شک $Z_3 = 1$ خواهد بود، به عبارتی متغیر Z_3 به مقدار Z_1 و Z_2 وابسته است و نیازی نیست Z_3 را بررسی کنیم و مدل را به صورت زیر تعریف می‌کنیم.

$$Y = \beta_0 + x\beta_1 + Z_1\beta_2 + Z_2\beta_3$$

پس می‌توان به یک نکته تپل دست یافت که برای r سطح یا متغیر ظاهری، بایستی $(r-1)$ متغیر ظاهری وارد مدل کنیم البته اگر $r=1$ باشد همان یک متغیر را وارد می‌کنیم.

گزینش بهترین معادله رگرسیون

در حالت کلی هرچه متغیرهای مستقل بیش‌تری در مدل حضور داشته باشد، دقت پیشگویی مدل بالا خواهد بود ولی طبیعتاً به هزینه و پیچیدگی مدل افزوده خواهد شد.

می‌توان تمام مدل‌های رگرسیونی ممکن را برازش داد، که اگر r متغیر مستقل داشته باشیم تعداد کل مدل‌ها 2^r خواهد بود، و از طریق ملاک‌های زیر، مدل‌ها را ارزیابی نمود و بهترین را انتخاب کرد:

الف) ضریب تعیین یا R_p^2 (که $P = k + 1$ و k متغیر مستقل است)

ب) $MSEP$ یا R_p^2

ج) آماره مالوس یا C_p

بخش دوم

روش‌های ناپارامتری

آماره‌های ترتیبی

در این فصل به مفهوم آماره ترتیبی و کاربرست آن می‌پردازیم. همان‌طور که می‌دانید متغیر تصادفی با حرف بزرگ و معمولاً با X نشان می‌دهیم و نمونه تصادفی را به صورت $X_1, \dots, X_n$ و مقدار مشاهده شده آن‌ها را به صورت $x_1, \dots, x_n$ نشان خواهیم داد.

می‌دانید که هر صفت یا ویژگی مورد بررسی را با متغیر تصادفی X نشان می‌دهیم، در حالت پارامتریک هر ویژگی در جامعه یا X ، از یک توزیع مانند $F(x)$ با تابع چگالی $f(x)$ است. متغیرهای تصادفی $X_1, X_2, \dots, X_n$ را یک نمونه تصادفی از توزیع $F(x)$ می‌نامند، هرگاه این متغیرها دارای توزیع $F(x)$ باشند و حتماً مستقل باشند.

هر تابعی از نمونه تصادفی را که در آن پارامتر نباشد آماره گویند. البته در اکثر کتاب‌ها در تعریف آماره از اصطلاح «پارامتر مجهول» استفاده می‌کنند که این اصطلاح اشتباه است زیرا در تعریف پارامتر داریم: هر ویژگی مجهول از جامعه را پارامتر گوئیم، پس پارامتر به نوبه خود مجهول است. به عنوان مثال $\bar{X}$ یک آماره بوده و برآوردگر میانگین جامعه است ولی $\mu + \bar{x}$ آماره نیست چون پارامتر μ در آن وجود دارد.

آماره‌های ترتیبی

دنباله نمونه تصادفی را به این دلیل با حروف بزرگ به صورت $X_1, X_2, \dots, X_n$ نشان می‌دادیم که مجهول بودند. حال اگر آن‌ها را به صورت غیرنزولی مرتب کنیم و به صورت:

$$Y_1 \leq Y_2 \leq \dots \leq Y_n$$

یا

$$X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$$

نشان دهیم آن‌گاه Y_1 یا $X_{(1)}$ کوچک‌ترین و Y_n یا $X_{(n)}$ بزرگ‌ترین آماره مرتب هستند. به عنوان مثال اگر یک تاس را ۵ بار پرتاب کنیم و یافته‌های ما $x_1 = 6, x_2 = 3, x_3 = 4, x_4 = 5, x_5 = 2$ باشند، آن‌گاه $Y_5 = 6$ است.

نکته ۱-۱: اگر X_i ها $\frac{\text{گسسته}}{\text{پیوسته}}$ باشد آن‌گاه Y_i ها نیز $\frac{\text{گسسته}}{\text{پیوسته}}$ هستند.

آماره‌های رتبه‌ای

فرض کنید متغیر تصادفی X را سن افراد تعریف کنیم و با زحمت و اصرار سن ۵ تن از بانوان اقوام را یادداشت کرده‌ایم و نتایج به‌صورت زیر بوده است.

$$x_1 = 32, x_2 = 21, x_3 = 44, x_4 = 25, x_5 = 24$$

در این صورت:

$$Y_1 = 21, Y_2 = 24, Y_3 = 25, Y_4 = 32, Y_5 = 44$$

اگر رتبه فرد i ام را با R_i نشان دهیم آن‌گاه:

شماره فرد (i)	نفر اول	نفر دوم	نفر سوم	نفر چهارم	نفر پنجم
رتبه فرد (R_i)	۴	۱	۵	۳	۲
	R_1	R_2	R_3	R_4	R_5

$R = (R_1, \dots, R_5)$ را بردار رتبه‌ها گوئیم و قصد داریم قضایا و نکاتی درباره توزیع تکی و توأم R_i ها و بردار R بیان کنیم.

قضیه ۱: اگر نمونه تصادفی $X_1, \dots, X_n$ را از توزیع پیوسته $F(x)$ با چگالی $f(x)$ برداشت کنیم. آن‌گاه:

$$P(X_1 < X_2 < \dots < X_n) = \frac{1}{n!}$$

توجه: اگر یک نمونه سه تایی داشته باشیم، منظور از قضیه ۱ آن است که:

$$P(X_1 < X_2 < X_3) = P(X_2 < X_3 < X_1) = \dots = P(X_3 < X_2 < X_1) = \frac{1}{6}$$

آزمون‌های معروف براساس رتبه

در این فصل به بیان برخی آزمون‌های معروف می‌پردازیم که «رتبه» نقش مهمی در آن‌ها دارد.

(I) آزمون جمعی رتبه‌ای ویلکاکسون

دو نمونه تصادفی زیر را در نظر بگیرید.

متغیرها	X	Y
نمونه‌ها	X_1	Y_1
	$\vdots$	$\vdots$
	X_n	Y_n
توزیع‌ها	$F_X(x)$	$G_Y(y)$

می‌توان رابطه $Y \stackrel{d}{=} X + C$ را در نظر گرفت و سه آزمون زیر را انجام داد.

$$\text{I)} \begin{cases} H_0: C = 0 \\ H_1: C > 0 \end{cases} \quad \text{II)} \begin{cases} H_0: C = 0 \\ H_1: C < 0 \end{cases} \quad \text{III)} \begin{cases} H_0: C = 0 \\ H_1: C \neq 0 \end{cases}$$

که در هریک از سه آزمون در صورتی که فرض صفر پذیرفته شود می‌توان به یکسانی دو گروه یا بی‌اثری مورد پژوهش، باتوجه به نمونه مشاهده شده در سطح معناداری α رسید.

به‌عنوان مثال در آزمون شماره (II) اگر فرض صفر پذیرفته نشود و رد شود، یعنی $C < 0$ خواهد بود و به‌عبارتی $Y > X$ ها خواهد بود.

$$i = 1, \dots, m, \text{ ها } X_i \xrightarrow{\text{رتبه } X_i} R_i \quad j = 1, \dots, n, \text{ ها } Y_j \xrightarrow{\text{رتبه } Y_j} S_j$$

اگر $N = m + n$ تعریف کنیم، N یعنی تعداد کل داده‌ها، آن‌گاه مجموع کل رتبه‌ها به‌صورت زیر است:

$$\sum_{i=1}^m R_i + \sum_{j=m+1}^n S_j = \sum_{k=1}^N k = \frac{N(N+1)}{2}$$

$$W_S = \sum_{j=1}^n S_j, \quad W_R = \sum_{i=1}^m R_i$$

آماره‌های جمعی رتبه‌ای و ویلکاکسون به‌صورت زیر هستند.

برای انجام آزمون I ناحیه بحرانی به‌صورت $\{W_S \geq k\}$ ، برای آزمون II به‌صورت $\{W_S \leq k\}$ و برای آزمون III به‌صورت زیر است:

$$P\text{-Value} = 2 \min\{P(W_S \geq k), P(W_S \leq k)\}$$

شدت رابطه بین دو متغیر تصادفی را و به عبارتی همبستگی بین دو متغیر تصادفی را با ضریب همبستگی می‌توان محاسبه کرد. با توجه به جنس و نوع متغیرهای تصادفی می‌توان نسبت به محاسبه آن اقدام کرد.

I) ضریب همبستگی پیرسون (ساده)

دو متغیر تصادفی X و Y را با میانگین‌های μ_X و μ_Y و انحراف معیارهای σ_X و σ_Y و کوواریانس σ_{XY} باشد آن‌گاه میزان شدت رابطه خطی بین دو متغیر

$$\rho = P(X, Y) = \text{Corr}(X, Y) = \frac{\text{Cov}(X, Y)}{\sqrt{\text{Var}(X)\text{Var}(Y)}} = \frac{\sigma_{XY}}{\sigma_X \sigma_Y}$$

نکته ۴-۱: اگر a و b و c و d اعداد ثابتی باشند به طوری که $ac > 0$ آن‌گاه:

$$\rho(aX + b, cY + d) = \rho(X, Y)$$

نکته ۴-۲: همواره داریم:

$$-1 \leq \rho \leq +1$$

نکته ۴-۳:

$$\rho = \pm 1 \Leftrightarrow \frac{X - \mu_X}{\sigma_X} = \mp \frac{Y - \mu_Y}{\sigma_Y}$$

نکته ۴-۴: اگر $Z = \frac{X - \mu_X}{\sigma_X} \pm \frac{Y - \mu_Y}{\sigma_Y}$ آن‌گاه:

$$\begin{cases} E(Z) = 0 \\ \text{Var}(Z) = 2 \pm 2\rho \geq 0 \end{cases}$$

نکته ۴-۵:

$$r = \frac{S_{XY}}{S_X S_Y} = \frac{\overline{XY} - \bar{X}\bar{Y}}{\sqrt{\bar{X}^2 - \bar{X}^2} \sqrt{\bar{Y}^2 - \bar{Y}^2}}$$

آزمون دوها

دوها و کاربرهای آن

فرض کنید اعضای تیم استقلال را با x و اعضای تیم پیروزی را با y نشان دهیم. فرض کنید مربی وقت تیم ملی آقای کیروش افراد را به صورت زیر از چپ به راست صدا زده باشد:

$\underline{xx} \underline{yyy} \underline{x} \underline{yy} \underline{x} \underline{y} \underline{xxxx}$

آن دسته از اعضای هم نوع را که پشت سرهم آمده اند رایک یا دو گردش می نامند. شماره تمام دوها را با R و تعداد دوهایی از نوع x را با R_1 و تعداد دوهایی از نوع y را با R_2 نشان می دهیم که

$$R = R_1 + R_2$$

است. در این مثال $R_1 = 4$ و $R_2 = 3$ است و $R = 7$ خواهد بود.

$$\mu_R = E(R) = 1 + \frac{2n_1n_2}{(n_1 + n_2)} \quad \sigma_R^2 = \frac{(\mu_R - 1)(\mu_R - 2)}{n_1 + n_2 - 1} \quad \frac{R - \mu_R}{\sigma_R} \xrightarrow{d} Z \sim N(0, 1)$$

توجه داشته باشید چون R یک متغیر گسسته است بایستی تصحیح پیوستگی را (عدد $\frac{1}{2}$) در نظر گرفت. اگر r مقدار مشاهده شده

برای R باشد آن گاه:

$$\text{if } r < \mu_R \Rightarrow P\text{-Value} = P\left(R \leq \left(r + \frac{1}{2}\right)\right) \quad \text{if } r > \mu_R \Rightarrow P\text{-Value} = P\left(R \geq \left(r - \frac{1}{2}\right)\right)$$

مثال: در مثال ارائه شده می خواهیم بدانیم آیا جناب کیروش افراد را به صورت تصادفی صدا زده است؟ فرض تصادفی بودن را با میزان $\alpha = 0.05$ بسنجید.

$$n_1 = 8, n_2 = 6, r = 7 \quad \mu_R = 1 + \frac{2(8)(6)}{(6+8)} = 9 \quad \sigma_R = \sqrt{\frac{(9-1)(9-2)}{8+6-1}} = \sqrt{\frac{56}{13}} \approx 2.07$$

$$\xrightarrow{r < R} P(R \leq 7.5) = P\left(\frac{R - \mu_R}{\sigma_R} \leq \frac{7.5 - 9}{2.07}\right) = P(2 \leq -0.72) \approx 0.2358$$

چون $0.23 > 0.05$ پس فرض صفر یعنی تصادفی بودن را می پذیریم.

آزمون‌های نیکویی برازش

در این بخش می‌خواهیم آزمون کنیم آیا متغیر تصادفی X دارای توزیع $F_X(x)$ است؟ یا به بیان دیگر آیا توزیع $F(x)$ برازنده بر X است؟

(I) آزمون χ^2 برای برازندگی

فرض کنید با انجام O_i بار نتیجه C_i را مشاهده کنیم ولی فراوان مورد انتظار ما e_i باشد آن‌گاه $e_i = np_i$ و داریم:

$$\sum_{i=1}^k O_i = \sum_{i=1}^k e_i = n$$

که در آن k تعداد پیشامدها است. فرض H_0 بیان می‌کند آیا آزمایش تصادفی دارای مدل احتمال زیر است؟

مجموعه	C_1	C_2	C_3	$\dots$	C_k
احتمال	P_1	P_2	P_3	$\dots$	P_k

آماره این آزمون به‌صورت زیر است:

$$V_k = \sum_{i=1}^n \frac{(O_i - e_i)^2}{e_i}$$

اگر n بزرگ باشد، تحت فرض صفر آن‌گاه V_k دارای توزیع χ^2_{k-1} درجه آزادی است و آماره پیرسون نامیده می‌شود.

$$\text{Reject } H_0 \text{ if } V_k > \chi^2_{1-\alpha, k-1}$$

توجه: فراوانی هر دسته نباید کمتر از ۵ باشد وگرنه مجبوریم با سایر رده‌ها ادغامش کنیم.

نکته ۱-۶: V_k به‌طور مجانبی با $\chi^2_{(k-1)}$ هم‌توزیع است.

در چهار راه ولیعصر یک دستگاه کار گذاشتیم و تعداد تصادف‌ها در ۵۰ هفته را ثبت کردیم و نتایج به‌صورت جدول زیر خلاصه کرده‌ایم.

تعداد تصادف‌ها در هفته	۰	۱	۲	≥ 3
فراوانی هفته‌ها	۳۲	۱۲	۶	۰

توزیع تجربی و سه آزمون معروف

می‌دانید که $P(X \leq x) = F_X(x)$ و منظور از $X \sim F_X(x)$ آن است که X دارای توزیع $F_X(x)$ است و در این صورت انتظار داریم که $[nF_X(x)]$ از X_i ها از x کوچکتر باشند.

$$\hat{F}_X(X) = \frac{\text{تعداد } X_i \text{ های کوچکتر یا مساوی } X}{n}$$

این برآورد تنها به n بستگی دارد و با $F_n(x)$ نشان داده می‌شود و تابع توزیع تجربی نام دارد.

توزیع تجربی با استفاده از متغیرهای برنولی

اگر $X \sim F_X(x)$ ، آن‌گاه برای هر عدد حقیقی ثابت x ، متغیر تصادفی برنولی $I_{X_i}(x)$ به صورت زیر تعریف می‌شود:

$$I_{X_i}(x) = \begin{cases} 1 & X_i \leq x \\ 0 & X_i > x \end{cases}$$

$$P(I_{X_i}(x) = 1) = P(X_i \leq x) = F_{X_i}(x)$$

و

$$E(I_{X_i}(x)) = F_X(x)$$

$$\text{Var}(I_{X_i}(x)) = F(x)[1 - F(x)]$$

حال توزیع تجربی براساس متغیر تصادفی برنولی به صورت زیر است:

$$F_n(x) = \frac{1}{n} \sum_{i=1}^n I_{X_i}(x)$$

$$E(F_n(x)) = F(x)$$

$$\text{Var}(F_n(x)) = \frac{F(x)[1 - F(x)]}{n}$$

بخش سوم

نمونه گیری

آشنایی با برخی مفاهیم پایه

جامعه: مجموعه‌ای است، از تمام مشاهدات ممکن، که می‌تواند با تکرار یک آزمایش رخ بدهد.

جامعه آماری: مجموعه‌ای از افراد یا اشیاء یا به بیان ساده‌تر جامعه‌ای که اعضای آن حداقل در یک ویژگی، مشترک باشند.

سرشماری: روشی است که طی آن تمام اعضای جامعه آماری مورد بررسی قرار می‌گیرد.

نمونه: در این روش تنها بخشی از جامعه آماری مورد بررسی قرار می‌گیرد و در وهله اول بایستی تصادفی انتخاب شود و در وهله بعد بایستی این نمونه تصادفی معرف جامعه‌ی آماری باشد تا بتوان نتایج حاصل آن را به جامعه تعمیم داد.

علم آمار یا به بیان دیگر بررسی‌های نمونه‌ای به دو بخش توصیفی و استنباطی تقسیم می‌شود، که در بخش توصیفی ویژگی‌های مشاهدات را بررسی می‌کنیم و در بخش آمار استنباطی، استنباط‌های لازم مرتبط با جامعه را از نمونه استخراج می‌کنیم.

مزایای نمونه‌گیری در مقایسه با سرشماری

۱. کاهش هزینه
۲. افزایش سرعت در تلخیص داده‌ها
۳. قدرت عمر بیش‌تر (آموزش تمام آمارگیران امکان‌پذیر نیست).
۴. صحت عمل بیش‌تر (افرادی که نمونه را بررسی می‌کنند حرفه‌ای‌تر هستند).
۵. حفظ واحدهای جامعه (تمام اعضا را مورد آزمایش قرار نمی‌دهیم).

مراحل یک بررسی نمونه‌ای

۱. مشخص نمودن هدف از انجام کار (اهداف بررسی)
۲. مشخص نمودن جامعه مورد نمونه‌گیری
۳. جمع‌آوری داده‌ها
۴. مشخص نمودن میزان عدم دقت (تعیین درجه دقت مطلوب)
۵. تعیین روش اندازه‌گیری
۶. فهرست نمودن واحدهای نمونه (چارچوب)
۷. تعیین حجم نمونه
۸. پیش آزمون (اصلاح اشتباهات ممکن در بخش‌های قبل و اتخاذ تصمیم متناسب با آن)
۹. آموزش آمارگیران

بانک سوالات

جمع آوری سوالات کنکور کاردانی به کارشناسی
کارشناسی ارشد و دکتری به صورت فصل بندی شده

سوالات تألیفی

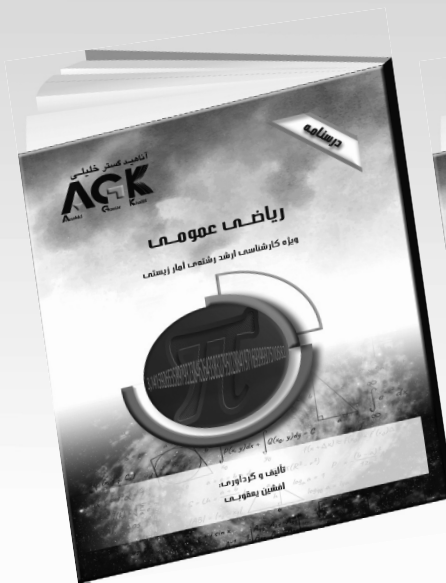
تألیف سوالات مشابه کنکور

جامع

حاوی تمامی مطالب و نکات لازم
برای کنکور براساس منابع

الگوریتم

چکیده‌ی تمامی مطالب و نکات لازم
برای کنکور براساس منابع



دریافت نمونه‌ی کتاب به صورت رایگان



www.agkins.com

جامع

حاوی تمامی مطالب و نکات لازم
برای کنکور براساس منابع

بانک سوالات

جمع آوری سوالات کنکور کاردانی به کارشناسی
کارشناسی ارشد و دکتری به صورت فصل بندی شده

الگوریتم

چکیده‌ی تمامی مطالب و نکات لازم
برای کنکور براساس منابع

سوالات تألیفی

تألیف سوالات مشابه کنکور



دریافت نمونه‌ی کتاب به صورت رایگان



www.agkins.com

شماره تماس با نمایندگی‌های فعال و رسمی آناهید گستر خلیلی

شماره تماس	نمایندگی	شماره تماس	نمایندگی
۰۹۱۰۱۷۱۱۸۷۲	سبزوار (آقای شریفان)	۰۹۹۱۳۷۱۳۷۵۰	اراک (آقای حسینی)
۰۹۹۱۳۷۱۳۷۴۱	شاهرود (آقای سلمانی)	۰۹۱۹۹۱۰۱۲۴۸	اردبیل (خانم عاصمی‌زاده)
۰۹۱۹۹۱۰۱۲۴۹	شهرکرد (خانم تقی‌پور)	۰۹۱۹۹۱۰۱۲۴۱	ارومیه (آقای آرمیون)
۰۹۱۹۵۷۳۰۱۵۲	شیراز (آقای فروردین / خانم هوشمندی)	۰۹۱۹۶۳۲۱۸۵۲	اسفراین (خانم اسماعیل‌زاده)
۰۹۱۹۳۷۱۳۷۴۲	فسا (خانم شعبانی)	۰۹۱۹۵۷۳۰۱۵۰	اصفهان (آقای کیانی)
۰۹۱۹۵۷۳۰۱۴۹	قزوین (خانم چگینی)	۰۹۱۹۵۷۳۰۱۵۴	اهواز (آقای دکتر رضازاده)
۰۹۱۹۷۷۸۱۹۴۷۱	کرج (آقای دکتر علیرضاپور)	۰۹۱۹۰۹۸۰۳۱۸	ایلام (آقای دکتر بوچانی)
۰۹۱۳۱۹۸۱۲۷۰	کرمان (خانم ضمانتی‌یار)	۰۹۱۹۶۸۵۳۱۱۶	بروجرد (آقای احمدی‌صدر)
۰۹۱۹۵۷۳۰۱۴۸	کرمانشاه (آقای ابراهیمی)	۰۹۹۱۳۷۱۳۷۴۹	بم (خانم سرحدی‌نژاد)
۰۹۱۹۵۳۷۱۸۹۰	کوهدشت (خانم یاریان)	۰۹۱۹۵۹۰۷۲۰۳	بیرجند (آقای دکتر بهروان)
۰۹۱۹۹۱۰۱۲۴۷	گرگان (آقای شاه‌علی)	۰۹۱۹۵۷۳۰۱۴۷	تبریز (خانم عاصمی‌زاده)
۰۹۱۹۵۷۳۳۱۷۹	مراغه (آقای صمدی)	۰۹۰۱۳۷۳۷۸۹۶	جهرم (آقای یاعلی جهرمی)
۰۹۹۱۳۷۱۳۷۴۴	مشهد (خانم دکتر شجاعی)	۰۹۱۹۹۱۰۱۲۴۰	جیرفت (خانم دکتر محمدی)
۰۹۱۹۵۷۳۳۱۷۹	میاندوآب (آقای صمدی)	۰۹۱۹۲۷۰۵۸۷۸	خرم‌آباد (خانم میر)
۰۹۱۹۶۳۵۰۷۶۸	نیشابور (آقای موسوی)	۰۹۱۹۵۷۳۰۱۵۳	رشت (خانم تنها)
۰۹۱۹۵۷۳۳۲۸۴	هرمزگان (آقای زارعی)	۰۹۱۹۶۸۲۹۲۸۰	رفسنجان (آقای یوسفی)
۰۹۱۹۵۷۳۰۱۵۵	همدان (آقای سوری)	۰۹۱۹۹۱۰۱۲۴۵	زاهدان (آقای مهمان‌دوست)
۰۹۱۹۵۹۰۷۲۰۴	یاسوج (آقای مرتضوی)	۰۹۱۹۲۷۰۵۸۷۱	زنجان (خانم دکتر هوشیار)
۰۹۱۹۹۱۰۱۲۴۳	یزد (خانم آزاد)	۰۹۱۹۷۷۸۱۹۴۴	ساری (آقای دکتر اکبری)

بانک کتاب ناهید



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

✓ جامع‌ترین بانک کتاب

✓ تحویل روزانه

✓ ارسال به تمامی نقاط کشور

✓ سفارش کتاب به صورت 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 - روبه روی درب اصلی دانشگاه تهران

پاساژ فروزنده - طبقه همکف - پلاک ۳۳۱

تلفن: ۶۶۴۸۹۳۷۵ - ۰۲۱ - ۶۶۴۸۹۳۴۹ - ۰۲۱