

آمار و کاربرد آن در مدیریت ۲

جهت آزمون کارشناسی ارشد در گرایش های رشته مدیریت و حسابداری

شامل آموزش دروس

تستهای آمار رشته های مدیریت از سال ۱۳۹۸ الی ۱۴۰۲

مؤلف : علی زارعی

۶	مقدمه
۷	فصل اول
۸	آمار توصیفی
۸	تعریف آمار
۸	مقیاس ها
۱۰	فراوانی مطلق
۱۰	فراوانی تجمعی
۱۱	نمودارهای آماری
۱۴	معیارهای پراکندگی
۲۱	مراحل محاسبه میانه در جدول توزیع فراوانی
۲۳	محاسبه‌ی مد (نما)
۲۸	تبدیل‌های خطی
۲۸	تشخیص چولگی (کجی)
۳۱	فصل دوم
۳۲	تعریف احتمال
۳۳	اصول شمارش
۳۴	ترتیب
۳۴	ترکیب
۳۶	متمم یک پیشامد
۳۶	پیشامدهای ناسازگار
۳۶	پیشامدهای مستقل
۳۹	احتمال شرطی
۴۲	قضیه بیز (احتمال پسین)
۴۵	فصل سوم
۴۶	متغیرهای تصادفی و توزیع‌های احتمال
۴۶	تابع احتمال
۴۶	امید ریاضی
۵۳	توزیع‌های احتمال گسسته
۶۸	توابع چگالی پیوسته خاص
۷۰	توزیع نمایی
۷۱	تابع چگالی احتمال خی دو
۷۲	توزیع نرمال
۷۹	فصل چهارم

نمونه گیری.....	۸۰
روشهای نمونه گیری برآوردگرها.....	۸۰
مصاحبه ساختار یافته.....	۸۱
فصل پنجم.....	۸۵
توزیع های نمونه گیری.....	۸۶
توزیع دو جمله ای.....	۸۹
تقریب نرمال بر توزیع دو جمله ای.....	۸۹
توزیع نسبت:.....	۹۰
توزیع اختلاف میانگینها.....	۹۰
توزیع نمونه ای اختلاف بین میانگینهای دو جامعه نرمال مستقل وقتی واریانس جامعه معلوم باشد.....	۹۰
برآورد فاصله ای میانگین واقعی جامعه وقتی واریانس جامعه معلوم باشد.....	۹۱
چون حجم نمونه بیشتر از ۳۰ است از توزیع نرمال استفاده می کنیم.....	۹۳
برآورد فاصله ای یکطرفه.....	۹۳
برآورد فاصله ای نسبتها.....	۹۴
برآورد فاصله ای واریانس و انحراف معیار.....	۹۴
برآورد فاصله ای اختلاف میانگین دو جامعه نرمال مستقل وقتی واریانس جوامع معلوم باشد.....	۹۵
برآورد فاصله اختلاف میانگین دو جامعه نرمال مستقل وقتی واریانس جوامع نامعلوم باشد.....	۹۵
برآورد فاصله ای نسبت دو واریانس.....	۹۶
آزمون فرض های آماری.....	۹۹
آزمون فرض میانگین با مقدار ثابت از جامعه نرمال و واریانس جامعه معلوم باشد.....	۱۰۰
P_ مقدار (P-VALUE).....	۱۰۰
آزمون فرض میانگین با مقدار ثابت از جامعه نرمال وقتی واریانس جامعه نامعلوم باشد.....	۱۰۱
حجم نمونه بیشتر از ۳۰ باشد.....	۱۰۲
آزمون فرض نسبت.....	۱۰۳
آزمون فرض واریانس با مقدار ثابت.....	۱۰۴
آزمون فرض اختلاف میانگین دو جامعه نرمال مستقل وقتی واریانس جوامع معلوم باشد.....	۱۰۵
آزمون فرض اختلاف میانگین دو جامعه نرمال مستقل وقتی واریانس جوامع نامعلوم باشد.....	۱۰۵
مقایسه های زوج شده.....	۱۰۷
فصل ششم.....	۱۱۳
ضریب همبستگی.....	۱۱۴
ضریب تعیین.....	۱۱۴
رگرسیون خطی.....	۱۱۸
مراحل نوشتن معادله خط رگرسیون.....	۱۱۹
پیش فرض های رگرسیون.....	۱۲۱
متغیر تعدیل کننده.....	۱۲۹

ضریب رگرسیون تفکیکی استاندارد د شده	۱۲۹
فرض های رگرسیون چند متغیری	۱۳۸
محاسبات رگرسیون چند متغیری	۱۳۸
ضرایب استاندارد شده	۱۳۸
ضریب تعیین تعدیل شده	۱۳۸
آماره آزمون F	۱۳۹
انواع روش های ورود داده ها در رگرسیون	۱۳۹
رگرسیون سلسله مراتبی	۱۳۹
تحلیل مسیر	۱۳۹
فصل هفتم	۱۴۱
تحلیل واریانس	۱۴۲
تحلیل واریانس یک عامله	۱۴۲
تحلیل واریانس دو عامله یک مشاهده به ازای هر ترکیب (بدون تاثیر متقابل)	۱۴۴
تحلیل واریانس دو عامله یا m مشاهده به ازای هر ترکیب (طرح فاکتوریل)	۱۴۶
متغیر کنترل	۱۵۰
فصل هشتم	۱۵۱
ضریب توافقی C	۱۵۵
ضریب فی	۱۵۵
ضریب کریمر	۱۵۵
آزمون رتبه علامتدار	۱۵۹
آزمون مجموع رتبه ها	۱۶۰
آزمون مجموع رتبه ها آزمون H	۱۶۰
تجزیه و تحلیل فریدمن	۱۶۱
آزمون ردیف	۱۶۲
ضریب همبستگی رتبه ای اسپیرمن	۱۶۳
آزمون کولموگروف اسمیرنوف	۱۶۳
منابع	۱۶۶

مقدمه

با توجه به توسعه روز افزون علم و آشنایی با روش های آماری و محاسبات آماری از جمله نیازهای علوم تجربی است. این کتاب بر اساس تجربه ی سال ها تدریس در دروس مختلف جهت سهولت در درس آموزش آمار و کاربرد آن در مدیریت نوشته شده است. این مجموعه شامل تدریس کلیه مطالبی است که در سالهای متفاوت در آزمون های مختلف مطرح شده می باشد و با توضیح ساده هر مطلب درسی به همراه مثال های کاربردی و نکات مهم می باشد. در این اثر سعی شده است تا پاسخ هر مثال به صورت کاملاً ساده و روان بیان شود تا کلیه دانشجویانی که نیاز به منبعی در مورد آمار و کاربرد آن در مدیریت دارند بتوانند از مطالب آن استفاده کنند. خواهشمند است هرگونه پیشنهاد در مورد بهبود این کتاب را به اینجانب با شماره ۰۹۱۲۳۲۱۱۰۷۸ اطلاع دهید.

با تشکر

علی زارعی

فصل چهارم

نمونه گیری

نمونه گیری

روشهای نمونه گیری برآوردها : به طور کلی انواع روش های نمونه گیری را می توان در دو بخش نمونه گیری احتمالی و غیر احتمالی جای داد.

۱- نمونه گیری غیر احتمالی

(a) در دسترس (b) داوطلب (c) هدفمند (d) سهمی (e) بدون نظم (دیمی)

(f) شبکه ای (گلوله برفی)

۲- نمونه گیری احتمالی

(a) تصادفی ساده (b) سیستماتیک (منظم) (c) طبقه بندی (d) خوشه ای

در نمونه گیری احتمالی هر یک از واحدهای جامعه می توانند با احتمالی مشخص در نمونه قرار گیرند. در حالیکه در نمونه گیری غیر احتمالی انتخاب نمونه براساس قوانین احتمالات صورت نمی گیرد و نمونه به کمک قضاوت انسانی تهیه می شود. بنابراین اشتباهات برآوردهای غیر احتمالی، اغلب غیر تصادفی و غیرقابل اندازه گیری است. در این روش ها هر قدر که حجم نمونه را بزرگ اختیار کنیم، نمونه ها اغلب نمی توانند معرف واقعی جامعه باشند. اما با این حال گاهی اوقات نمونه گیری غیراحتمالی بهترین روش نمونه گیری می باشد. مانند زمانی که امکان تهیه چارچوب نمونه گیری وجود نداشته باشد. یکی از مهم ترین کاربردهای نمونه گیری غیراحتمالی، استفاده در مطالعات مقدماتی پیمایش های بزرگ است. در چنین مواردی که هدف تعیین واریانس صفت مورد اندازه گیری (برای تعیین حجم نمونه و روش نمونه گیری) یا تست پایایی ابزار اندازه گیری و... می باشد، احتیاجی به روش های نمونه گیری احتمالی نیست.

۱- نمونه گیری غیر احتمالی

(a) در دسترس : این روش ضعیف ترین نوع نمونه گیری غیر احتمالی است. که در آن نمونه ها داوطلبانه و هدفمند هستند مانند گروهی از دانشجویان که در یک درس انتخاب واحد کرده اند این نمونه ها در جمع آوری اطلاعات اکتشافی مفید است ولی امکان برآورد خطا ($e = \bar{X} - \mu$) در آن وجود ندارد. و چون معرف کل جامعه نیست روایی بیرونی ندارد

(b) داوطلب : این روش نیز غیر احتمالی است چون نمونه ها بطور تصادفی انتخاب نشده اند و امکان دارد جهت دار باشد و برآورد با آنها برآورد صحیحی از جامعه ندهد.

مصاحبه بدون ساختار : این گونه مصاحبه به طور معمول از پژوهش هایی که در رابطه با سرگذشت زندگی افراد باشد بهره گیری می شود و کاربرد فراوانی در روش تحقیق جهت تحلیل آماری دارد. در روش مصاحبه بدون ساختار پژوهشگر سعی می کند که نقاط عقیده ها یا شرایط مصاحبه شوندگان را درک کند

مصاحبه ساختار یافته: در مصاحبه ساختار یافته سوالات از پیش تعیین شده و برنامه‌ریزی شده است. سوالات به ترتیب مشخصی پرسیده می‌شود. مصاحبه شونده در پاسخگویی به سوالات کم‌ترین آزادی عمل ممکن را در مقایسه با انواع دیگر مصاحبه دارد. از انعطاف پذیری پایینی برخوردار است.

مصاحبه تاییدی: در مواقعی استفاده می‌شود که حجم نمونه بسیار زیاد بوده و شواهدی برای بررسی داده‌هایی که از قبل جمع‌آوری شده فراهم می‌کند.

(c) قضاوتی (هدفمند): معمولاً در نمونه‌گیری قضاوتی انتخاب سنجیده واحدها به طریقی صورت می‌گیرد که هریک معرف بخشی از جامعه مورد نظر باشند. به عنوان مثال فرض کنید قصد داریم به بررسی وضعیت نوجوانانی که از خانه فرار کرده‌اند بپردازیم و مصاحبه با چنین افرادی ضروری است. از آنجایی که تهیه لیستی از این افراد ممکن نیست، محقق سعی می‌کند به مکان‌هایی مانند پارک‌ها، ترمینال‌ها و... مراجعه کرده و به طور عمدی افراد مورد نظر را در مکان‌های خاص شناسایی کند. این روش نمونه‌گیری با توجه به موضوع مورد بررسی براساس فهم نظری و تجربه پیشین محقق از جمعیت مورد مطالعه است. در نمونه‌گیری هدفمند ملاک تعیین حجم نمونه اشباع شدن است.

(d) سهمیه‌ای: در این روش جامعه آماری به چند طبقه تقسیم شده سپس به اختیار سهمی به هر طبقه اختصاص داده می‌شود و در ادامه نمونه‌هایی را که دسترسی به آن‌ها ساده‌تر است به دلخواه انتخاب می‌کنند. البته ممکن است جامعه مورد مطالعه طبقه‌بندی نشود، در این حالت نیز آنقدر نمونه اختیار می‌شود تا حجم نمونه مورد نظر حاصل شود.

برای مثال فرض کنید می‌خواهیم به بررسی رضایت شغلی کارکنان صنایع بزرگ استان تهران بپردازیم و قبل از مرحله اصلی گردآوری داده‌ها احتیاج به یک مطالعه مقدماتی برای سنجش واریانس صفت مورد نظر، روایی و پایایی پرسشنامه داریم و از طرفی براساس تحقیقات قبلی می‌دانیم رضایت شغلی کارکنان صنایع مختلف با یکدیگر تفاوت دارد. با فرض این که حجم نمونه مقدماتی ۵۰ نفر در نظر گرفته شده باشد و با در نظر گرفتن نسبتی از کل کارکنان که در هر صنعت مشغول به کار هستند، سهمیه هر گروه را تعیین می‌کنیم.

در نمونه‌گیری سهمیه‌ای غالباً آن دسته از ویژگی‌هایی که در تحقیقات پیشین هم تشخیص داده شده‌اند مورد مطالعه قرار می‌گیرند به همین دلیل نمونه‌گیری سهمیه‌ای تا حد زیادی به تحقیقات پیشین و یافته‌های بدست آمده در میدان تحقیق وابسته است. مشکل دیگر این روش عدم توجه به برخی ویژگی‌های جمعیت مورد نظر است. اما در مجموع با توجه به تشابه ساختار نمونه و جمعیت در نمونه‌گیری سهمیه‌ای، احتمال معرف بودن برآوردهای حاصل از آن بیش از برآوردهایی است که از نمونه‌گیری اتفاقی یا قضاوتی به دست می‌آید. و خطای آن کمتر است.

(e) اتفاقی (بدون نظم یا دیمی): در این روش محقق افرادی را مورد مطالعه قرار می‌دهد که در دسترس هستند و مصاحبه‌گر در چارچوب تعداد و حجم نمونه افرادی را به طور اتفاقی انتخاب کرده و با آن‌ها مصاحبه می‌کند. به عنوان مثال محقق می‌خواهد نگرش مردم را نسبت به یک نمایشگاه بسنجد، از میان افرادی که از نمایشگاه دیدن کرده‌اند به طور اتفاقی چند نفر را انتخاب می‌کند. با توجه به موضوع تحقیق مکان‌هایی که این نمونه‌گیری می‌تواند در آن صورت گیرد مانند ایستگاه اتوبوس، بازار، فروشگاه،

بیمارستان و... می باشد. نکته قابل توجه این است که در نمونه گیری قضاوتی محقق افراد خاص را انتخاب می کند اما در نمونه گیری اتفاقی، هرچند محقق در مکان های خاصی قرار می گیرد، اما به طور اتفاقی افرادی را در آن مکان ها انتخاب می کند. نمونه گیری اتفاقی عمدتاً با مشاهده مشارکتی و آزمایشی توأم است و از طریق آن محقق همواره موضوعات و محیط هایی را که به آسانی قابل دسترسی باشند، مورد مطالعه قرار می دهد. نمونه گیری به روش اتفاقی راحت تر است.

f) شبکه ای (گلوله برفی یا زنجیره ای): این روش در مقایسه با دیگر روش های ذکر شده چندان منظم نیست، کاربرد کمتری دارد و بیشتر در تحقیقات قوم نگاری مورد استفاده قرار می گیرد تا در مطالعات پیمایشی.

این فن شامل شناسایی برخی افراد مهم یک جمعیت و مصاحبه با آن ها است، سپس محقق به پیشنهاد این افراد برای مصاحبه به سراغ افراد دیگر می رود. در این روش هسته کوچک اصلی، با افزایش مرحله ای، رشد می کند و مانند گلوله برفی که با غلتیدن بر روی زمین بزرگ می شود، نمونه تحقیق نیز افزایش می یابد. این روش معمولاً زمانی مورد استفاده قرار می گیرد که چارچوب نمونه گیری وجود ندارد و از طرفی افراد نمونه نسبت به یکدیگر شناخت دارند. و میتوان اعضای پنهان را جامعه را پیدا کرد. به عنوان مثال فرض کنید می خواهیم به علل اعتیاد از دیدگاه معتادان بپردازیم. در این صورت با شناسایی تعدادی از معتادان که حاضر به مصاحبه هستند، از آن ها نام نشانی افراد دیگر را که می شناسند جویا شده و به این ترتیب به تعداد مورد نیاز نمونه تهیه کنیم. از این روش هنگامی استفاده می شود که اعضای گروه هدف به نوعی با یکدیگر در ارتباط بوده و دارای ویژگی های مشترک باشند. مزیت این روش در این است که محقق می تواند به بررسی کل شبکه بپردازد اما از طرف دیگر او تنها محدود به کسانی خواهد بود که در این شبکه مرتبط به هم قرار دارند.

g) گروهی : در این روش تنها با گروه های معینی تماس حاصل شده و اطلاعات لازم گرفته می شود. در این روش محقق براساس تجربه شخصی و یا تجارب مشابه دیگران، یک گروه اجتماعی را معرف جامعه ای که به آن تعلق دارند، می یابد و با آن ها مصاحبه می کند. گروه مورد مطالعه می تواند اعضای شورای اسلامی روستا، ریش سفیدان و... باشند. نتایج این نمونه ها هرچند ممکن است معرف جامعه باشند، اما نمی توان با اطمینان آماری به جامعه مورد مطالعه تعمیم داد.

۲- نمونه گیری احتمالی

a) تصادفی ساده : در این روش نمونه گیری واحدهای مورد انتخاب دارای شانس مساوی $\frac{n}{N}$ برای انتخاب شدن هستند. که در آن **n** تعداد نمونه و **N** تعداد جامعه است در اینجا قوانین احتمال است که معین می کند کدام واحدها یا افراد از جامعه انتخاب شوند. انتخاب یا از طریق قرعه کشی است و یا از طریق استفاده از جدول اعداد تصادفی. در نمونه گیری تصادفی افراد از فهرست جامعه انتخاب می شوند در روش قرعه کشی ابتدا کلیه واحدها یا افراد شماره بندی شده و یا اسامی آنها تهیه می شود و سپس به قید قرعه از بین آنها تعداد لازم برای نمونه انتخاب می شود. برای کنترل واریانس در گمارش تصادفی اثر یک متغیر در بین گروه های مورد بررسی توزیع می شود. این نمونه گیری معمولاً به یکی از روشهای زیر انجام می شود

الف) نمونه گیری تصادفی ساده بدون جایگذاری : در این روش نمونه انتخابی کنار گذاشته شده و حذف می‌شود یک ویژگی مهم نمونه گیری تصادفی ساده بدون جایگذاری این است که احتمال استخراج هر واحد مشخص از جامعه در هر استخراجی مساوی با احتمال استخراج آن واحد مشخص در استخراج اول است. تعداد نمونه انتخابی از رابطه $\binom{N}{n} = \frac{N!}{n!(N-n)!}$ بدست می آید. و احتمال انتخاب از رابطه

$$p(A) = \frac{1}{\binom{N}{n}}$$

بدست می آید

ب) نمونه گیری تصادفی ساده با جایگذاری: اگر در انتخاب n واحد نمونه پس از انتخاب هر واحد آن را به جامعه بر گردانیم و انتخاب بعدی انجام دهیم نمونه گیری تصادفی ساده با جایگذاری می نامند در این روش انتخاب هر واحد مستقل از انتخاب واحدهای دیگر است. همچنین در این طرح نمونه گیری احتمال انتخاب مجدد و دوباره نمونه ها وجود دارد.

(b) سیستماتیک (منظم): نمونه گیری سیستماتیک مشتمل بر گزینش واحدها به روشی سیستماتیک و در نتیجه به صورتی غیر تصادفی است. منظور از این فن نمونه گیری معمولاً پخش کردن واحدها به طور یکنواخت بر روی چارچوب است عنصر تصادفی بودن اغلب به این ترتیب دخالت داده می شود که اولین واحد را به طور تصادفی انتخاب می کنند. در این صورت گزینش اولین واحد، بقیه واحدهای نمونه را معین می کنند. یعنی ابتدا $I = \frac{N}{n}$ را بدست می آوریم سپس یک عدد بطور تصادفی کمتر از I را انتخاب می‌کنیم. سپس به اندازه I به نمونه انتخابی اضافه می‌کنیم تا تعداد n نمونه انتخاب شود. مثلاً از ۱۰۰ نفر بخواهیم نمونه ۵ نفری انتخاب کنیم اول از ۱ تا ۱۰۰ افراد را شماره گذاری می‌کنیم سپس $I = \frac{N}{n} = \frac{100}{5} = 20$ را محاسبه می‌کنیم آنگاه عددی بین ۰ و ۵ انتخاب می‌کنیم مثلاً ۲ سپس ۲۰ تا ۲۰ تا جلو می‌رویم یعنی انتخابهای ما شماره های ۲-۲۲-۴۲-۶۲-۸۲ میشود. هر گاه عنصرهای جامعه بطور تصادفی شماره گذاری شوند دقت نمونه گیری سیستماتیک (منظم) با نمونه گیری تصادفی ساده برابر است.

(c) طبقه بندی: از این روش نمونه گیری در مواقعی استفاده میشود که ترکیب جامعه همگن نباشد و افراد دارای صفات غیر مشترکی باشند. در این روش نمونه گیری برای اجتناب از اشکالاتی که ممکن است در روش قبلی با آن مواجه شویم، افراد جامعه آماری را بسته به خصوصیات که آنها را از یکدیگر متمایز می سازد به طبقات مختلف تقسیم می کنیم. در نمونه گیری تصادفی طبقه ای، واحدهای جامعه مورد مطالعه در طبقه هایی که از نظر صفت متغیر همگن تر هستند، گروه بندی می شوند. به این ترتیب تغییرات در درون گروه ها حداقل می شود. در نمونه گیری طبقه بندی تغییرات داخل طبقات کم و تغییرات بین طبقات زیاد است. معمولاً برای طبقه بندی واحدهای جامعه، متغیری به عنوان ملاک در نظر گرفته می شود که با صفت متغیر مورد مطالعه بستگی داشته باشد. سپس به تعداد مورد نیاز و متناسب با جمعیت هر یک از طبقات افراد نمونه را انتخاب می کنیم. انتخاب افراد می تواند هم به روش تصادفی باشد و هم به روش تصادفی سیستماتیک. بنابر این روش در جمعیت‌های نا همگن که توزیع جمعیت در گروه‌ها و طبقات مختلف متفاوت است، استفاده میشود. نمونه گیری تصادفی طبقه ای هنگامی که بخواهیم نمونه

ای نسبتاً کوچک از جامعه ای بزرگ که شامل چند زیر گروه عمده باشد انتخاب کنیم که نشانده دهنده هدف نمونه باشد و انتخاب بر اساس اصول احتمال باشد استفاده می شود در این روش اطلاعات از یک طبقه با طبقه دیگر متفاوت است مانند انتخاب دانشجویان یک دانشگاه در این روش واریانس (تغییر پذیری) هر طبقه باید از واریانس کل کمتر باشد در غیر اینصورت دقت کاهش پیدا می کند.

برای مثال به منظور بررسی نسبت قبول شدگان در مدارس تیز هوشان پایه ششم ابتدایی در شهر تهران و رابطه آن با محل جغرافیایی دبستان، می توان دبستان های شهر تهران را برحسب محل آن ها به پنج طبقه تقسیم کرد. طبقه اول دبستان های شمال غربی، طبقه دوم دبستان های شمال شرقی، طبقه سوم دبستان های مرکزی، طبقه چهارم دبستان های جنوب غربی و طبقه پنجم دبستان های جنوب شرقی. پس از آن از هر طبقه تعدادی دبستان به روش تصادفی ساده انتخاب می شود. در نمونه گیری طبقه ای حجم نمونه n را به شیوه های مختلف می توان میان طبقات تقسیم کرد. ساده ترین شیوه تقسیم مساوی تعداد نمونه میان طبقات است. سایر شیوه ها شامل انتساب بهینه و انتساب متناسب می باشند.

d) خوشه ای: نمونه گیری خوشه ای یک نوع نمونه گیری تصادفی است در بسیاری از مواقع می توان با اجرای روش انتخاب تصادفی گروهها یا خوشه هایی از واحدهای نمونه گیری به جای گرفتن نمونه تصادفی ساده از کل جامعه است. یکی از مزیت های این طرح نمونه گیری این می باشد که در میزان هزینه به طور اساسی صرفه جویی صورت می پذیرد. نمونه گیری خوشه ای ما را از ساختن چارچوب برای تمامی جامعه بی نیاز می کند که این چارچوب خود اغلب یک کار پر خرج و خسته کننده ای است. به علاوه چون واحدهای یک خوشه، مجاور هم هستند و بنابراین دسترسی به آنها آسان است، فرایند نمونه گیری خوشه ای بطور قابل توجهی به مقرون به صرفه است. و کاربرد آن در مواقعی است که اگر جامعه دارای تجانس (همگن) کامل باشد و توزیع جغرافیایی افراد یا جامعه مورد مطالعه پراکنده باشد و یا فهرست افراد در دسترس نباشد هزینه ها افزایش می یابد که با این روش در هزینه صرفه جویی میشود. در نمونه گیری خوشه ای تغییرات داخل خوشه ها زیاد و تغییرات بین خوشه ها کم است. به عبارتی در نمونه گیری خوشه ای واریانس (تغییر پذیری) درون طبقات زیاد و بین طبقات کم است. برای مثال فرض کنید می خواهیم درآمد هر خانوار را در یک شهر بزرگ برآورد کنیم، اگر از نمونه گیری تصادفی ساده استفاده کنیم نیاز به فهرست تمام خانوارهای شهر داریم. یافتن این چارچوب می تواند بسیار پرهزینه و یا غیرممکن باشد، استفاده از نمونه گیری طبقه ای نیز در چنین جامعه ای مستلزم داشتن فهرستی از خانوارها در هر طبقه می باشد. در این حالت می توانیم شهر را به نواحی معینی مانند بلوک ها (خوشه هایی از اعضا) تقسیم نماییم و یک نمونه تصادفی ساده را از بلوک های این جامعه انتخاب کرده و در آمد هر خانوار را در بلوک هایی که در نمونه واقع شده است اندازه بگیریم. این کار با استفاده از یک چارچوب که کلیه بلوک های شهر را فهرست نموده، به اجرا در می آید. دلیل دیگری که استفاده از نمونه گیری خوشه ای را پیشنهاد می دهد جمع آوری داده ها با هزینه ی کمتر است. برای مثال فرض کنید فهرستی از خانوارهای شهر در دسترس است و می خواهیم نمونه ای به روش تصادفی ساده از خانوارهای مختلف که در سطح شهر پراکنده اند تهیه کنیم. هزینه ی مصاحبه با خانوارها با توجه به هزینه رفت و آمد پرسشگر و دیگر هزینه ها بالا است. یک روش کاهش هزینه های رفت و آمد استفاده از روش نمونه گیری خوشه ای می باشد زیرا افراد درون یک خوشه باید از نظر جغرافیایی به هم نزدیک باشند. در نمونه گیری خوشه ای واحد تحلیل فرد نیست.

فصل پنجم

توزیع های نمونه گیری

برآورد فاصله ای

آزمون فرض

توزیع‌های نمونه‌گیری

برآوردگر: هر تابعی از نمونه که به پارامتر یا پارامترهای جامعه بستگی نداشته باشد برآوردگر یا آماره گویند مقدار عددی برآوردگر را برآورد گویند.

ویژگیهای برآوردگر:

۱- نااریب باشد یعنی $E(\hat{\theta}) = \theta$ که $b = \theta - E(\hat{\theta})$ را مقدار اریبی گویند. که در این صورت

$$E(\hat{\theta} - \theta)^2 = E\left[\hat{\theta} - E(\hat{\theta})\right]^2 + E\left[\theta - E(\hat{\theta})\right]^2 \Rightarrow \text{MSE}(\theta) = \text{var}(\theta) + b^2(\theta)$$

۲- کارایی داشته باشد یعنی دارای کمترین واریانس باشد به عنوان مثال اگر θ_1 و θ_2 برآوردگرهای یک آماره باشد و $\text{var}(\theta_1) \leq \text{var}(\theta_2)$ باشد آنگاه θ_1 از θ_2 کاراتر است به طور کلی $e = \frac{\text{MSE}(\theta_2)}{\text{MSE}(\theta_1)}$ و اگر

$$e = \frac{\text{var}(\theta_2)}{\text{var}(\theta_1)} \text{ نا اریب باشد}$$

اگر $e > 1$ یعنی θ_1 کاراتر از θ_2
 $e = 1$ یعنی ارجحیت نسبت به یکدیگر ندارند
 $e < 1$ یعنی θ_2 کاراتر از θ_1

۳- سازگاری: برآورد کننده $\hat{\theta}$ را سازگار گویند. اگر با بزرگ شدن اندازه نمونه از لحاظ احتمال به سمت پارامتر θ میل کنند یعنی

$$\lim_{n \rightarrow \infty} |P_{\theta} - \hat{P}_{\theta}| = 0$$

مثال: اگر X_1, X_2, X_3, X_4 یک نمونه ۴ تایی از جامعه‌ای با $\mu = 20$ و $\sigma^2 = 10$ باشد کدام یک از برآوردگرها

$$\bar{X}_1 = \frac{1}{4} \sum_{i=1}^4 X_i \quad \bar{X}_2 = \frac{1}{4} (X_1 + 2X_2 + X_3) \quad \text{مناسب است؟}$$

$$E(\bar{X}_1) = E\left(\frac{1}{4} \sum_{i=1}^4 X_i\right) = \frac{1}{4} \sum_{i=1}^4 E(X_i) = \frac{4 \times 20}{4} = 20$$

$$E(\bar{x}_2) = E\left[\frac{1}{4}(x_1 + 2x_2 + x_3)\right] = \frac{1}{4}[E(x_1) + 2E(x_2) + E(x_3)] = \frac{4 \times 20}{4} = 20$$

بنابراین نااریب هستند

$$\text{var}(\bar{x}_1) = \text{var}\left[\frac{1}{4}\sum_{i=1}^4 x_i\right] = \frac{1}{16}\left[\sum_{i=1}^4 \text{var}(x_i)\right] = \frac{4 \times 10}{16} = 2/5$$

$$\begin{aligned}\text{var}(\bar{x}_2) &= \text{var}\left[\frac{1}{4}(x_1 + 2x_2 + x_3)\right] = \frac{1}{16}[\text{var}(x_1) + 4\text{var}(x_2) + \text{var}(x_3)] \\ &= \frac{1}{16}[10 + 40 + 10] = 3/4\end{aligned}$$

بنابراین $\text{var}(\bar{x}_1) < \text{var}(\bar{x}_2)$ یعنی $\bar{x}_1$ کاراتر از $\bar{x}_2$ است.

ارشد مدیریت ۱۴۰۰

۶۰ به منظور برآورد پارامتری از جامعه از دو برآوردکننده استفاده شده است. برآوردکننده اول دارای اریبی صفر و انحراف معیار ۲ و برآوردکننده دوم دارای اریبی ۱ و انحراف معیار ۲ می باشد. کارایی برآوردکننده دوم به اول چقدر است؟

(۱) ۵/۸

(۲) ۵/۹

(۳) ۱

(۴) ۱/۲۵

$$\begin{aligned}\sigma_{\theta_1} &= 2 \quad b_{\theta_1} = 0 \quad \sigma_{\theta_2} = 2 \quad b_{\theta_2} = 1 \\ \text{MSE}(\theta_1) &= \text{var}(\theta_1) + b^2(\theta_1) = 2^2 + 0 = 4 \\ \text{MSE}(\theta_2) &= \text{var}(\theta_2) + b^2(\theta_2) = 2^2 + 1 = 5 \\ e &= \frac{\text{MSE}(\theta_1)}{\text{MSE}(\theta_2)} = \frac{4}{5} = 0.8\end{aligned}$$

گزینه ۱ صحیح است

توزیع میانگین نمونه‌ای:

۱- جامعه متناهی باشد یعنی اگر $\frac{n}{N} > 0.05$ باشد جامعه را متناهی گویند که در این صورت نیاز به تصحیح واریانس است و $\frac{N-n}{N-1}$ که ضریب تصحیح واریانس گویند.

$$E(\bar{x}) = \mu, \quad \text{var}(\bar{x}) = \frac{N-n}{N-1} \cdot \frac{\sigma^2}{n}$$

۲- جامعه نامتناهی باشد: اگر $\frac{n}{N} \leq 0.05$ باشد جامعه را نامتناهی گویند و ضریب تصحیح واریانس تقریباً

یک است.

$$E(\bar{x}) = \mu, \quad \text{var}(\bar{x}) = \frac{\sigma_x^2}{n}$$

شرایط تقریب خوب به نرمال:

هرگاه توزیع فراوانی متقارن باشد اگر $n \geq 30$ باشد توزیع $\bar{x}$ را تقریباً نرمال در نظر می‌گیریم و در صورتی که توزیع فراوانی چوله باشد اگر $n \geq 100$ باشد توزیع $\bar{x}$ را تقریباً نرمال در نظر می‌گیریم.

نکته: چون میانگین توزیع نمونه‌گیری $\bar{X}$ برابر μ است $\bar{X}$ یک برآورد گر بدون سوگیری از μ است مثال: یک نمونه به اندازه $n=100$ به طور تصادفی از جامعه‌ای شامل یک میلیون عنصر با میانگین ۱۰ و واریانس ۳۶ انتخاب می‌کنیم.

الف) احتمال اینکه میانگین نمونه‌ای بزرگتر یا مساوی ۱۱/۲ باشد را بدست آورید.

ب) اگر حجم جامعه ۱۰۰۰ عنصر باشد احتمال آن که میانگین نمونه‌ای بیشتر از ۱۱/۲ باشد را بدست آورید.

$$\text{الف) } \frac{n}{N} = \frac{100}{1000000} = 0.0001 < 0.05 \Rightarrow \bar{z} = \frac{\bar{x} - \mu}{\frac{\sigma_x}{\sqrt{n}}} \text{ نیاز به تصحیح ندارد}$$

$$P(\bar{x} \geq 11/2) = P\left(z \geq \frac{11/2 - 10}{\frac{6}{\sqrt{100}}}\right) = P(z > 2) = 0.0228$$

$$\text{ب) } \frac{n}{N} = \frac{100}{1000} = 0.1 > 0.05 \Rightarrow \bar{z} = \frac{\bar{x} - \mu}{\sqrt{\frac{N-n}{N-1}} \cdot \frac{\sigma_x}{\sqrt{n}}} \text{ تصحیح دارد}$$

$$P(\bar{x} \geq 11/2) = P\left(z \geq \frac{11/2 - 10}{\sqrt{\frac{1000-100}{1000-1}} \times \frac{6}{\sqrt{100}}}\right) = P(Z \geq 2/1) = 1 - \phi(2/1) =$$

$$1 - 0.9821 = 0.0179$$

مثال: در یک نمونه تصادفی ۴۰۰ نفری توزیع نمرات خلاقیت کلامی دانش آموزان دارای توزیع نرمال با میانگین ۱۳۵ و انحراف استاندارد ۱۸ است تقریباً نمره خلاقیت چند دانش آموز کمتر یا مساوی ۱۷۱ است؟

۳۹۲(۴)

۳۳۶(۳)

۲۲۹(۲)

۱۸۹(۱)

$$n=400, \sigma=18, \mu=135$$

$$p(x \leq 171) = p(z \leq \frac{171-135}{18}) = p(z \leq 2) = 0.9772$$

$$p = \frac{x}{n} \Rightarrow x = p.n = 0.9772 \times 400 = 390.88 \approx 391$$

گزینه ۴ صحیح است

ارشد مدیریت ۱۴۰۱

۶۲- میانگین توزیع نمونه‌گیری واریانس نمونه‌های ۳ تایی از جامعه‌ای نرمال با ۴ عضو برابر ۱۶ است. انحراف معیار توزیع نمونه‌گیری میانگین نمونه‌های ۳ تایی از این جامعه چقدر است؟

$$\sqrt{12} \quad (1)$$

$$12 \quad (2)$$

$$4 \quad (3)$$

$$2 \quad (4)$$

$$\text{var}(\bar{x}) = \sigma_{\bar{x}}^2 = \frac{\sigma_x^2}{n} = \frac{16}{4} = 4 \Rightarrow \sigma_{\bar{x}} = \sqrt{4} = 2$$

گزینه ۴ صحیح است

ارشد مدیریت ۱۴۰۲

۵۶- میانگین توزیع نمونه‌گیری واریانس نمونه‌های تصادفی ۳ تایی از جامعه‌ای نرمال برابر ۲۰ و انحراف معیار توزیع نمونه‌گیری میانگین نمونه‌های ۴ تایی از این جامعه ۲ است. تعداد اعضای این جامعه، کدام است؟

$$4 \quad (1)$$

$$5 \quad (2)$$

$$16 \quad (3)$$

$$25 \quad (4)$$

$$\sigma_{\bar{x}}^2 = \frac{\sigma_x^2}{n} \Rightarrow n = \frac{\sigma_x^2}{\sigma_{\bar{x}}^2} = \frac{20}{2^2} = 5$$

گزینه ۲ صحیح است

توزیع دو جمله‌ای: اگر x دارای توزیع دو جمله‌ای با پارامترهای P و n باشد در صورتی که n بزرگ باشد x دارای توزیع تقریبی نرمال با میانگین $\mu = np$ و واریانس $\text{var}(x) = npq$ است.

تقریب نرمال بر توزیع دو جمله‌ای:

شرایط خوب برای استفاده از این تقریب آن است که $n > 20$ (۱) و $nP > 5$ (۲) و $nq > 5$ که در این صورت باید تصحیح پیوستگی انجام گیرد. چون توزیع دو جمله‌ای گسسته و توزیع نرمال پیوسته است.

$$P(a \leq x \leq b) = P(a - \cdot / \Delta < x < b + \cdot / \Delta) P = \left(\frac{a - \cdot / \Delta - nP}{\sqrt{nPq}} \right) < z < \left(\frac{b + \cdot / \Delta - nP}{\sqrt{nPq}} \right)$$

مثال: احتمال آن که در ۲۵ بار پرتاب یک سکه بین هفت تا ۱۵ بار شیر ظاهر شود را بدست آورید.

الف) با استفاده از توزیع دو جمله‌ای

ب) با استفاده از تقریب نرمال

$$P(7 < x < 15) = P\left(\frac{7 - \cdot / \Delta - 12 / \Delta}{\sqrt{25 \times \cdot / \Delta \times \cdot / \Delta}}\right) < z < \left(\frac{15 + \cdot / \Delta - 12 / \Delta}{\sqrt{25 \times \cdot / \Delta \times \cdot / \Delta}}\right)$$

$$= P(-2 / 4 < z < 1 / 2) = \Phi(1 / 2) - \Phi(-2 / 4) = \cdot / 8767$$

توزیع نسبت: اگر $\bar{P} = \frac{x}{n}$ را به عنوان تبدیل خطی در نظر بگیریم هر n به اندازه کافی بزرگ باشد در این

$$\text{var}(x) = \frac{\bar{P}q}{n} \text{ و } E(x) = \mu = P \text{ صورت}$$

مثال: اگر نسبت موفقیت در جامعه‌ای ۰/۸ باشد و از این جامعه یک نمونه ۱۰۰ تایی انتخاب کنیم احتمال اینکه نسبت موفقیتها بیش از ۰/۷۴ باشد را بدست آورید.

$$P(\bar{P} \geq \cdot / 74) = P\left(z > \frac{\bar{P} - P}{\sqrt{\frac{Pq}{n}}}\right) = P\left(z > \frac{\cdot / 74 - \cdot / 8}{\sqrt{\frac{\cdot / 8 \times \cdot / 2}{100}}}\right) = P(z > 1 / 5) = \cdot / 668$$

توزیع اختلاف میانگینها:

$$E(\bar{x}_1 - \bar{x}_2) = \mu_1 - \mu_2, \quad \sigma_{\bar{x}_1 - \bar{x}_2}^2 = \sigma_{\bar{x}_1}^2 + \sigma_{\bar{x}_2}^2 - 2 \text{cov}(\bar{x}_1 - \bar{x}_2)$$

$$\text{cov}(\bar{x}_1 - \bar{x}_2) = 0 \Rightarrow \sigma_{\bar{x}_1 - \bar{x}_2}^2 = \sigma_{\bar{x}_1}^2 + \sigma_{\bar{x}_2}^2 \Rightarrow \sigma_{\bar{x}_1 - \bar{x}_2} = \sqrt{\sigma_{\bar{x}_1}^2 + \sigma_{\bar{x}_2}^2} \Rightarrow \sigma_{\bar{x}_1 - \bar{x}_2} = \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$$

توزیع نمونه‌ای اختلاف بین میانگینهای دو جامعه نرمال مستقل وقتی واریانس جامعه معلوم باشد

در این صورت کوواریانس بین دو جامعه صفر است بنابراین:

$$E(\bar{x}_1 - \bar{x}_2) = \mu_1 - \mu_2$$

$$\text{var}(\bar{x}_1 - \bar{x}_2) = \frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2} \Rightarrow z = \frac{(\bar{x}_1 - \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$$

مثال: از دو جامعه آماری A و B با مشخصات $\mu_1 = 100$ و $\mu_2 = 105$ و $\sigma_1^2 = \sigma_2^2 = 200$ و نمونه‌های $n_1 = 40$ و $n_2 = 50$ گرفته‌ایم احتمال آنکه اختلاف میانگین کمتر از یک باشد را بدست آورید.

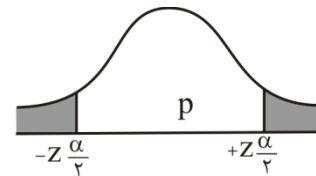
$$P(|\bar{x}_2 - \bar{x}_1| < 1) = P(-1 < \bar{x}_2 - \bar{x}_1 < +1) = P\left(\frac{(-1) - (105 - 100)}{\sqrt{\frac{200}{40} + \frac{200}{50}}} < z < \frac{(+1) - (105 - 100)}{\sqrt{\frac{200}{40} + \frac{200}{50}}}\right)$$

$$= P(-2 < z < +1/33) = \Phi(+1/33) - \Phi(-2) = 1 - \Phi(-2) = 1 - (1 - \Phi(2)) = \Phi(2) - \Phi(1/33) =$$

$$= 0.9772 - 0.9082 = 0.069$$

بر آورد فاصله‌ای میانگین واقعی جامعه وقتی واریانس جامعه معلوم باشد.

در این صورت از توزیع نرمال (z) استفاده می‌کنیم که در سطح اطمینان P می‌توان به صورت زیر نوشت (خطا و $\alpha + P = 1$)



$$\pm Z_{\frac{\alpha}{2}} = \frac{\bar{x} - \mu}{\frac{\sigma}{\sqrt{n}}} \Rightarrow \mu: \bar{x} \pm Z_{\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}$$

نکته ۱: $L = 2Z_{\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}$ طول فاصله اطمینان و $d = |\bar{x} - \mu| = Z_{\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}$ خطای بر آورد

نکته ۲: رابطه انحراف معیار و دامنه تغییرات $R = X_{\max} - X_{\min}$ $\sigma_x \approx \frac{R}{\sqrt{6}}$ & $\sigma_x \approx \frac{R}{\sqrt{4}}$

نکته ۳: روابط حدود بالا و پایین

$$d = \frac{UCL - LCL}{2}, \mu = \frac{UCL + LCL}{2} \Leftrightarrow UCL = \bar{x} + z_{\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}, LCL = \bar{x} - z_{\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}$$

مثال: از جامعه با انحراف استاندارد ۱۰ یک نمونه ۲۵ گرفته ایم که میانگین نمونه ۷۸ بدست آمده است یک برآورد فاصله‌ای ۹۵٪ برای میانگین واقعی جامعه بنویسید.

$$\alpha = 1 - P = 1 - 0.95 = 0.05, \quad \frac{\alpha}{2} = \frac{0.05}{2} = 0.025, \quad Z_{0.025} = 1.96$$

$$\mu: \bar{x} \pm Z_{\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}} = 78 \pm 1.96 < \frac{10}{\sqrt{25}} \Rightarrow 74.08 < \mu < 81.92$$

$$Z_{0.025} = 1.96 \approx 2 \quad Z_{0.05} = 1.645 \text{ حفظ شود}$$

ارشد مدیریت ۱۴۰۰

۵۹- فرض کنید توزیع ضخامت پاک‌کن‌های تولیدی یک دستگاه به صورت $X \sim N(11, 2)$ است. ضخامت ۹۰ درصد

پاک‌کن‌های تولیدی در چه دامنه‌ای در دو طرف میانگین قرار می‌گیرد؟

(در توزیع نرمال استاندارد: $P(Z \leq -1.64) = 0.05$, $P(Z \geq 1.64) = 0.05$)

$$[7.72, 14.28] \quad (1)$$

$$[8.44, 13.56] \quad (2)$$

$$[9.36, 12.64] \quad (3)$$

$$[9.72, 12.28] \quad (4)$$

$$\mu = 11 \quad \sigma = 2 \quad p = 90\%$$

$$\alpha = 1 - P = 1 - 0.9 = 0.1, \quad \frac{\alpha}{2} = \frac{0.1}{2} = 0.05, \quad Z_{0.05} = 1.645$$

$$\mu: \bar{x} \pm Z_{\frac{\alpha}{2}} \sigma = 11 \pm 1.645 \times 2 \Rightarrow \mu: 11 \pm 3.29 \Rightarrow 7.71 < \mu < 14.29$$

گزینه ۱ صحیح است

برآورد فاصله‌ای میانگین واقعی واریانس جامعه نامعلوم باشد در این صورت با استفاده از واریانس نمونه و با توجه به حجم نمونه از توزیع Z یا t استفاده می‌شود.

الف) برآورد فاصله‌ای میانگین واقعی جامعه وقتی واریانس جامعه نامعلوم باشد و حجم نمونه بیش‌تر از

۳۰ باشد در این صورت طبق قضیه حد مرکزی از توزیع Z استفاده می‌شود.

$$\mu: \bar{x} \pm Z_{\frac{\alpha}{2}} \frac{s}{\sqrt{n}}$$

ب) حجم نمونه کمتر از ۳۰ باشد در این صورت به جای توزیع Z از توزیع t با درجه آزادی $v = df = n - 1$

$$\mu: \bar{x} \pm t_{\frac{\alpha}{2}, df} \cdot \frac{s}{\sqrt{n}} \quad (\text{استفاده می‌شود.})$$

مثال: در یک نمونه ۱۶ تایی میانگین نمونه ۶۵ و واریانس نمونه ۸۱ بدست آمده است یک برآورد فاصله‌ای ۹۵٪ برابری میانگین واقعی جامعه بنویسید.

$$n=16 < 30 \rightarrow t$$

$$s^2 = 81 \Rightarrow s = 9$$

$$\bar{x} = 65$$

$$n = 16$$

$$\mu: \bar{x} \pm t_{\alpha, df} \cdot \frac{s}{\sqrt{n}} = 65 \pm t_{0.05, 15} \cdot \frac{9}{\sqrt{16}} = 65 \pm 1.753 \times \frac{9}{4} \Rightarrow 61.055 < \mu < 68.944$$

ارشد مدیریت ۹۸

۴۷- نمونه تصادفی از قیمت کالایی در ۴۹ فروشگاه دارای میانگین ۱۵۰ و انحراف معیار ۱۴ واحد پول است. میانگین

قیمت این کالا با ضریب اطمینان ۰/۹۵ در کدام بازه است؟ ($S_{-\infty}^{-1/96} = 0.025$)

$$(148.08, 151.92) \quad (2)$$

$$(146.08, 153.92) \quad (1)$$

$$(142.26, 157.74) \quad (4)$$

$$(149.02, 150.98) \quad (3)$$

چون حجم نمونه بیشتر از ۳۰ است از توزیع نرمال استفاده می کنیم

$$\int_{-\infty}^{-1/96} f(z) = 0.025 \Rightarrow Z_{0.025} = 1/96$$

$$\mu: \bar{x} \pm Z_{\alpha} \frac{s}{\sqrt{n}} = 150 \pm 1/96 \times \frac{14}{\sqrt{49}} \Rightarrow \mu: 150 \pm 3/92 \Rightarrow 146.08 < \mu < 153.92$$

گزینه ۱ صحیح است

برآورد فاصله‌ای یک طرفه

کافی است به جای $\frac{\alpha}{2}$ از α استفاده کنیم

$$\mu: \bar{x} + z_{\alpha} \frac{\sigma}{\sqrt{n}} \quad \text{یک طرفه بالا}$$

$$\mu: \bar{x} - z_{\alpha} \frac{\sigma}{\sqrt{n}} \quad \text{یک طرفه پائین}$$

ارشد مدیریت ۱۴۰۰

۶۲- از یک نمونه تصادفی ۱۶ تایی از دستگاه پرکننده شیشه‌های نوشابه، میانگین و واریانس به ترتیب ۲۵۰ و ۱۰۰ میلی لیتری حاصل شده است. با توجه به کارکرد غیر نرمال دستگاه، در سطح اطمینان ۹۶٪ میانگین مقدار

نوشابه در شیشه‌های پر شده توسط دستگاه چقدر برآورد می‌شود؟

$$[235.8, 264.2] \quad (1)$$

$$[245, 265] \quad (2)$$

$$[150, 350] \quad (3)$$

$$[237.5, 262.5] \quad (4)$$

چون گفته شده غیر نرمال از قضیه چپیشف حل می کنیم

$$P(|x - \mu| < k\sigma) \geq 1 - \frac{1}{k^2} \Rightarrow P\left(|\bar{x} - \mu| < k \frac{s}{\sqrt{n}}\right) \geq 1 - \frac{1}{k^2}$$

$$1 - \frac{1}{k^2} = 0.96 \Rightarrow \frac{1}{k^2} = 1 - 0.96 = 0.04 \Rightarrow k^2 = \frac{1}{0.04} = \frac{100}{4} = 25 \Rightarrow k = 5$$

$$\mu: \bar{x} \pm k \frac{s}{\sqrt{n}} \Rightarrow \mu: 250 \pm 5 \times \frac{10}{\sqrt{16}} = 250 \pm 12.5 \Rightarrow 237.5 \leq \mu \leq 262.5$$

برآورد فاصله‌ای نسبتها:

برای نسبتها اغلب از توزیع Z استفاده می‌شود.

$$\bar{P} = \frac{x}{n}, \quad \bar{q} = 1 - \bar{P}, \quad \sigma_p = \sqrt{\frac{\bar{P}\bar{q}}{n}}$$

$$P: \bar{P} \pm Z_{\frac{\alpha}{2}} \sqrt{\frac{\bar{P}\bar{q}}{n}}$$

مثال: از بین ۱۰۰ کالا ۲۰ کالا معیوب یک برآورد فاصله‌ای در سطح خطای ۵٪ برای نسبت کالاهای معیوب بنویسید

$$\bar{P} = \frac{x}{n} = \frac{20}{100} = 0.2 \quad q = 1 - \bar{P} = 1 - 0.2 = 0.8$$

$$P: \bar{P} \pm Z_{\frac{\alpha}{2}} \sqrt{\frac{\bar{P}\bar{q}}{n}} = 0.2 \pm 1.96 \sqrt{\frac{0.2 \times 0.8}{100}} \Rightarrow 0.1216 < P < 0.2784$$

برآورد فاصله‌ای واریانس و انحراف معیار:

همانگونه که می‌دانیم $\frac{(n-1)s^2}{\sigma^2}$ دارای توزیع خی دو با درجه آزادی $df = n-1$ است

$$\frac{(n-1)s^2}{\sigma^2} \sim \chi^2$$

$$\sigma_u^2 = \frac{(n-1)s^2}{\chi_{1-\frac{\alpha}{2}, df}^2}, \quad \sigma_L^2 = \frac{(n-1)s^2}{\chi_{\frac{\alpha}{2}, df}^2}, \quad \sigma_L^2 < \sigma^2 < \sigma_u^2$$

$$\sigma_u = s \sqrt{\frac{n-1}{\chi_{1-\frac{\alpha}{2}, df}^2}}, \quad \sigma_L = s \sqrt{\frac{n-1}{\chi_{\frac{\alpha}{2}, df}^2}}, \quad \sigma_L < \sigma < \sigma_u$$

مثال: در یک نمونه ۱۰ تایی $S^2 = ۸۰$ بدست آمده است یک برآورد فاصله‌ای ۹۰٪ برای واریانس واقعی جامعه بنویسید.

$$df = n - 1 = 10 - 1 = 9$$

$$\alpha = 1 - 0.9 = 0.1$$

$$\chi^2_{0.05,9} = 16/919$$

$$\chi^2_{0.95,9} = 3/325$$

$$\frac{\alpha}{2} = \frac{0.1}{2} = 0.05 \Rightarrow \chi^2_{0.05,9} = 16/919 \Rightarrow \sigma_L^2 = \frac{(n-1)s^2}{\chi^2_{\frac{\alpha}{2},df}} = \frac{9 \times 80}{16/919} = 42/555$$

$$1 - \frac{\alpha}{2} = 1 - 0.05 = 0.95 \Rightarrow \chi^2_{0.95,9} = 3/325 \Rightarrow \sigma_U^2 = \frac{(n-1)s^2}{\chi^2_{1-\frac{\alpha}{2},df}} = \frac{9 \times 80}{3/325} = 216/54$$

برآورد فاصله‌ای اختلاف میانگین دو جامعه نرمال مستقل وقتی واریانس جوامع معلوم باشد در این صورت بدون توجه به حجم نمونه از توزیع Z استفاده می‌شود.

$$\mu_2 - \mu_1 : \bar{X}_2 - \bar{X}_1 \pm Z_{\frac{\alpha}{2}} \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$$

و اگر عدد صفر در بازه بدست آمده قرار گیرد یعنی اختلاف بر اثر عامل تصادفی است

مثال: از دو جامعه نرمال مستقل با واریانس $\sigma_1^2 = 60$ و $\sigma_2^2 = 45$ نمونه‌های $n_1 = 10$ و $n_2 = 10$ گرفته‌ایم که $\bar{X}_2 = 45$ و $\bar{X}_1 = 40$ بدست آمده یک برآورد فاصله‌ای ۹۵٪ برای اختلاف میانگین دو جامعه بنویسید و توضیح دهد اختلاف بر اثر عامل تصادفی است.

$$\mu_2 - \mu_1 : \bar{X}_2 - \bar{X}_1 \pm Z_{\frac{\alpha}{2}} \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$$

$$\mu_2 - \mu_1 : 45 - 40 \pm 1/96 \sqrt{\frac{60}{10} + \frac{45}{10}} = 5 \pm 5/88 \Rightarrow -0/88 < \mu_2 - \mu_1 < 10/88$$

چون عدد صفر در بازه قرار می‌گیرد بنابراین اختلاف بر اثر عامل تصادفی است.

برآورد فاصله اختلاف میانگین دو جامعه نرمال مستقل وقتی واریانس جوامع نامعلوم باشد:

با فرض برابری واریانسها دو حالت امکان دارد رخ دهد

الف) حجم نمونه‌ها بیشتر از ۳۰ باشد: در این صورت طبق قضیه حد مرکزی می‌توان از توزیع Z با استفاده از واریانس نمونه‌ها برآورد فاصله‌ای را نوشت

$$\mu_2 - \mu_1 : \bar{X}_2 - \bar{X}_1 \pm Z_{\frac{\alpha}{2}} \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}$$

ب) حجم نمونه‌ها کمتر از ۳۰ باشد: در این از توزیع t و درجه آزادی $df = n_1 + n_2 - 2$ استفاده می‌شود

$$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}, \quad s_p = \sqrt{s_p^2}$$

$$\mu_2 - \mu_1: \bar{x}_2 - \bar{x}_1 \pm t_{\frac{\alpha}{2}, df} \cdot s_p \cdot \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}$$

مثال: یک برآورد فاصله‌ای ۹۰٪ برای اختلاف میانگین دو جامعه A و B که نمونه‌های زیر از آنها بدست آمده است بنویسید.

A	B
$\bar{X}_1 = 100$	$\bar{X}_2 = 95$
$s_1^2 = 90$	$s_2^2 = 100$
$n_1 = 8$	$n_2 = 14$

$$df = n_1 + n_2 - 2 = 8 + 14 - 2 = 20$$

$$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2} = \frac{7 \times 90 + 13 \times 100}{20} = 96/5$$

$$s_p = \sqrt{s_p^2} = \sqrt{96/5} = 9/82$$

$$\mu_2 - \mu_1: \bar{x}_2 - \bar{x}_1 \pm t_{\frac{\alpha}{2}, df} \cdot s_p \cdot \sqrt{\frac{1}{n_1} + \frac{1}{n_2}} \Rightarrow$$

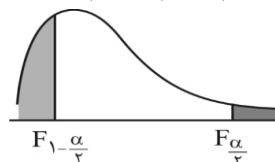
$$\mu_2 - \mu_1: 100 - 95 \pm 1/725 \times 9/82 \sqrt{\frac{1}{8} + \frac{1}{14}} = 5 \pm 3/279 \Rightarrow 1/72 < \mu_2 - \mu_1 < 8/279$$

بنابراین اختلاف بین میانگین دو جامعه واقعی است.

برآورد فاصله‌ای نسبت دو واریانس:

همان‌گونه که می‌دانیم نسبت دو واریانس دارای توزیع F یا فیشر است که چوله به راست است

با درجات آزادی $v_1 = df_1 = n_1 - 1$, $v_2 = df_2 = n_2 - 1$



$$v_1 = n_1 - 1, v_2 = n_2 - 1, F = \frac{\sigma_1^2}{\sigma_2^2} \cdot \frac{s_1^2}{s_2^2} \Rightarrow P \left[F_{1-\frac{\alpha}{2}, v_1, v_2} < F < F_{\frac{\alpha}{2}, v_1, v_2} \right] = 1 - \alpha$$

$$\Rightarrow \frac{1}{F_{\frac{\alpha}{2}, v_1, v_2}} \cdot \frac{s_1^2}{s_2^2} < \frac{\sigma_1^2}{\sigma_2^2} < \frac{s_1^2}{s_2^2} \cdot \frac{1}{F_{1-\frac{\alpha}{2}, v_1, v_2}}$$

$$\frac{1}{F_{\frac{\alpha}{2}, v_1, v_2}} \cdot \frac{s_1^2}{s_2^2} < \frac{\sigma_1^2}{\sigma_2^2} < \frac{s_1^2}{s_2^2} \cdot F_{\frac{\alpha}{2}, v_2, v_1} \quad \text{بنابراین:} \quad F_{1-\frac{\alpha}{2}, v_1, v_2} = \frac{1}{F_{\frac{\alpha}{2}, v_2, v_1}}$$

مثال: یک فاصله اطمینان ۹۰٪ برای نسبت واریانس جامعه A و B بدست آورید و توضیح دهید آیا واریانس دو جامعه با احتمال ۹۰٪ برابر است

A	B
$n_1 = 5$	$n_2 = 7$
$s_1^2 = 25$	$s_2^2 = 13/5$
$\bar{X}_1 = 13$	$\bar{X}_2 = 15/5$

$$\frac{1}{F_{\frac{\alpha}{2}, v_1, v_2}} \cdot \frac{s_1^2}{s_2^2} < \frac{\sigma_1^2}{\sigma_2^2} < \frac{s_1^2}{s_2^2} \cdot F_{\frac{\alpha}{2}, v_2, v_1}$$

$$\frac{1}{F_{0.05, 4, 6}} \cdot \frac{25}{13/3} < \frac{\sigma_1^2}{\sigma_2^2} < \frac{25}{13/3} \cdot F_{0.05, 6, 4}$$

$$\frac{1}{4/5337} \times \frac{25}{13/5} < \frac{\sigma_1^2}{\sigma_2^2} < \frac{25}{13/5} \cdot 6/1631$$

$$0/4084 < \frac{\sigma_1^2}{\sigma_2^2} < 11/4098$$

چون عدد یک در بازه بدست آمده قرار می گیرد بنابراین با احتمال ۹۰٪ و واریانس می تواند برابر باشد.
بر آورد فاصله ای اختلاف دو نسبت:

در صورتی که از n_1 عضو جامعه اول X_1 دارای خاصیتی و از n_2 عضو جامعه دوم X_2 عضو دارای آن خاصیت باشد در این به دو صورت می توان بر آورد فاصله ای برای اختلاف نسبتی نوشت

روش اول

$$\bar{P}_1 = \frac{X_1}{n_1}, \quad \bar{q}_1 = 1 - \bar{P}_1, \quad \bar{P}_2 = \frac{X_2}{n_2}, \quad \bar{q}_2 = 1 - \bar{P}_2$$

$$P_1 - P_2 : \bar{P} - \bar{P}_2 \pm Z_{\frac{\alpha}{2}} \sqrt{\frac{\bar{P}_1(1-\bar{P}_1)}{n_1} + \frac{\bar{P}_2(1-\bar{P}_2)}{n_2}}$$

روش دوم

$$P = \frac{X_1 + X_2}{n_1 + n_2}, \quad q = 1 - P$$

$$P_1 - P_2 : \bar{P} - \bar{P}_2 \pm Z_{\frac{\alpha}{2}} \sqrt{Pq \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}$$

مثال: کالایی توسط دو دستگاه A و B تولید می‌شود که از ۱۰۰ کالای دستگاه A ۱۵ کالا معیوب و ۶۰ کالای دستگاه B، ۶ کالا معیوب است یک برآورد فاصله‌ای ۸۵٪ برای اختلاف نسبت کالاهای معیوب توسط دو دستگاه A و B بنویسید

$$\bar{P}_A = \frac{x_A}{n_A} = \frac{15}{100} = 0.15, \quad \bar{P}_B = \frac{x_B}{n_B} = \frac{6}{60} = 0.1$$

$$P_A - P_B : \bar{P}_A - \bar{P}_B \pm Z_{\frac{\alpha}{2}} \sqrt{\frac{\bar{P}_A(1-\bar{P}_A)}{n_A} + \frac{\bar{P}_B(1-\bar{P}_B)}{n_B}}$$

$$P_A - P_B : 0.15 - 0.1 \pm 1.96 \sqrt{\frac{0.15 \times 0.85}{100} + \frac{0.1 \times 0.9}{60}}$$

$$0.05 - 0.1 < P_A - P_B < 0.05 + 0.1 \Rightarrow -0.05 < P_A - P_B < 0.15$$

بنابراین اختلاف نسبت کالاهای معیوب دو دستگاه با احتمال ۹۵٪ بر اثر عامل تصادفی است

آزمون فرض‌های آماری:

هرگاه ادعایی مطرح شود از آزمون فرض‌های آماری استفاده می‌شود آنچه ادعا شده اغلب H_1 در نظر گرفته می‌شود و نقیض آن در H_0 قرار می‌گیرد که در نتیجه تصمیم گرفته شده و واقعیت امکان دارد دو نوع خطا رخ دهد

۱- خطای نوع اول (α): احتمال رد H_0 در صورتی که H_0 درست باشد و با حرف α نشان می‌دهیم
 $\alpha = P(H_0 \text{ درست باشد} | H_0)$

۲- خطای نوع دوم (β): احتمال پذیرش H_0 در صورتی که H_0 نادرست باشد و با حرف β نشان می‌دهیم
 $\beta = P(H_0 \text{ نادرست باشد} | H_0)$

۳- توان آزمون (π): احتمال رد H_0 در صورتی که H_0 نادرست باشد و با حروف π نشان می‌دهیم. توان آزمون متمم خطای نوع دوم است $\pi = 1 - \beta$ و $\pi = P(H_0 \text{ نادرست باشد} | H_0)$

نکته ۱: با افزایش حجم نمونه خطای نوع اول و دوم کاهش می‌یابد
 نکته ۲: مجموع خطای نوع اول و دوم ثابت است یعنی با افزایش خطای نوع اول خطای نوع دوم کاهش می‌یابد و بالعکس
 نکته ۳: در ساخت فرضیه علامت مساوی همیشه در H_0 قرار دارد یعنی کلمات حداقل - حداکثر و برابر

نتایج ممکن آزمون فرض

وضعیت / تصمیم	H_0 پذیرفته می‌شود	H_0 رد می‌شود
اگر H_0 درست باشد	تصمیم درست سطح اطمینان $= 1 - \alpha$ = احتمال	خطای نوع اول سطح آزمون $= \alpha$ = احتمال
اگر H_0 نادرست باشد	خطای نوع دوم احتمال $= \beta$	تصمیم درست توان آزمون $= 1 - \beta = \pi$ = احتمال

توجه ۱: هدف از محاسبه آماره ها بر اساس داده های حاصل از نمونه توصیف نمونه و آزمون فرض پیرامون ویژگی های جامعه است.

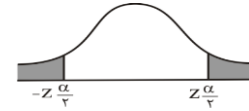
توجه ۲: اگر توان آزمون بالا برود خطای نوع اول افزایش می‌یابد چون توان آزمون از رابطه $1 - \beta$ بدست می‌آید بنابر این (توان آزمون) $\beta = 1 -$ در نتیجه اگر توان آزمون بالا برود خطای نوع دوم کاهش می‌یابد و چون در یک آزمون مجموع خطای نوع اول و دوم ثابت است با کاهش خطای نوع دوم خطای نوع اول افزایش می‌یابد
 مراحل انجام آزمون فرض:

- ۱-نوشتن فرض آماری
- ۲-تعیین آماره آزمون و محاسبه مقدار آن
- ۳-تعیین مقدار بحرانی از روی جدول
- ۴-تصمیم گیری

آزمون فرض میانگین با مقدار ثابت از جامعه نرمال و واریانس جامعه معلوم باشد.

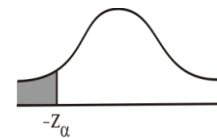
در این صورت از توزیع Z استفاده می شود و آماره آزمون از رابطه $Z = \frac{\bar{X} - \mu_0}{\frac{\sigma}{\sqrt{n}}}$ بدست می آید.

$$\text{آزمون دو دامنه} \begin{cases} H_0: \mu = \mu_0 \\ H_1: \mu \neq \mu_0 \end{cases} \quad Z > Z_{\frac{\alpha}{2}} \text{ یا } Z < -Z_{\frac{\alpha}{2}} \Rightarrow H_0 \text{ رد می شود}$$



$$H_0 \text{ رد می شود} \Rightarrow |Z| > Z_{\frac{\alpha}{2}} \text{ یا}$$

$$\text{آزمون یک طرفه چپ} \begin{cases} H_0: \mu \geq \mu_0 \\ H_1: \mu < \mu_0 \end{cases} \quad Z < -Z_{\alpha} \Rightarrow H_0 \text{ رد می شود}$$



$$\text{آزمون یک طرفه راست} \begin{cases} H_0: \mu \leq \mu_0 \\ H_1: \mu > \mu_0 \end{cases} \quad Z > Z_{\alpha} \Rightarrow H_0 \text{ رد می شود}$$



نکته: هنگام نوشتن فرضیه در صورتی که پیشینه کافی نداشته آزمون دودامنه می نویسیم

مثال: از یک جامعه نرمال با واریانس ۶۴ و میانگین ۸۰ یک نمونه ۲۵ تایی گرفته ایم که میانگین نمونه $\bar{X} = ۸۰$ بدست آمده است در سطح خطای ۵٪ آزمون کنید میانگین جامعه بیشتر از ۸۲ است.

$$\begin{cases} H_0: \mu \leq ۸۲ \\ H_1: \mu > ۸۲ \end{cases} \quad Z = \frac{\bar{X} - \mu_0}{\frac{\sigma}{\sqrt{n}}} = \frac{۸۰ - ۸۲}{\frac{۸}{\sqrt{۲۵}}} = -۱/۲۵$$

چون $Z = -۱/۲۵ < Z_{۰.۰۵} = ۱/۶۴۵$ بنابراین H_0 پذیرفته می شود و با احتمال ۹۵٪ حداکثر میانگین جامعه ۸۲ است

نکته: در آزمونهای یک طرفه به جای آلفا از نصف آلفا استفاده می شود

P - مقدار (P-VALUE): نشانه احتمالی است که ما مشاهده خواهیم کرد اگر فرض صفر صحیح باشد. در هنگام استفاده از P -مقدار اگر P -مقدار کمتر از آلفا تعیین شده باشد فرض صفر رد میشود

نکته: در استفاده از فاصله اطمینان برای بررسی معنی داری آزمون های برابری میانگین دو جامعه اگر عدد صفر در فاصله اطمینان ۹۵٪ قرار گیرد فرض تایید می شود و $p > ۰/۰۵$ است

$$\text{محاسبه حجم نمونه: } n = \frac{(Z_{\frac{\alpha}{2}})^2 \times \sigma^2}{d^2} \Leftrightarrow n = \left(\frac{(Z_{\frac{\alpha}{2}}) \times \sigma}{\frac{d}{2}} \right)^2 \Rightarrow d = Z_{\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}} \text{ درحالتی که}$$

$$n = \frac{(Z_{\alpha} + Z_{1-\beta})^2 * \sigma^2}{d^2}$$

خطای نوع دوم در نظر گرفته شود

مثال: در صورتیکه خطای نوع اول و دوم هر دو برابر ۰/۵ باشد توان آزمون را بدست آورید

۰/۵ (۱) ۰/۹۵ (۲) ۰/۲۵ (۳) ۰/۰۵ (۴)

گزینه ۱ صحیح است چون $1 - \beta = 1 - 0.5 = 0.5$ توان آزمون

ارشد مدیریت ۹۹

۵۵- در یک جامعه آماری با واریانس ۱، اگر هدف برآورد میانگین جامعه آماری باشد، با دقت برآورد ۰/۱۰، حجم نمونه

تحقیق در سطح خطای ۵٪ چقدر است؟

$$\left(\int_{-\infty}^{-1/96} f(z) dz = 0.025 \right)$$

۳۸۴ (۱)

۱۹۶ (۲)

۹۵ (۳)

۴۹ (۴)

$$n = \left(\frac{(Z_{\alpha}) \times \sigma}{d} \right)^2 \Rightarrow n = \left(\frac{1/96 \times 1}{0.1} \right)^2 = 384$$

گزینه ۱ صحیح است

ارشد مدیریت ۱۴۰۰

۶۱- فاصله اطمینان برای میانگین وزن محصولات یک شرکت با احتمال ۹۵٪ براساس نمونه‌های تصادفی،

(۴۹/۸، ۳۰/۲) محاسبه شده است. اگر توزیع وزن محصولات این شرکت نرمال باشد و انحراف معیار وزن

محصولات این شرکت ۳۰ باشد، حجم نمونه تصادفی این تخمین چقدر بوده است؟

(راهنمایی: $[(z | \alpha = 0.05) = 1.96 ; (z | \alpha = 0.025) = 1.96]$)

۲۵ (۱)

۳۶ (۲)

۶۴ (۳)

۱۴۴ (۴)

$$UCL = 49/8 \quad LCL = 30/2 \quad \sigma = 30$$

$$d = \frac{UCL - LCL}{2} = \frac{49/8 - 30/2}{2} = 9/8$$

$$n = \left(\frac{(Z_{\alpha}) \times \sigma}{d} \right)^2 \Rightarrow n = \left(\frac{1/96 \times 30}{9/8} \right)^2 = 36$$

گزینه ۲ صحیح است

آزمون فرض میانگین با مقدار ثابت از جامعه نرمال وقتی واریانس جامعه نامعلوم باشد: در این صورت از

واریانس نمونه به جای واریانس جامعه استفاده می‌کنیم و بسته شرایط از توزیع z یا t استفاده می‌شود.

الف) حجم نمونه بیشتر از ۳۰ باشد:

$$\begin{cases} H_0: \mu = \mu_0 \\ H_1: \mu \neq \mu_0 \end{cases} \quad |z| > z_{\frac{\alpha}{2}} \Rightarrow \text{رد می شود } H_0$$

$$\begin{cases} H_0: \mu \geq \mu_0 \\ H_1: \mu < \mu_0 \end{cases} \quad z < -z_{\alpha} \Rightarrow \text{رد می شود } H_0 \quad \text{آماره آزمون } z = \frac{\bar{x} - \mu_0}{\frac{s}{\sqrt{n}}}$$

$$\begin{cases} H_0: \mu \leq \mu_0 \\ H_1: \mu > \mu_0 \end{cases} \quad z > z_{\alpha} \Rightarrow \text{رد می شود } H_0$$

ب) حجم نمونه کمتر از ۳۰ باشد: در این صورت از توزیع t استیودنت با درجه آزادی $df = n - 1$ استفاده می کنیم

$$\begin{cases} H_0: \mu = \mu_0 \\ H_1: \mu \neq \mu_0 \end{cases} \quad |t| > t_{\frac{\alpha}{2}, df} \Rightarrow \text{رد می شود } H_0$$

$$\begin{cases} H_0: \mu \geq \mu_0 \\ H_1: \mu < \mu_0 \end{cases} \quad t < -t_{\alpha, df} \Rightarrow \text{رد می شود } H_0 \quad \text{آماره آزمون } t = \frac{\bar{x} - \mu_0}{\frac{s}{\sqrt{n}}}$$

$$\begin{cases} H_0: \mu \leq \mu_0 \\ H_1: \mu > \mu_0 \end{cases} \quad t > t_{\alpha, df} \Rightarrow \text{رد می شود } H_0$$

مثال: در یک نمونه ۱۶ تایی میانگین نمونه ۴۳ و واریانس نمونه ۴۹ بدست آمده است در سطح خطای ۵٪ آزمون کنید آیا میانگین جامعه ۴۰ است. $df = n - 1 = 16 - 1 = 15$ $n = 16 < 30 \rightarrow t$

$$\begin{cases} H_0: \mu = 40 \\ H_1: \mu \neq 40 \end{cases} \quad t = \frac{\bar{x} - \mu_0}{\frac{s}{\sqrt{n}}} = \frac{43 - 40}{\frac{9}{\sqrt{16}}} = 1/333$$

چون $t = 1/333 \leq t_{0.025, 15} = 2/131$ ، بنابراین H_0 پذیرفته می شود و با توجه به نمونه با احتمال ۹۵٪ میانگین جامعه ۴۰ است.

ارشد مدیریت ۱۴۰۱

۶۴- برای آزمون $H_0: \mu \leq 15$ در یک جامعه نرمال از یک نمونه تصادفی $n = 25$ ، $\bar{x} = 16$ و $S^2 = 0.16$ به دست آمده است. مقدار آماره آزمون کدام است؟

۰/۴ (۱)

۵ (۲)

۷/۵ (۳)

۱۲/۵ (۴)

$$n=25 \quad \bar{x}=16 \quad s^2=0.16 \Rightarrow s=0.4$$

$$\begin{cases} H_0: \mu \leq 15 \\ H_1: \mu > 15 \end{cases} \quad t = \frac{\bar{x} - \mu_0}{\frac{s}{\sqrt{n}}} = \frac{16 - 15}{\frac{0.4}{\sqrt{25}}} = 12.5$$

گزینه ۴ صحیح است

ارشد مدیریت ۱۴۰۲

۵۸- ادعا شده است: «میانگین جامعه‌ای برابر θ است». اگر از این جامعه نرمال نمونه تصادفی به حجم n انتخاب کنیم، آماره آزمون مناسب دارای چه توزیعی است؟

$$\chi_n^2 \quad (1)$$

$$\chi_{n-1}^2 \quad (2)$$

$$t_{n-1} \quad (3)$$

$$t_n \quad (4)$$

گزینه ۳ صحیح است

آزمون فرض نسبت: اغلب از توزیع z استفاده می‌شود

$$\begin{cases} H_0: P = P_0 \\ H_1: P \neq P_0 \end{cases} \quad |z| > z_{\frac{\alpha}{2}} \Rightarrow H_0 \text{ رد می‌شود}$$

$$\begin{cases} H_0: P \geq P_0 \\ H_1: P < P_0 \end{cases} \quad z < -z_{\alpha} \Rightarrow H_0 \text{ رد می‌شود} \quad \text{آماره آزمون } Z = \frac{P - P_0}{\sqrt{\frac{P_0(1-P_0)}{n}}}$$

$$\begin{cases} H_0: P \leq P_0 \\ H_1: P > P_0 \end{cases} \quad z > z_{\alpha} \Rightarrow H_0 \text{ رد می‌شود}$$

مثال: در یک نمونه ۱۰۰ تایی تعداد ۱۸ کالا معیوب است در سطح خطای ۵٪ آزمون کنید نسبت کالای معیوب کمتر از ۲۰٪ است؟

$$P = \frac{x}{n} = \frac{18}{100} = 0.18$$

$$\begin{cases} H_0: P \geq 20\% \\ H_1: P < 20\% \end{cases} \quad Z = \frac{P - P_0}{\sqrt{\frac{P_0(1-P_0)}{n}}} = \frac{0.18 - 0.2}{\sqrt{\frac{0.2 \times 0.8}{100}}} = -0.5$$

چون $z = -0.5 < -z_{0.05} = -1.645$ ، H_0 پذیرفته می‌شود و با توجه به نمونه با احتمال ۹۵٪ نسبت کالای معیوب حداقل ۲۰٪ است.

ارشد مدیریت ۱۴۰۰

۶۳ ادعا شده است که نرخ بیکاری جوانان کشور کمتر از ۱۴ درصد است، فرضیه آماری H_0 کدام است؟

$$P > 0.14 \quad (1)$$

$$P \leq 0.14 \quad (2)$$

$$P < 0.14 \quad (3)$$

$$P \geq 0.14 \quad (4)$$

گزینه ۴ صحیح است چون علامت مساوی همیشه در H_0 قرار می گیرد. $\begin{cases} H_0: P \geq 0.14 \\ H_1: P < 0.14 \end{cases}$

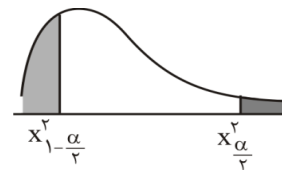
آزمون فرض واریانس با مقدار ثابت:

همانگونه که می دانیم $\frac{(n-1)s^2}{\sigma^2}$ دارای توزیع خی دو (χ^2) با درجه آزادی $v = df = n - 1$ است که این

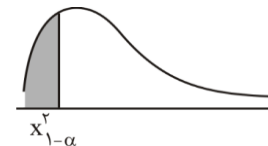
توزیع یک منحنی چوله به راست دارد آمار، آزمون در این آزمون از رابطه $\chi^2 = \frac{(n-1)s^2}{\sigma^2}$ بدست

می آید.

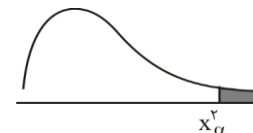
$$\begin{cases} H_0: \sigma^2 = \sigma_0^2 \\ H_1: \sigma^2 \neq \sigma_0^2 \end{cases} \quad \chi^2 > X^2_{\frac{\alpha}{2}, df} \Rightarrow \text{رد می شود } H_0, \quad \chi^2 \leq X^2_{1-\frac{\alpha}{2}, df} \Rightarrow \text{رد می شود } H_0.$$



$$\begin{cases} H_0: \sigma^2 \geq \sigma_0^2 \\ H_1: \sigma^2 < \sigma_0^2 \end{cases} \quad \chi^2 < X^2_{1-\alpha, df} \Rightarrow \text{رد می شود } H_0.$$



$$\begin{cases} H_0: \sigma^2 \leq \sigma_0^2 \\ H_1: \sigma^2 > \sigma_0^2 \end{cases} \quad \chi^2 > X^2_{\alpha, df} \Rightarrow \text{رد می شود } H_0.$$



مثال: در یک نمونه ۱۶ تایی میانگین نمونه ۷۵ و واریانس نمونه ۴۴ بدست آمده است در سطح خطای ۵٪.

آزمون کنید آیا واریانس جامعه بیش تر از ۴۰ است؟

$$\begin{cases} H_0: \sigma^2 \leq 40 \\ H_1: \sigma^2 > 40 \end{cases} \quad \chi^2 = \frac{(n-1)s^2}{\sigma_0^2} = \frac{15 \times 24}{40} = 16/5$$

چون $\chi^2 = 16/5 > \chi^2_{0.05, 15} = 24/996$ ، بنابراین با احتما ۹۵٪ با توجه به نمونه H_0 پذیرفته می شود و

واریانس جامعه حداقل ۴۰ است

ارشد مدیریت ۱۴۰۲

۵۹- به منظور آزمون فرضیه «پراکندگی وزن محصولات کارخانه، کمتر از ۶ است»، نمونه تصادفی ۲۰ تایی از این جامعه نرمال انتخاب شده است. با توجه به سطح زیر منحنی دنباله سمت راست توزیع χ^2 ارائه شده در جدول زیر، به ازای چه مقدار برای آماره آزمون فرضیه پژوهشی در سطح معناداری ۱۰ درصد رد می شود؟

$\alpha \backslash df$	۰/۹۵	۰/۹	۰/۱	۰/۵
۱۹	۱۰/۱	۱۱/۶	۲۷/۲	۳۰/۱
۲۰	۱۰/۸	۱۲/۴	۲۸/۴	۳۱/۴

$$\chi^2 \leq 28/4 \quad (1)$$

$$\chi^2 \leq 27/2 \quad (2)$$

$$\chi^2 \geq 12/4 \quad (3)$$

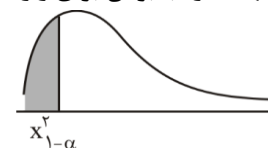
$$\chi^2 \geq 11/6 \quad (4)$$

جهت ناریب بودن برای واریانس آزمون انجام میشود.

$$n = 20 \Rightarrow df = n - 1 = 20 - 1 = 19, \quad 1 - \alpha = 1 - 0/1 = 0/9$$

$$\Rightarrow \chi^2_{0/9, 19} = 11/6$$

$$\begin{cases} H_0: \sigma \geq 6 \\ H_1: \sigma < 6 \end{cases} \Leftrightarrow \begin{cases} H_0: \sigma^2 \geq 36 \\ H_1: \sigma^2 < 36 \end{cases}$$



گزینه ۴ صحیح است

آزمون فرض اختلاف میانگین دو جامعه نرمال مستقل وقتی واریانس جوامع معلوم باشد: در این صورت

$$Z = \frac{\bar{X}_1 - \bar{X}_2 - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} \quad \text{بدون توجه به حجم نمونه از توزیع } Z \text{ استفاده می شود و آمار، آزمون از رابطه}$$

بدست می آید که اگر آزمون برابری میانگینها انجام شود به جای $\mu_1 - \mu_2$ عدد صفر را قرار می دهیم

$$\begin{cases} H_0: \mu_1 - \mu_2 = a \\ H_1: \mu_1 - \mu_2 \neq a \end{cases} \quad |Z| > \frac{Z_{\alpha}}{2} \Rightarrow \text{رد می شود } H_0$$

$$\begin{cases} H_0: \mu_1 - \mu_2 \geq a \\ H_1: \mu_1 - \mu_2 < a \end{cases} \quad Z < -Z_{\alpha} \Rightarrow \text{رد می شود } H_0 \quad \text{آماره آزمون } Z = \frac{\bar{X}_1 - \bar{X}_2 - a}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$$

$$\begin{cases} H_0: \mu_1 - \mu_2 \leq a \\ H_1: \mu_1 - \mu_2 > a \end{cases} \quad Z > Z_{\alpha} \Rightarrow \text{رد می شود } H_0$$

آزمون فرض اختلاف میانگین دو جامعه نرمال مستقل وقتی واریانس جوامع نامعلوم باشد. در این صورت با فرض برابر واریانسها دو حالت را در نظر می گیریم.

الف) حجم نمونه بیش از ۳۰ باشد. در این صورت از توزیع Z استفاده می شود.

$$\begin{cases} H_0: \mu_1 - \mu_2 = a \\ H_1: \mu_1 - \mu_2 \neq a \end{cases} \quad |Z| > \frac{Z_{\alpha}}{2} \Rightarrow \text{رد می شود } H_0$$

$$\begin{cases} H_0: \mu_1 - \mu_2 \geq a \\ H_1: \mu_1 - \mu_2 < a \end{cases} \quad Z = \frac{\bar{x}_1 - \bar{x}_2 - a}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} \quad \text{آماره آزمون}$$

$Z < -Z_{\alpha} \Rightarrow H_0$ رد می‌شود

$$\begin{cases} H_0: \mu_1 - \mu_2 \leq a \\ H_1: \mu_1 - \mu_2 > a \end{cases} \quad Z > Z_{\alpha} \Rightarrow H_0 \text{ رد می‌شود}$$

ب) حجم نمونه‌ها کمتر از ۳۰ باشد (t مستقل): در این صورت از توزیع t با درجه آزادی $df = n_1 + n_2 - 2$

$$\begin{cases} H_0: \mu_1 - \mu_2 = a \\ H_1: \mu_1 - \mu_2 \neq a \end{cases} \quad |t| > t_{\frac{\alpha}{2}, df} \Rightarrow H_0 \text{ رد می‌شود}$$

استفاده می‌شود.

$$df = n_1 + n_2 - 2$$

$$\begin{cases} H_0: \mu_1 - \mu_2 \geq a \\ H_1: \mu_1 - \mu_2 < a \end{cases} \quad t < -t_{\alpha, df} \Rightarrow H_0 \text{ رد می‌شود}$$

$$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}$$

$$t = \frac{\bar{x}_1 - \bar{x}_2 - a}{S_{\bar{x}_1 - \bar{x}_2}} \Rightarrow t = \frac{\bar{x}_1 - \bar{x}_2 - a}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$$

$$\begin{cases} H_0: \mu_1 - \mu_2 \leq a \\ H_1: \mu_1 - \mu_2 > a \end{cases} \quad t > t_{\alpha, df} \Rightarrow H_0 \text{ رد می‌شود}$$

یاد آوری: در استفاده از فاصله اطمینان برای بررسی معنی داری آزمون‌ها ی برابری میانگین دو جامعه اگر عدد صفر در فاصله اطمینان ۹۵٪ قرار گیرد فرض تایید می‌شود و $p > 0.05$ است

مثال: از دو جامعه نرمال مستقل نمونه‌های زیر بدست آمده است:

A	$\bar{x}_1 = 80$	$n_1 = 11$	$s_1^2 = 50$
B	$\bar{x}_2 = 85$	$n_2 = 16$	$s_2^2 = 60$

در سطح معنی داری ۵٪ آزمون کنید میانگین دو جامعه برابر است.

$$\begin{cases} H_0: \mu = \mu_0 \\ H_1: \mu \neq \mu_0 \end{cases}$$

$$df = n_1 + n_2 - 2 = 11 + 16 - 2 = 25$$

$$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2} = \frac{10 \times 50 + 15 \times 60}{25} = 56$$

$$t = \frac{\bar{x}_1 - \bar{x}_2 - a}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} = \frac{85 - 80}{\sqrt{56} \sqrt{\frac{1}{11} + \frac{1}{16}}} = +1.7$$

چون $|t| = 1.7 > t_{0.025, 25} = 2.06$ بنابراین H_0 پذیرفته می‌شود و با احتمال ۹۵٪ با توجه به نمونه میانگین دو جامعه برابر است.

مثال: در گزارش یک پژوهش نتایج آزمون t با دو گروه مستقل نتایج به صورت $t_{28} = 3.16$ گزارش شده است حجم نمونه چند نفر بوده است

(۱) ۱۲۷

(۲) ۱۲۸

(۳) ۱۲۹

(۴) ۱۳۰

گزینه ۴ صحیح است

$$df = n_1 + n_2 - 2 \Rightarrow 128 = n_1 + n_2 - 2 \Rightarrow n_1 + n_2 = 128 + 2 = 130$$

اندازه اثر: در صورتیکه حد بالای میانگین کوچکتر مثلاً گروه ۲ از حد پایین میانگین بزرگتر مثلاً گروه ۱ مقدار کوچکتری داشته باشد همپوشانی ندارند اندازه اثر بزرگتر نشان دهنده تفاوت بیشتر دو گروه است هر چه عدد اندازه اثر کوچکتر باشد همپوشانی بیشتر است یکی از مزیت‌های اندازه اثر مستقل بودن از حجم نمونه است.

آزمون فرض برابری میانگین‌های دو جامعه نرمال مستقل با فرض مجهول بودن واریانسها و برابر نبودن واریانسها در این صورت از توزیع t استفاده می‌شود با درجه آزادی

$$\begin{cases} H_0: \mu = \mu_0 \\ H_1: \mu \neq \mu_0 \end{cases} \quad |t'| > t_{\frac{\alpha}{2}, df} \Rightarrow \text{رد می‌شود } H_0$$

$$df = \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2} \right)^2}{\frac{\left(\frac{s_1^2}{n_1} \right)^2}{n_1 - 1} + \frac{\left(\frac{s_2^2}{n_2} \right)^2}{n_2 - 1}}$$

$$\begin{cases} H_0: \mu \geq \mu_0 \\ H_1: \mu < \mu_0 \end{cases} \quad t' < -t_{\alpha, df} \Rightarrow \text{رد می‌شود } H_0$$

$$t' = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$$

$$\begin{cases} H_0: \mu \leq \mu_0 \\ H_1: \mu > \mu_0 \end{cases} \quad t' > t_{\alpha, df} \Rightarrow \text{رد می‌شود } H_0$$

$$df = \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2} \right)^2}{\frac{\left(\frac{s_1^2}{n_1} \right)^2}{n_1 + 1} + \frac{\left(\frac{s_2^2}{n_2} \right)^2}{n_2 + 1}} - 2 \quad \text{در کتاب آمار مهندسی لیبرمن درجه آزادی به صورت}$$

مقایسه‌های زوج شده (t وابسته):

هرگاه افراد مورد آزمایش ثابت باشد و دو آزمایش بر روی این افراد انجام شود از مقایسه‌های زوج شده استفاده می‌شود مانند نمرات افراد قبل از رفتن به کلاس درس در یک درس خاص که از توزیع t با درجه آزادی $df = n - 1$ استفاده می‌شود برای انجام این آزمون اختلاف دو حالت را بدست آورده و بررسی

میانگینهای آزمون را انجام می دهیم

$$\begin{cases} H_0: \mu_d = a \\ H_1: \mu_d \neq a \end{cases} \quad |t| > t_{\alpha/2, df} \Rightarrow \text{رد می شود } H_0.$$

$$d_i = x_i - y_i \quad df = n - 1$$

$$\begin{cases} H_0: \mu_d \geq a \\ H_1: \mu_d < a \end{cases} \quad t < -t_{\alpha, df} \Rightarrow \text{رد می شود } H_0.$$

$$s_d^2 = \frac{\sum d_i^2 - n\bar{d}^2}{n-1}$$

$$t = \frac{\bar{d} - a}{\frac{s_d}{\sqrt{n}}}$$

$$\begin{cases} H_0: \mu_d \leq a \\ H_1: \mu_d > a \end{cases} \quad t > t_{\alpha, df} \Rightarrow \text{رد می شود } H_0.$$

مثال: سرعت خواندن ۶ دانشجو قبل از رفتن به کلاس تندخوانی و بعد از رفتن به کلاس در جدول زیر ثبت شده است آیا سرعت خواندن به طور متوسط افزایش یافته است. (در سطح خطای ۰.۰۵٪)

قبل (x)	۲	۲/۵	۴	۵	۷	۶
بعد (y)	۳	۳	۴/۵	۱۰	۱۲	۹

$d_x = y_i - x_i$	۱	۰/۵	۰/۵	۵	۵	۳	۱۵
d_i^2	۱	۰/۲۵	۰/۲۵	۲۵	۲۵	۹	۶۰/۵

$$df = n - 1 = 6 - 1 = 5 \quad \bar{d} = \frac{\sum x_i}{n} = \frac{15}{6} = 2.5$$

$$s_d^2 = \frac{\sum x_i^2 - n\bar{d}^2}{n-1} = \frac{60.5 - 6 \times (2.5)^2}{6-1} = 4/6$$

$$t = \frac{\bar{d} - a}{\frac{s_d}{\sqrt{n}}} = \frac{2.5}{\frac{\sqrt{4/6}}{\sqrt{6}}} = 2.87$$

چون $t = 2.87 > t_{0.05, 5} = 2.015$ بنابراین H_0 رد می شود و با احتمال ۰.۰۵٪ و با توجه به $\begin{cases} H_0: \mu_d \leq 0 \\ H_1: \mu_d > 0 \end{cases}$

نمونه رفتن به کلاس مؤثر بوده است.

مثال: در یک آزمون T وابسته به حجم ۹ اگر میانگین اختلاف نمرات قبل و بعد از آزمون ۰/۵ با انحراف

استاندارد اختلاف نمرات ۰/۵ باشد مقدار آماره آزمون را بدست آورید. $t = \frac{\bar{d} - a}{\frac{s_d}{\sqrt{n}}} = \frac{0.5}{\frac{0.5}{\sqrt{9}}} = \sqrt{9} = 3$

مثال: در یک نمونه ۱۰۰ تایی اگر میانگین اختلاف نمرات قبل و بعد از آزمون ۸ با انحراف استاندارد اختلاف

نمرات ۱۰ باشد یک برآورد فاصله ای ۰.۰۵٪ برای اختلاف نمرات بنویسید. $t_{0.025, 99} \approx 2$

$$t = \frac{\bar{d} - a}{\frac{s_d}{\sqrt{n}}} \Rightarrow d: \bar{d} \pm t_{\alpha/2, df} \cdot \frac{s_d}{\sqrt{n}} \Rightarrow d: 10 \pm 2 \times \frac{1}{\sqrt{100}} = 10 \pm 1/6 \Rightarrow 8/4 \leq d \leq 11/6$$

آزمون فرض اختلاف نسبتها:

از توزیع z استفاده می‌کنیم

$$\begin{cases} H_0: P_1 - P_2 = a \\ H_1: P_1 - P_2 \neq a \end{cases} \quad |z| > z_{\frac{\alpha}{2}} \Rightarrow \text{رد می‌شود } H_0$$

$$P_1 = \frac{x_1}{n_1} \quad P_2 = \frac{x_2}{n_2}$$

$$z = \frac{P_1 - P_2 - a}{\sqrt{\frac{P_1(1-P_1)}{n_1} + \frac{P_2(1-P_2)}{n_2}}}$$

$$z = \frac{\hat{p}_1 - \hat{p}_2 - a}{\sqrt{\hat{p}(1-\hat{p}) \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}$$

$$\begin{cases} H_0: P_1 - P_2 \geq a \\ H_1: P_1 - P_2 < a \end{cases} \quad z < -z_{\alpha} \Rightarrow \text{رد می‌شود } H_0$$

$$\begin{cases} H_0: P_1 - P_2 \leq a \\ H_1: P_1 - P_2 > a \end{cases} \quad z > z_{\alpha} \Rightarrow \text{رد می‌شود } H_0$$

ارشد مدیریت ۱۴۰۱

۶۳- ادعا شده است، نسبت ضایعات در کارخانه (۱) کمتر از نسبت ضایعات در کارخانه (۲) است. فرضیه H_0 کدام است؟

$$H_0: p_1 \geq p_2 \quad (۱)$$

$$H_0: \mu_1 \geq \mu_2 \quad (۲)$$

$$H_0: p_1 \leq p_2 \quad (۳)$$

$$H_0: \mu_1 \leq \mu_2 \quad (۴)$$

گزینه ۱ صحیح است علامت مساوی همیشه در H_0 قرار می‌گیرد. $\begin{cases} H_0: P_1 - P_2 \geq 0 \\ H_1: P_1 - P_2 < 0 \end{cases} \Rightarrow \begin{cases} H_0: P_1 \geq P_2 \\ H_1: P_1 < P_2 \end{cases}$

آزمون فرض نسبت دو واریانس:

همانگونه که می‌دانیم نسبت دو واریانس دارای توزیع F یا فیشر با درجات آزادی $v_1 = n_1 - 1$ و

$v_2 = n_2 - 1$ است و آماره آزمون از رابطه $F = \left(\frac{s_1^2}{s_2^2} \right) \frac{s_1^2}{s_2^2}$ بدست می‌آید.

$$\begin{cases} H_0: \frac{\sigma_1^2}{\sigma_2^2} = a \\ H_1: \frac{\sigma_1^2}{\sigma_2^2} \neq a \end{cases} \quad F_{1-\frac{\alpha}{2}, v_1, v_2} < FF < F_{\frac{\alpha}{2}, v_1, v_2} \Rightarrow \text{رد می‌شود } H_0$$

$$\begin{cases} H_0: \frac{\sigma_1^2}{\sigma_2^2} \geq a \\ H_1: \frac{\sigma_1^2}{\sigma_2^2} < a \end{cases} \quad F < F_{1-\alpha, v_1, v_2} \Rightarrow \text{رد می شود } H_0.$$

$$\begin{cases} H_0: \frac{\sigma_1^2}{\sigma_2^2} \leq a \\ H_1: \frac{\sigma_1^2}{\sigma_2^2} > a \end{cases} \quad F > F_{\alpha, v_1, v_2} \Rightarrow \text{رد می شود } H_0.$$

نکته: آماره آزمون $F = a \cdot \frac{s_1^2}{s_2^2}$ در آزمون برابر واریانس‌ها $a = 1$ است.

مثال: در سطح خطای ۱۰٪ آزمون کنید آیا واریانس دو جامعه A و B برابر است

A	$n_1 = 10$	$\bar{X}_1 = 10.2$	$s_1^2 = 80$
B	$n_2 = 6$	$\bar{X}_2 = 10.5$	$s_2^2 = 50$

$$\begin{cases} H_0: \sigma_1^2 = \sigma_2^2 \\ H_1: \sigma_1^2 \neq \sigma_2^2 \end{cases} \quad F = \frac{s_1^2}{s_2^2} = \frac{80}{50} = 1.6$$

$$F_{\frac{\alpha}{2}, v_1, v_2} = F_{0.05, 9, 5} = 4.7725$$

$$F_{1-\frac{\alpha}{2}, v_1, v_2} = \frac{1}{F_{\frac{\alpha}{2}, v_2, v_1}} = F_{0.05, 5, 9} = \frac{1}{3.4817} = 0.2872 = F_{0.95, 9, 5}$$

چون $F = 1.6 < F_{0.95, 9, 5} = 0.2872$ یا $F = 1.6 > F_{0.05, 9, 5} = 4.7725$ بنابراین H_0 با توجه به نمونه با احتمال ۹۵٪ پذیرفته می‌شود و واریانس دو جامعه برابر است.

ارشد مدیریت ۹۹

۵۶- به منظور آزمون فرضیه «ریسک سرمایه‌گذاری در بورس کمتر از مسکن است»، نمونه‌هایی تصادفی از هر گروه انتخاب شده است که اطلاعات آن به صورت جدول زیر است:

انحراف معیار بازدهی	میانگین بازدهی	تعداد نمونه	نوع سرمایه‌گذاری
۴	۱۰۰	۲۵	بورس
۸	۶۴	۱۶	مسکن

با فرض نرمال بودن توزیع سود سرمایه‌گذاری در بورس و مسکن، آماره آزمون این فرضیه کدام است؟

(۱) ۰/۵۰

(۲) ۰/۲۵

(۳) ۱/۵۶

(۴) ۲

$$F = \frac{s_1^2}{s_2^2} = \frac{4^2}{8^2} = 0.25$$

گزینه ۲ صحیح است

آزمون کوکران برای همگنی واریانسها:

همانگونه که می دانیم برای آزمون ۲ واریانس از آزمون F استفاده می شود ولی اگر تعداد واریانس های مورد مقایسه بیش از ۲ باشد از آزمون کوکران استفاده می شود.

$$\begin{cases} H_0 : \sigma_1^2 = \sigma_2^2 = \dots = \sigma_K^2 \\ H_1 : \sigma_i^2 \neq \sigma_j^2 \end{cases}$$

$$G = \frac{s_{\max}^2}{s_1^2 + s_2^2 + \dots + s_k^2}$$

در صورتی که $G > g_\alpha$ باشد فرض H_0 رد می شود و حداقل یکی از σ_i^2 ها متفاوت است.

فصل ششم

ضریب همبستگی
رگرسیون

ضریب همبستگی: معیاری است برای تعیین ارتباط بین یک یا چند متغیر مستقل (x) با یک متغیر وابسته (y) در صورتی که ارتباط بین دو متغیر همسو باشد مانند قد و وزن ضریب همبستگی مثبت و در صورتی که ارتباط بین دو متغیر غیر همسو باشد مانند رابطه فشار هوا و فاصله از سطح زمین ضریب همبستگی منفی می شود در آمار پارامتریک از ضریب همبستگی پیرسن استفاده می کنیم که عددی بین -۱ و +۱ است و بدون واحد است ضریب همبستگی نمونه ای را با حرف r و ضریب همبستگی جامعه را با حرف ρ نشان می دهیم.

$$r = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2 \cdot \sum (y_i - \bar{y})^2}} \Rightarrow r = \frac{n \sum xy - (\sum x \cdot \sum y)}{\sqrt{\sum (x_i - \bar{x})^2 \cdot \sum (y_i - \bar{y})^2}}$$

$$S_{pxy} = S_{xy} = \sum (x_i - \bar{x})(y_i - \bar{y}) = \sum x_i y_i - n\bar{x}\bar{y} = \sum x_i y_i - \frac{\sum x_i \sum y_i}{n}$$

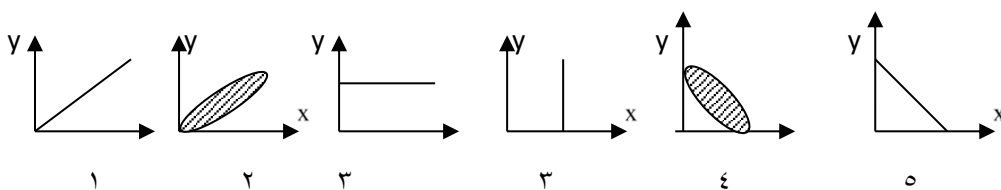
$$S_{xx} = \sum (x_i - \bar{x})^2 = \sum x_i^2 - n\bar{x}^2 = \sum x_i^2 - \frac{(\sum x_i)^2}{n}$$

$$S_{yy} = \sum (y_i - \bar{y})^2 = \sum y_i^2 - n\bar{y}^2 = \sum y_i^2 - \frac{(\sum y_i)^2}{n}$$

با جایگذاری در رابطه ضریب همبستگی می توان نوشت

$$r = \frac{S_{xy}}{\sqrt{S_{xx}S_{yy}}} \Rightarrow r = \frac{S_{xy}}{\sqrt{N\sigma_x^2 N\sigma_y^2}} \Rightarrow r = \frac{S_{xy}}{N\sigma_x\sigma_y}, \text{cov}_{xy} = \frac{S_{xy}}{n} \Rightarrow r = \frac{\text{cov}_{xy}}{\sigma_x\sigma_y}$$

$$|r| \leq 1 \Rightarrow -1 \leq r \leq +1 \Rightarrow \begin{cases} r = +1 & \text{رابطه خطی مستقیم کامل} \\ 0 < r < 1 & \text{رابطه مستقیم مستقیم} \\ r = 0 & \text{عدم همبستگی خطی} \\ -1 < r < 0 & \text{رابطه معکوس} \\ r = -1 & \text{رابطه خطی معکوس کامل} \end{cases}$$



توجه: متغیری که تغییرات آن اندازه گیری می شود را متغیر وابسته گویند.

ضریب تعیین: مجذور ضریب همبستگی را ضریب تعیین گویند که نشان دهنده آن است که چند درصد تغییرات متغیر وابسته توسط متغیر مستقل بیان می شود. $V = r_{xy}^2$ ضریب تعیین که اغلب بر حسب درصد بیان میشود یعنی $\%V = (r_{xy}^2) \cdot (100)$

مثال: ضریب همبستگی بین دو متغیر x, y را تعیین کنید و نوع آن را مشخص کنید

x	y	$x - \bar{x}$	$y - \bar{y}$	$(x - \bar{x})^2$	$(y - \bar{y})^2$	$(x - \bar{x}) \times (y - \bar{y})$
۲	۱۵	$۲ - ۶ = -۴$	$۱۵ - ۵ = ۱۰$	$(-۴)^2 = ۱۶$	$(۱۰)^2 = ۱۰۰$	$-۴ \times ۱۰ = -۴۰$
۴	۱۰	$۴ - ۶ = -۲$	$۱۰ - ۵ = ۵$	$(-۲)^2 = ۴$	$(۵)^2 = ۲۵$	$-۲ \times ۵ = -۱۰$
۶	۵	$۶ - ۶ = ۰$	$۵ - ۵ = ۰$	$(۰)^2 = ۰$	$(۰)^2 = ۰$	$۰ \times ۰ = ۰$
۸	۰	$۸ - ۶ = ۲$	$۰ - ۵ = -۵$	$(۲)^2 = ۴$	$(-۵)^2 = ۲۵$	$۲ \times -۵ = -۱۰$
۱۰	-۵	$۱۰ - ۶ = ۴$	$-۵ - ۵ = -۱۰$	$(۴)^2 = ۱۶$	$(-۱۰)^2 = ۱۰۰$	$۴ \times -۱۰ = -۴۰$
$\sum x = ۳۰$	$\sum y = ۲۵$	$\cdot$	$\cdot$	$SS_x = ۴۰$	$SS_y = ۲۵۰$	$s_{xy} = -100$

$$\bar{x} = \frac{\sum x}{n} = \frac{۳۰}{۵} = ۶, \quad \bar{y} = \frac{\sum y}{n} = \frac{۲۵}{۵} = ۵$$

$$r = \frac{s_{xy}}{\sqrt{SS_x \times SS_y}} = \frac{-۱۰۰}{\sqrt{۴۰ \times ۲۵۰}} = \frac{-۱۰۰}{\sqrt{۱۰۰۰۰}} = \frac{-۱۰۰}{۱۰۰} = -۱$$

همبستگی معکوس کامل

نکته ۱: به طور کلی اگر فواصل x, y مساوی باشد ضریب همبستگی یک است و علامت آن بستگی به حاصل ضرب تغییرات دارد.

x	y	x	y	x	y	x	y
۲	۱۵	۳	۴	۲	۱۰	۳	۵
۴	۱۰	۶	۸	۲	۱۵	۶	۵
۶	۵	۹	۱۲	۲	۲۰	۴۵	۵
۸	۰	۱۲	۱۶	۲	۳۰	۱۰	۵
۱۰	-۵	۱۵	۲۰	۲	۵۰	۷۰	۵
$r = -۱$		$r = ۱$		$r = ۰$		$r = ۰$	

آزمون فرض ضریب همبستگی با مقدار ثابت: این آزمون جهت نشان دادن میزان تاثیر متغیر مستقل بر متغیر وابسته را نشان می دهد که به دو روش صورت می گیرد

۱- با استفاده از آزمون t : آماره آزمون از رابطه $t = \frac{r\sqrt{n-2}}{\sqrt{1-r^2}}$ که نقطه بحرانی با درجه آزادی

$df = n - 2$ بدست می آید یعنی اگر $t > t_{\frac{\alpha}{2}, n-2}$ باشد فرض صفر رد می شود و ضریب همبستگی معنی دار است با اطمینان $p = 1 - \alpha$ و متغیر مستقل و وابسته رابطه دارند. در صورتیکه فرض صفر پذیرفته شود ($H_0: \rho = 0$) رابطه میان ضریب همبستگی و توان آماری ناودیسی هست.

$$\begin{cases} H_0: \rho = \rho_0 \\ H_1: \rho \neq \rho_0 \end{cases} \quad |t| > t_{\frac{\alpha}{2}, df} \Rightarrow \text{رد می شود } H_0$$

$$\begin{cases} H_0: \rho \geq \rho_0 \\ H_1: \rho < \rho_0 \end{cases} \quad t < -t_{\alpha, df} \Rightarrow \text{رد می شود } H_0$$

$$df = n - 2$$

$$\text{متغیر آزمون} = \frac{(r - \rho_0) \sqrt{n - 2}}{\sqrt{1 - r^2}}$$

$$\begin{cases} H_0: \rho \leq \rho_0 \\ H_1: \rho > \rho_0 \end{cases} \quad t > t_{\alpha, df} \Rightarrow \text{رد می شود } H_0$$

۲- با استفاده از توزیع z:

در این صورت با استفاده از تبدیل فیشر برای $n > 3$ ، $z = \frac{1}{2} \ln \frac{1+r}{1-r}$ دارای توزیع نرمال با میانگین

$$E(z) = \frac{1}{2} \ln \frac{1+\rho}{1-\rho} \quad \text{و واریانس} \quad \text{var}(z) = \frac{1}{n-3}$$

$$\begin{cases} H_0: \rho = \rho_0 \\ H_1: \rho \neq \rho_0 \end{cases} \quad |z^*| > z_{\frac{\alpha}{2}} \Rightarrow \text{رد می شود } H_0$$

$$\begin{cases} H_0: \rho \geq \rho_0 \\ H_1: \rho < \rho_0 \end{cases} \quad z^* < -z_{\alpha} \Rightarrow \text{رد می شود } H_0$$

$$z^* = \frac{\frac{1}{2} \ln \frac{1+r}{1-r} - \frac{1}{2} \ln \frac{1+\rho_0}{1-\rho_0}}{\frac{1}{\sqrt{n-3}}}$$

$$\begin{cases} H_0: \rho \leq \rho_0 \\ H_1: \rho > \rho_0 \end{cases} \quad z^* > z_{\alpha} \Rightarrow \text{رد می شود } H_0$$

مثال: در صورتی که همبستگی هوش و عملکرد تحصیلی ۰/۹ باشد فرض صفر بودن همبستگی در جامعه به کدام صورت نوشته می شود؟

$$H_0: \rho_{xy} \leq 0 \quad (۴) \quad H_0: \rho_{xy} \geq 0 \quad (۳) \quad H_0: \rho_{xy} = 0 \quad (۲) \quad H_0: \rho_{xy} \neq 0 \quad (۱)$$

گزینه ۲ صحیح است چون علامت مساوی در H_0 قرار میگیرد و گزینه ۳ جهت حداقل ۰/۹ است

مثال: در صورتی که در یک نمونه ۱۸ تایی ضریب همبستگی دو متغیر X و Y برابر 0.8 باشد آماره آزمون را بدست آورید و در سطح معنی داری 5% آزمون کنید معنی دار بودن ضریب همبستگی را.

$$r = 0.8, N = 18 \Rightarrow df = n - 2 = 18 - 2 = 16$$

$$t = \frac{r_{xy} \sqrt{N-2}}{\sqrt{1-r^2}} = \frac{0.8 \times \sqrt{18-2}}{\sqrt{1-0.8^2}} = \frac{0.8 \times \sqrt{16}}{\sqrt{0.36}} = \frac{3.2}{0.6} = 5.33$$

چون $t = 5.33 > t_{0.025, 18} = 2.101$ فرض صفر رد می شود ضریب همبستگی معنی دار است با احتمال 95% .

مثال: اگر به ازای سه نمره X_1, X_2, X_3 در متغیر X و Y_1, Y_2, Y_3 در متغیر Y روابط زیر برقرار باشد همبستگی چند خواهد بود؟

$$(x_1 - \bar{x})(y_1 - \bar{y}) = +1$$

$$(x_2 - \bar{x})(y_2 - \bar{y}) = 0$$

$$(x_3 - \bar{x})(y_3 - \bar{y}) = +1$$

$$+1(4)$$

$$0(3)$$

$$+0.5(2)$$

$$-1(1)$$

گزینه ۳ صحیح است. در واقع جمع اعداد $S_{xy} = 0$ برابر صفر است بنابراین ضریب همبستگی صفر است

نکته: آزمون ضریب همبستگی به شدت تابع حجم نمونه است. تا جایی که امکان دارد حجم نمونه را افزایش می دهیم تا آزمون با صحت بالاتری باشد

$$r = \frac{\sum Z_x \cdot Z_y}{N} \quad \text{نکته ۲: اگر مقادیر } X \text{ و } Y \text{ استاندارد شده باشد}$$

$$t = \frac{r\sqrt{n-2}}{\sqrt{1-r^2}} \xrightarrow{t^2=F} F = \left(\frac{r\sqrt{n-2}}{\sqrt{1-r^2}} \right)^2 \Rightarrow F = \frac{r^2 \cdot (n-2)}{1-r^2} \quad \text{نکته ۳: چون } t_{df}^2 = F_{df, df}$$

$$COV(ax, by) = a.b.CO V(x, y) \quad \text{نکته ۴:}$$

$$Z_{y \cdot x} = Z_x \cdot r_{xy} \quad \text{نکته ۵: پیش بینی نمره استاندارد}$$

نکته ۶: همیشه علامت S_{xy}, r, COV, b یکی است

مثال: متغیر X دارای توزیع نرمال $\bar{x} = 36, S_x = 4$ و متغیر Y دارای توزیع نرمال با $\bar{y} = 70, S_y = 7$ است. اگر همبستگی دو متغیر $r_{xy} = -1$ باشد. کواریانس بین دو متغیر کدام است؟

$$-36(4)$$

$$-28(3)$$

$$-10(2)$$

$$-9(1)$$

گزینه ۳ درست است

$$S_x = 4, S_y = 7, r_{xy} = -1$$

$$r_{xy} = \frac{COV(X, Y)}{S_x \cdot S_y} \Rightarrow -1 = \frac{COV(X, Y)}{4 \times 7} \Rightarrow COV(X, Y) = -1 \times 4 \times 7 = -28$$

ارشد مدیریت ۱۳۹۹

۵۷- اگر $\bar{x} = 4$ ، $\bar{y} = 5$ ، $s_x^2 = 16$ ، $s_y^2 = 25$ و $COV(x, y) = -15$ باشد، چند درصد تغییرات y به وسیله x

بیان می‌شود؟

(۱) ۷۵٪

(۲) ۴۷٪

(۳) ۵۶٪

(۴) ۲۵٪

$$s_x^2 = 16 \Rightarrow s_x = 4, s_y^2 = 25 \Rightarrow s_y = 5$$

$$r_{xy} = \frac{COV(X, Y)}{S_x \cdot S_y} \Rightarrow r_{xy} = \frac{-15}{4 \times 5} \Rightarrow r_{xy} = -0.75$$

$$r_{xy}^2 = (-0.75)^2 \times 100 = 56.25\%$$

گزینه ۳ صحیح است.

ارشد مدیریت ۱۴۰۰

۶۴- اگر $\bar{x} = 16$ ، $\bar{y} = 16$ ، $\sigma_x^2 = 9$ ، $\sigma_y^2 = 16$ و $Cov(x, y) = -9$ باشد، چند درصد تغییرات y به وسیله x بیان

می‌شود؟

(۱) ۷۵٪

(۲) ۶۳٪

(۳) ۵۰٪

(۴) ۵۶٪

$$s_x^2 = 9 \Rightarrow s_x = 3, s_y^2 = 16 \Rightarrow s_y = 4, COV(X, Y) = -9$$

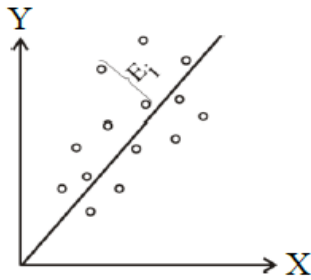
$$r_{xy} = \frac{COV(X, Y)}{S_x \cdot S_y} \Rightarrow r_{xy} = \frac{-9}{3 \times 4} = -0.75$$

$$r_{xy}^2 = (-0.75)^2 \times 100 = 56.25\%$$

رگرسیون خطی:

رگرسیون به معنی بازگشت برای تعیین رابطه بین یک چند متغیر مستقل (پیش بین) با یک متغیر وابسته (ملاک) است به عبارتی رگرسیون مطالعه چگونگی بازگشت نمره های y به x می باشد. در رگرسیون خطی ساده رابطه بین مستقل و یک متغیر وابسته را بررسی می کنیم در رگرسیون دو متغیر باید کمی (مقیاس فاصله ای یا نسبی) باشند و کاربرد رگرسیون خطی پیش بینی متغیر ملاک بر اساس متغیر پیش بین است. پارامترهای مدل به روش کمترین مربعات (L.S) بدست می آید رگرسیون خطی به صورت

$$y_i = a + bx_i + e_i \text{ است که } e_i \sim N(0, \sigma_e^2) \text{ هستند.}$$



$$e_i = y_i - \hat{y}_i$$

$$Q = \sum E_i^2 = \sum (y_i - a - b_i x_i)^2$$

$$\begin{cases} \frac{\partial Q}{\partial a} = -2 \sum (y_i - a - b_i x_i) = 0 \\ \frac{\partial Q}{\partial b} = -2 \sum (y_i - a - b_i x_i) x_i = 0 \end{cases} \Rightarrow \begin{cases} na + b \sum_{i=1}^n x_i = \sum_{i=1}^n y_i \\ a \sum_{i=1}^n x_i + b \sum_{i=1}^n x_i^2 = \sum_{i=1}^n x_i y_i \end{cases} \Rightarrow \begin{cases} \sum_{i=1}^n x_i y_i - n \bar{x} \bar{y} \\ \sum_{i=1}^n x_i^2 - n \bar{x}^2 \end{cases} \Rightarrow b = \frac{\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}}{\sum_{i=1}^n x_i^2 - n \bar{x}^2} \Rightarrow b = \frac{n \sum xy - (\sum x)(\sum y)}{n \sum x^2 - (\sum x)^2}$$

توجه ۱: در صورتیکه $\sum (y_i - \hat{y}_i)^2 = 0$ باشد بیشترین رابطه بین متغیر مستقل و وابسته وجود دارد

توجه ۲: تکرار آزمایش بهترین روش تشخیص برای ثبات تاثیر متغیر مستقل بر متغیر وابسته است.

مراحل نوشتن معادله خط رگرسیون :

$$1) b_{xy} = \frac{S_{xy}}{S_{xx}} = \frac{n \sum xy - (\sum x)(\sum y)}{n \sum x^2 - (\sum x)^2} \quad 2) a = \bar{y} - b \bar{x} \quad 3) y = a + b x$$

$$S_{xy} = \sum (x_i - \bar{x})(y_i - \bar{y}) = \sum x_i y_i - n \bar{x} \bar{y} = \sum x_i y_i - \frac{\sum x_i \sum y_i}{n}$$

$$S_{xx} = \sum (x_i - \bar{x})^2 = \sum x_i^2 - n \bar{x}^2 = \sum x_i^2 - \frac{(\sum x_i)^2}{n}$$

یادآوری

$$S_{yy} = \sum (y_i - \bar{y})^2 = \sum y_i^2 - n \bar{y}^2 = \sum y_i^2 - \frac{(\sum y_i)^2}{n}$$

محاسبه مجموع مربعات خطا و واریانس خطا به صورت زیر محاسبه می شود.

$$\sum e_i^2 = SS_E = S_{yy} - b^2 S_{xx} = S_{yy} - \frac{S_{xy}^2}{S_{xx}}$$

$$\sigma_{\text{خطا}}^2 = \frac{SS_E}{n-2}$$

نکته ۲: برای پیش بینی کردن میانگین پاسخ به جای متغیر مستقل مقدار داده شده را قرار دهیم

$$E(y|x=x^*) = \hat{y}^* = \hat{a} + \hat{b} x^*$$

ارشد مدیریت ۱۳۹۸

۵۶- قیمت (x) و تقاضای کالایی (y) در ۵ بازار مختلف جمع آوری شده است. شیب خط رگرسیون کدام است؟

x	۱۲	۱۰	۱۳	۱۶	۹
y	۲۰	۱۵	۱۱	۷	۲۲

-۱/۴ (۴)

-۱/۹ (۳)

-۱/۶ (۲)

-۱/۷ (۱)

x	y	$x - \bar{x}$	$y - \bar{y}$	$(x - \bar{x})^2$	$(x - \bar{x}) \times (y - \bar{y})$
۱۲	۲۰	$12 - 12 = 0$	$20 - 15 = 5$	$0^2 = 0$	$0 \times 5 = 0$
۱۰	۱۵	$10 - 12 = -2$	$15 - 15 = 0$	$(-2)^2 = 4$	$-2 \times 0 = 0$
۱۳	۱۱	$13 - 12 = 1$	$11 - 15 = -4$	$1^2 = 1$	$1 \times (-4) = -4$
۱۶	۷	$16 - 12 = 4$	$7 - 15 = -8$	$4^2 = 16$	$4 \times -8 = -32$
۹	۲۲	$9 - 12 = -3$	$22 - 15 = 7$	$(-3)^2 = 9$	$-3 \times 7 = -21$
$\sum x = 60$	$\sum y = 75$	.	.	$S_{xx} = 40$	$S_{xy} = -57$

$$\bar{x} = \frac{\sum x}{n} = \frac{60}{4} = 12, \quad \bar{y} = \frac{\sum y}{n} = \frac{75}{5} = 15$$

$$b = \frac{S_{xy}}{S_{xx}} = \frac{-57}{40} = -1.425$$

گزینه ۳ صحیح است

ارشد مدیریت ۱۴۰۰

۶۵- اگر معادله رگرسیونی دو متغیر x و y به صورت $y = 3x + 56$ حاصل شده باشد و به ازای مقدار $x = 10$ مقدارواقعی $y = 90$ باشد، کدام مورد درباره همبستگی دو متغیر درست است؟۱) $0 < \rho < 1$ ۲) $-1 < \rho < 0$ ۳) $\rho = 1$ ۴) $\rho = -1$

$$y = 3x + 56$$

$$E(y|x=x^*) = y^* = \hat{a} + \hat{b}x^* \Rightarrow y^* = 3 \times 10 + 56 = 86$$

گزینه ۱ صحیح است چون مقدار برآورد شده با مقدار واقعی برابر نیست و همچنین شیب خط +۳ است

بنابراین همبستگی عددی بین صفر و یک است.

ارشد مدیریت ۱۴۰۱

۶۵- براساس ۵ نقطه (x, y) داریم: $\sum x = 25$ و $\sum y = 45$ ، $\sum xy = 195$ ، $\sum x^2 = 145$ ، $\sum y^2 = 451$ شیب معادله خط رگرسیون چقدر است؟

(۱) $-1/25$ (۲) $-1/5$ (۳) $2/75$ (۴) $1/25$

$$\bar{x} = \frac{\sum x}{n} = \frac{25}{5} = 5, \bar{y} = \frac{\sum y}{n} = \frac{45}{5} = 9$$

$$s_{xy} = \sum xy - n\bar{x}\bar{y} = 195 - (5 \times 5 \times 9) = -30$$

$$s_{xx} = \sum x^2 - n\bar{x}^2 = 145 - 5 \times 5^2 = 20$$

$$b = \frac{s_{xy}}{s_{xx}} = \frac{-30}{20} = -1/5$$

گزینه ۲ صحیح است

ارشد مدیریت ۱۴۰۲

۶۰- به منظور بررسی وجود رابطه خطی بین دو متغیر، ۱۶ نمونه تصادفی انتخاب شده است که براساس آنها، میانگین مجذورات تغییرات ناشی از خطا و آماره آزمون به ترتیب ۶ و ۲۱ به دست آمده است. چند درصد از تغییرات متغیر وابسته، به وسیله متغیر مستقل تبیین می شود؟

(۱) ۳۶

(۲) ۴۵

(۳) ۶۰

(۴) ۶۷

$$F = 21 \quad SS_E = 6 \quad n = 16$$

$$F = \frac{R^2 \cdot (n-2)}{1-R^2} \Rightarrow 21 = \frac{R^2 \cdot (16-2)}{1-R^2} \Rightarrow \frac{21}{14} = \frac{R^2}{1-R^2}$$

$$1/5 = \frac{R^2}{1-R^2} \Rightarrow 1/5 - 1/5 R^2 = R^2 \Rightarrow 1/5 = 2/5 R^2 \Rightarrow R^2 = \frac{1/5}{2/5} = 0/6 = 60\%$$

گزینه ۳ صحیح است

پیش فرض های رگرسیون:

۱- بین متغیرهای پیش بین و ملاک رابطه خطی وجود داشته باشد. که با نمودار پراکنش یا ضریب همبستگی پیرسون

- ۲- متغیرهای مستقل و وابسته باید دارای مقیاس فاصله ای باشند
- ۳- واریانس خطا نرمال یا ثابت باشد با بررسی نمودار صورت می گیرد
- ۴- توزیع فراوانی باقیمانده ها نرمال باشد با آزمون کولموگروف اسمیرنوف
- ۵- باقیمانده ها از هم مستقل باشند تست دوربین واتسون (عدد نزدیک به ۲)
- ۶- بین متغیرهای مستقل رابطه همخطی وجود نداشته باشد با مقدار تلرانس یا VIF بررسی می شود
- مثال: داده های زیر مربوط به نتایج یک بررسی در زمینه انجماد شتاب داده شده شمشها است که Y عمق انجماد و X معرف زمان سرد شدن تعیین کنید (الف ضریب همبستگی ب) معادله خط رگرسیون ج) امید ریاضی عمق انجماد پس از ۵۰ دقیقه د) واریانس خطا

عمق انجماد (Y)	۵	۶	۷/۵	۸/۵	۱۰	۱۲
Y زمان سرد شدن	۸	۱۴	۱۸/۵	۲۳/۵	۳۰	۳۸

جواب

x	Y	x^2	y^2	xy
۸	۵	۶۴	۲۵	۴۰
۱۴	۶	۱۹۶	۳۶	۸۴
۱۸/۵	۷/۵	۳۴۲/۲۵	۵۶/۲۵	۱۳۸/۷۵
۲۳/۵	۸/۵	۵۵۲/۲۵	۷۲/۲۵	۱۹۹/۷۵
۳۰	۱۰	۹۰۰	۱۰۰	۳۰۰
۳۸	۱۲	۱۴۴۴	۱۴۴	۴۵۶
۱۳۲	۴۹	۳۴۹۸/۵	۴۳۳/۵	۱۲۱۸/۵

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n} = \frac{132}{6} = 22$$

$$\bar{y} = \frac{\sum_{i=1}^n y_i}{n} = \frac{49}{6} = 8.17$$

$$s_{xy} = \sum_{i=1}^n x_i y_i - n \bar{x} \bar{y} = 1218/5 - 6 \times 22 \times 8.17 = 140.06$$

$$s_{xx} = \sum_{i=1}^n x_i^2 - n \bar{x}^2 = 3498/5 - 6 \times 22^2 = 594/5$$

$$s_{yy} = \sum_{i=1}^n y_i^2 - n \bar{y}^2 = 433/5 - 6 \times 8.17^2 = 33$$

$$\text{الف) } r = \frac{s_{xy}}{\sqrt{s_{xx}s_{yy}}} = \frac{140/0.6}{\sqrt{594/5 \times 33}} = \frac{140/0.6}{140/0.66} = 0.999$$

$$\hat{b} = \frac{s_{xy}}{s_{xx}} = \frac{140/0.6}{594/5} = 0.2356$$

$$\text{ب) } \hat{a} = \bar{y} - b\bar{x} = 8/17 - 0.2356 \times 22 = 2/98$$

$$\hat{y} = \hat{a} + \hat{b}x = 2/98 + 0.2356x$$

$$\text{ج) } E(y|x=x^*) = \hat{a} + \hat{b}x^* = 2/98 + 0.2356 \times 50 = 14/76$$

$$ss_E = s_{yy} - b^2 s_{xx} = 33 - 0.2356^2 \times 594/5 = 0.0009$$

$$\text{د) } s_{y|x}^2 = \sigma_e^2 = \frac{ss_E}{n-2} = \frac{0.0009}{6-2} = 0.000225$$

بر آورد فاصله‌ای برای شیب خط رگرسیون

$$\sigma_b^2 = \frac{\sigma_e^2}{s_{xx}} = \frac{\sigma_e^2}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

$$B: \hat{b} \pm t_{\alpha/2, df} \sigma_b \Rightarrow \beta: \hat{b} \pm t_{\alpha/2, df} \frac{\sigma_e}{\sqrt{s_{xx}}}$$

بر آورد فاصله‌ای برای عرض از مبدا

$$\sigma_a^2 = \sigma_e^2 \left(\frac{1}{n} + \frac{\bar{x}^2}{s_{xx}} \right) \Rightarrow \sigma_a^2 = \sigma_e^2 \left(\frac{1}{n} + \frac{\bar{x}^2}{\sum (x - \bar{x})^2} \right)$$

$$a: \hat{a} \pm t_{\alpha/2, df} \sqrt{\sigma_e^2 \left(\frac{1}{n} + \frac{\bar{x}^2}{s_{xx}} \right)} \Rightarrow a: \hat{a} \pm t_{\alpha/2, df} \sigma_e \sqrt{\frac{1}{n} + \frac{\bar{x}^2}{s_{xx}}}$$

بر آورد نقطه‌ای و فاصله‌ای برای میانگین به ازای مقدار مشخص x

$$E(y|x=x^*) = y^* = \hat{a} + \hat{b}x^*$$

$$y^*: \hat{a} + \hat{b}x^* \pm t_{\alpha/2, df} \sigma_e \sqrt{\frac{1}{n} + \frac{(x^* - \bar{x})^2}{s_{xx}}}$$

فاصله اطمینان پیش‌بین برای مشاهده‌آتی از متغیر وابسته

$$\hat{y}' : \hat{a} + \hat{b}x^* \pm t_{\frac{\alpha}{2}, df} \sigma_e \sqrt{1 + \frac{1}{n} + \frac{(x' - \bar{x})^2}{s_{xx}}}$$

بر آورد نقطه‌ای و فاصله‌ای برای متغیر مستقل x به ازای یک مشاهده از متغیر وابسته y

$$x' = \frac{y' - a}{b}$$

$$x' = \frac{y' - a}{b} \pm t_{\frac{\alpha}{2}, df} \cdot \frac{\sigma_e}{b} \sqrt{1 + \frac{1}{n} + \frac{(x' - \bar{x})^2}{s_{xx}}}$$

آزمون فرض برای شیب خط رگرسیون:

$$\begin{cases} H_0: \beta = \beta_0 \\ H_1: \beta \neq \beta_0 \end{cases} \quad |t| > t_{\frac{\alpha}{2}, df} \Rightarrow \text{رد می‌شود } H_0$$

$$\begin{cases} H_0: \beta \geq \beta_0 \\ H_1: \beta < \beta_0 \end{cases} \quad t < -t_{\alpha, df} \Rightarrow \text{رد می‌شود } H_0$$

$$df = n - 2$$

$$t = \frac{b - B_0}{\frac{\sigma_e}{\sqrt{s_{xx}}}}$$

$$\begin{cases} H_0: \beta \leq \beta_0 \\ H_1: \beta > \beta_0 \end{cases} \quad t > t_{\alpha, df} \Rightarrow \text{رد می‌شود } H_0$$

آزمون فرض برای عرض از مبدا

$$\begin{cases} H_0: A = A_0 \\ H_1: A \neq A_0 \end{cases} \quad |t| > t_{\frac{\alpha}{2}, df} \Rightarrow \text{رد می‌شود } H_0$$

$$\begin{cases} H_0: A \geq A_0 \\ H_1: A < A_0 \end{cases} \quad t < -t_{\alpha, df} \Rightarrow \text{رد می‌شود } H_0$$

$$df = n - 2$$

$$t = \frac{a - A_0}{\sigma_e \sqrt{\frac{1}{n} + \frac{\bar{x}^2}{s_{xx}}}}$$

$$\begin{cases} H_0: A \leq A_0 \\ H_1: A > A_0 \end{cases} \quad t > t_{\alpha, df} \Rightarrow \text{رد می‌شود } H_0$$

مثال: به منظور به دست آوردن رسوب از یک ماده خاص γ در یک محلول مفروض از معرف شیمیایی، x استفاده می شود خلاصه اطلاعات به شرح زیر است.

$$\sum_{i=1}^{15} x_i = 9, \quad \sum_{i=1}^{15} (x_i - \bar{x})^2 = 21$$

$$\sum_{i=1}^{15} y_i = 135, \quad \sum_{i=1}^{15} (y_i - \bar{y})^2 = 119$$

$$\sum_{i=1}^{15} (x_i - \bar{x})(y_i - \bar{y}) = 42$$

تعیین کنید.

الف) ضریب همبستگی

ب) معادله خط رگرسیون

پ) امید ریاضی مقدار رسوب به ازای $x = \gamma$

ت) فاصله اطمینان ۹۵٪ برای امید ریاضی مقدار رسوب نظیر $x = \gamma$

ج) فاصله پیش بینی ۹۵٪ مقدار رسوب نظیر $x = \gamma$

د) آزمون فرض معنادار بودن ضریب همبستگی در سطح خطای ۵٪

ذ) آزمون فرض معنادار شیب خط در سطح خطای ۵٪

ه) آزمون فرض معنادار بودن عرض از مبدا در سطح خطای (۵٪/۱۶/۲۵,۱۳) (t)

الف) $r = \frac{s_{xy}}{\sqrt{s_{xx}s_{yy}}} = \frac{42}{\sqrt{21 \times 119}} = 0.84$

$$\hat{b} = \frac{s_{xy}}{s_{xx}} = \frac{42}{21} = 2$$

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n} = \frac{9}{15} = 0.6$$

$$\bar{y} = \frac{\sum_{i=1}^n y_i}{n} = \frac{135}{15} = 9$$

ب) $\hat{a} = \bar{y} - \hat{b}\bar{x} = 9 - 2 \times 0.6 = -3$

$$\hat{y} = \hat{a} + \hat{b}x = -3 + 2x$$

پ) $\hat{y}^* = \hat{a} + \hat{b}x^* = -3 + 2 \times 7 = 11$

$$SS_E = s_{yy} - b^2 s_{xx} = 119 - 2^2 \times 21 = 35$$

$$s_{y|x}^2 = \sigma_e^2 = \frac{SS_E}{n-2} = \frac{35}{15-2} = 2/69 \Rightarrow \sigma_e = 1/64$$

$$\text{ت) } \hat{y}^* = \hat{a} + \hat{b}x^* \pm t_{\frac{\alpha}{2}, df} \sigma_e \sqrt{\frac{1}{n} + \frac{(x^* - \bar{x})^2}{s_{xx}}}$$

$$y^* = (-3 + 2 \times 7) \pm 2/16 \times 1/64 \sqrt{\frac{1}{15} + \frac{(7-6)^2}{21}} = 11 \pm 1/197$$

$$9/80.3 < y^* < 12/197$$

$$y^* = a + bx^* \pm t_{\frac{\alpha}{2}, df} \sigma_e \sqrt{\frac{1}{n} + \frac{(x^* - \bar{x})^2}{s_{xx}}}$$

$$\text{ج) } y^* = (-3 + 2 \times 7) \pm 2/16 \times 1/64 \sqrt{\frac{1}{15} + \frac{(7-6)^2}{21}} = 11 \pm 3/74$$

$$7/26 < y^* < 14/74$$

$$\text{د) } \begin{cases} H_0: \rho = 0 \\ H_1: \rho \neq 0 \end{cases} \quad t = \frac{(r - \rho_0) \sqrt{n-2}}{\sqrt{1-r^2}} = \frac{0/84 \sqrt{15-2}}{\sqrt{1-0/84^2}} = 5/5819$$

چون $t = 5/5819 > t_{0.25, 13} = 2/16$ است بنابراین H_0 رد می شود و ضریب همبستگی در سطح اطمینان ۹۵٪ معنادار است.

$$\text{ذ) } \begin{cases} H_0: \beta = 0 \\ H_1: \beta \neq 0 \end{cases} \quad t = \frac{b - \beta_0}{\frac{\sigma_e}{\sqrt{s_{xx}}}} = \frac{2-0}{1/64 \sqrt{21}} = 5/5885$$

چون $t = 5/5885 > t_{0.25, 13} = 2/16$ است بنابراین H_0 رد می شود و شیب خط در سطح اطمینان ۹۵٪ معنادار است.

$$\text{ه) } \begin{cases} H_0: A = 0 \\ H_1: A \neq 0 \end{cases} \quad t = \frac{a - A_0}{\sigma_e \sqrt{\frac{1}{n} + \frac{x^2}{s_{xx}}}} = \frac{-3-0}{1/64 \sqrt{\frac{1}{15} + \frac{62}{21}}} = -1/37$$

چون $|t| = 1/37 > t_{0.25, 13} = 2/16$ است بنابراین H_0 تایید می شود و عرض از مبدا در سطح اطمینان ۹۵٪ معنی دار نیست و برابر صفر است.

نکته: برای تعیین معادله خط رگرسیون ابتدا دیاگرام پراکنش را رسم می‌کنیم و مدل را تشخیص و با تغییر متغیر مناسب آن به مدل خطی تبدیل می‌کنیم به عنوان مثال اگر تشخیص دهیم رابطه بین X و Y نمایی است با لگاریتم گرفتن از طرفین به مدل خطی تبدیل می‌شود و مدل بدست می‌آوریم.

$$y = ab^x \Rightarrow \ln y = \ln a + x \ln b \quad y^* = a^* + b^* x$$

$$\Rightarrow a^* = e^{a^*}, \quad b^* = e^{b^*} \Rightarrow y = (e^{a^*}) \cdot (e^{b^*})^x$$

نکته: همیشه یک تابع خطی و درجه یک بر حسب X ، برای Y وجود دارد یعنی:

$$(x, y) \sim (\mu_x, \mu_y, \sigma^2_x, \sigma^2_y, \rho)$$

$$\Rightarrow y - \mu_y = \rho \frac{\sigma_y}{\sigma_x} (x - \mu_x)$$

نکته: در خروجی SPSS در ستون **Sigt** در آزمون فرض‌ها می‌توان بدون داشتن جدول **t** با مقایسه عدد **Sigt** و α در مورد B_j تصمیم‌گیری کرد یعنی اگر این مقدار احتمال کوچکتر یا مساوی α باشد فرض صفر را رد می‌کنیم و $B_j \neq 0$ است.

مثال: داده‌های مقابل مفروض است اگر $E(e_i) = 0, \text{Var}(e_i) = \sigma^2$ باشد رگرسیون خطی Y بر حسب X کدام است؟

Y	X
۱	۲
۳	۰
۵	۴

$$y = x + 1 \quad (1) \quad y = \frac{1}{4}(x+1) \quad (2) \quad y = \frac{1}{4}(x-1) \quad (3) \quad y = -\frac{1}{4}(x-1) \quad (4)$$

جواب: گزینه ۲ صحیح است

x	y	xy	x^2	y^2
۱	۲	۲	۱	۴
۳	۰	۰	۹	۰
۵	۴	۲۰	۲۵	۱۶
$\sum x = 9$	$\sum y = 6$	$\sum xy = 22$	$\sum x^2 = 35$	$\sum y^2 = 20$

$$b_{xy} = \frac{n \sum xy - (\sum x \cdot \sum y)}{n \sum x^2 - (\sum x)^2} = \frac{(3 \times 22) - (9 \times 6)}{3 \times 35 - 9^2} = \frac{66 - 54}{24} = \frac{12}{24} = \frac{1}{2}$$

$$\bar{x} = \frac{\sum x}{n} = \frac{9}{3} = 3, \quad \bar{y} = \frac{\sum y}{n} = \frac{6}{3} = 2$$

$$2) a = \bar{y} - b\bar{x} = 2 - \left(\frac{1}{2} \times 3\right) = \frac{4}{2} - \frac{3}{2} = \frac{1}{2}$$

$$3) y = a + bx = \frac{1}{2} + \frac{1}{2}x = \frac{1}{2}(1+x)$$

مثال: اصل حداقل مجذورات (LS) تفاوت کدام دو نمره را کمینه می سازد؟

$$X, \varepsilon \quad (2) \quad \varepsilon, y \quad (1)$$

$$\hat{y}, \varepsilon \quad (4) \quad \hat{y}, \varepsilon \quad (3)$$

گزینه ۴ صحیح است چون $e_i = y_i - \hat{y}_i$ خطا باید حداقل شود

نکته ۷: خط رگرسیون از نقطه $\bar{x}$ و $\bar{y}$ می گذرد.

$$b = r \cdot \frac{S_y}{S_x} \quad \text{نکته ۸: رابطه شیب خط رگرسیون و ضریب همبستگی}$$

$$(y - \bar{y}) = r \cdot \frac{S_y}{S_x} (x - \bar{x}) \quad \text{با توجه به نکته ۱ و ۲ نتیجه می گیریم}$$

$$S_{xy} = S_y \sqrt{1 - r_{xy}^2} \quad \text{نکته ۹: رابطه واریانس خطای پیش بینی (استاندارد) و ضریب همبستگی}$$

نکته ۱۰: ضریب رگرسیون تفکیکی استاندارد شده که معیاری برای تعیین مقدار تاثیر متغیر پیش بین

$$\beta = \frac{S_x}{S_y} b \quad \text{روی متغیر ملاک است و عددی بین -۱ و +۱ است}$$

توجه ۱: در رگرسیون یک متغیره $df = n - 2$ و در حالت کلی $df = n - k - 1$ است که k تعداد متغیر مستقل می باشد

توجه ۲: در صورتی که بخواهیم معادله خط x بر حسب y را بنویسیم به ترتیب زیر عمل می کنیم.

$$1) \hat{b}^* = \frac{S_{xy}}{S_{yy}} \quad 2) \hat{a}^* = \bar{x} - \hat{b}^* \bar{y} \quad 3) \hat{x} = \hat{a}^* + \hat{b}^* y$$

$$r = \sqrt{b \cdot b^*} \quad \text{نکته ۱۱: ضریب همبستگی واسطه هندسی } b \text{ و } b^* \text{ است یعنی}$$

مثال: اگر $r_{xy} = 0.5$ باشد و تمام نمرات دو متغیر X و Y به صورت نمره های استاندارد باشند عرض از مبدا و شیب خط رگرسیون برابر چند است؟

$$1) -0.5 \quad 2) 0.5-1 \quad 3) 0.5-0 \quad 4) 1-0.5$$

گزینه ۳ صحیح است چون نرمال استاندارد است بنابر این میانگین صفر و واریانس یک است

$$s_x = s_y = 1 \Rightarrow b = r \frac{s_y}{s_x} \Rightarrow b = 0.5 \times \frac{1}{1} = 0.5 \Rightarrow a = \bar{y} - b \bar{x} = 0 - (0.5 \times 0) = 0$$

نکته: $E(ax + by) = aE(x) + bE(y)$ و $\text{var}(ax + by) = \text{var}(x) + \text{var}(y) + 2a.b.\text{cov}(x, y)$

مثال: اگر $\sigma_x^2 = 2$, $\sigma_y^2 = 3$, $\sigma_{xy} = 1$ باشد. حاصل σ_{x+y}^2 چند است؟

$$\begin{array}{cccc} 3(1) & 5(2) & 6(3) & 7(4) \end{array}$$

گزینه ۴ صحیح است. $\text{var}(ax + by) = \text{var}(x) + \text{var}(y) + 2a.b.\text{cov}(x, y) = 2 + 3 + 2 \times 1 = 7$

متغیر تعدیل کننده: متغیری است که به صورت مستقیم بر جهت یا میزان رابطه متغیرهای مستقل و وابسته می تواند موثر باشد. اثرات این متغیر قابل مشاهده و اندازه گیری است. مانند رابطه هوش با افزایش میزان یادگیری دانش آموزان یا متغیر جنسیت در بررسی رابطه ی روش تدریس و یادگیری دانشجویان یک متغیر تعدیل کننده است

خطای پیش بینی: اختلاف مقدار واقعی (y) و مقدار پیش بینی را خطای پیش بینی گویند و با y' نشان می دهیم

ضریب رگرسیون تفکیکی استاندارد شده: اگر به جای نمره های خام از داده های استاندارد شده z استفاده کنیم از ضریب رگرسیون تفکیکی استاندارد شده استفاده می کنیم و با حرف β نشان می دهیم.

و از رابطه $\beta = b \frac{s_y}{s_x}$ بدست می آید. که در آن b شیب خط رگرسیون و s_x انحراف معیار متغیر پیش

بین (x) و s_y انحراف معیار متغیر ملاک (y) است

$$s_y = \sqrt{\frac{\sum y^2 - \frac{(\sum y)^2}{n}}{n}}, \quad s_x = \sqrt{\frac{\sum x^2 - \frac{(\sum x)^2}{n}}{n}}$$

مثال: در صورتی که شیب خط رگرسیون ۴ و انحراف معیار متغیر پیش بین ۱۰ و انحراف معیار متغیر ملاک ۵ است. ضریب رگرسیون تفکیکی استاندارد شده کدام است؟

$$\begin{array}{cccc} 1(الف) & 2(ب) & 3(ج) & 4(د) \end{array}$$

$$\beta = b \frac{s_y}{s_x} = 4 \times \frac{5}{10} = 2$$

جواب گزینه ۲ صحیح است.

آزمون خطی بودن (Lack of fit): در صورتی که به ازای یک x چند y وجود داشته باشد از آزمون خطی یا Lack of fit استفاده می شود.

$$\begin{cases} H_0: y = a + bx \\ H_1: y \neq a + bx \end{cases}$$

که در این حالت متوسط y مقادیر X_i را با $\bar{X}$ نشان می دهیم.

$$\begin{array}{ccc}
 x_1 \rightarrow y_{11} & x_2 \rightarrow y_{21} & x_n \rightarrow y_{n1} \\
 & y_{22} & y_{n2} \\
 & \vdots & \vdots \\
 & y_{2k_2} & y_{nk_n} \\
 & \bar{y}_2 & \bar{y}_n
 \end{array}$$

$$\begin{aligned}
 SS_E &= \sum_{i=1}^n \sum_{j=1}^{k_i} \left(y_{ij} - \hat{y}_i \right)^2 \\
 SS_{PE} &= \sum_{i=1}^n \sum_{j=1}^{k_i} \left(y_{ij} - \hat{y}_i \right)^2 \\
 SS_{Lof} &= \sum_{i=1}^n k_i \left(\bar{y}_i - \hat{y}_i \right)^2, \quad SS_{Lof} = SS_E - SS_{PE}
 \end{aligned}$$

که جدول آنالیز واریانس آن به صورت زیر است.

منبع تغییرات	SS	df	$MS_R = \frac{SS}{df}$	F
رگرسیون	SS_R	۱	$MS_R = SS_R$	$F = \frac{MS_{Lof}}{MS_{PE}}$
Lof	SS_{Lof}	$n - 2$	$MS_{Lof} = \frac{SS_{Lof}}{n - 2}$	
PE	SS_{PE}	$N - n$	$MS_{PE} = \frac{SS_{PE}}{N - n}$	
کل	S_{yy}	$N - 1$		

اگر $F > F_{\alpha, n-2, N-n}$ باشد H_0 رد می شود و آزمون خطی بودن مردود است و در حالت خاص اگر

$k_1 = k_2 = \dots = k_n$ باشد در این صورت رابطه F را می توان به صورت زیر نوشت

$$F = \frac{k \sum_{i=1}^n \frac{\left(\bar{y}_i - \hat{y}_i \right)^2}{n - 2}}{\sum_{i=1}^n \sum_{j=1}^k \frac{\left(\bar{y}_j - \hat{y}_i \right)^2}{n(k-1)}}$$

که در این صورت اگر $F > F_{\alpha, n-2, n(k-1)}$ باشد فرض H_0 رد می شود.

شرط کافی برای آن که دو متغیر تصادفی x و y دارای توزیع نرمال باشند آن است که دو متغیر رابطه خطی داشته باشند.

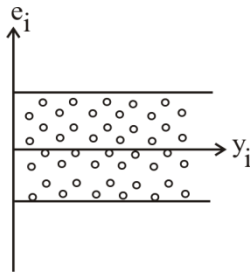
$$f_{xy}(u, v) = \frac{1}{\sqrt{2\pi}\sigma_x\sigma_y\sqrt{1-\rho^2}} \exp \left[-\frac{\frac{(u-\mu_x)^2}{\sigma_x^2} - 2\rho\frac{(u-\mu_x)(v-\mu_y)}{\sigma_x\sigma_y} + \frac{(v-\mu_y)^2}{\sigma_y^2}}{2(1-\rho^2)} \right]$$

که در این صورت $\sigma^2 = \sigma_y^2(1-\rho^2)$ ، $E(y|x) = \mu_x + \rho\frac{\sigma_y}{\sigma_x}(x-\mu_x)$

تشخیص در تحلیل رگرسیون از روی نمودارها:

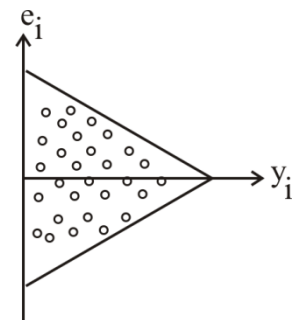
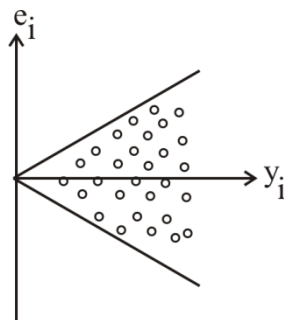
در صورتی که $\left(\hat{y}_i, e_i\right)$ را رسم کنیم حالات زیر امکان دارد رخ دهد.

(۱) داده ها بین دو خط موازی قرار گیرد در این صورت

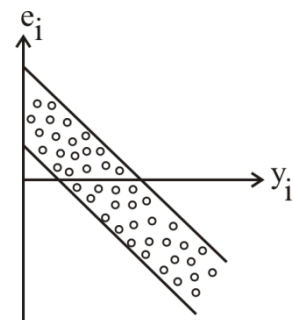
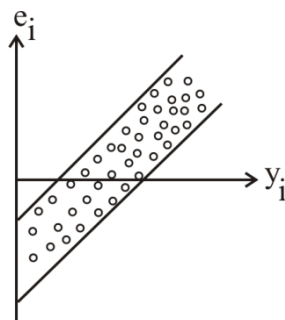


$\hat{y}_i \Rightarrow e_i = 0 \Rightarrow$ فرضیه های مدل برقرار است

(۲) داده ها به صورت دوکی باشد در این صورت واریانس خطا ثابت نیست



(۳) داده بین دو خط موازی شیب دار باشد در این صورت مدل باید بازسازی شود



تحلیل پس مانده‌ها:

همانگونه که می‌دانیم $e_i \sim N(0, \sigma_e^2)$ می‌باشد بنابراین $\frac{e_i}{\sigma}$ دارای توزیع نرمال با میانگین صفر و واریانس یک است در صورتی که $d_i = \frac{e_i}{\sqrt{MSE}}$ باشد اگر $P\left(-z_{\frac{\alpha}{2}} < d_i < z_{\frac{\alpha}{2}}\right)$ باشد نرمال بودن را می‌پذیریم یعنی اگر $P(-1/96 < d_i < 1/96) = 0/95$ باشد با احتمال ۹۵٪ نرمال بودن را می‌پذیریم

رگرسیون چند متغیره:

مدل رگرسیونی که بیش از یک متغیر مستقل داشته باشد رگرسیون چند متغیره گویند که مدل آن به صورت $y = B_0 + B_1x_1 + B_2x_2 + B_3x_3 + \dots + B_kx_k + \varepsilon$ می‌باشد که برای برآورد پارامترها آن به فرم ماتریسی نمایش می‌دهیم.

$$e = y - y' \Rightarrow SS_{res} = \sum (y - y')^2 = \sum e^2 \Rightarrow SS_{res} = \sum y'^2 - \frac{(\sum y)^2}{n}$$

$$B = \begin{bmatrix} B_0 \\ B_1 \\ \vdots \\ B_k \end{bmatrix} \quad \varepsilon = \begin{bmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \vdots \\ \varepsilon_k \end{bmatrix} \quad y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix} \quad x = \begin{bmatrix} 1 & x_{11} & x_{12} & x_{1k} \\ 1 & x_{21} & x_{22} & x_{2k} \\ \vdots & \vdots & \vdots & \vdots \\ 1 & x_{n1} & x_{n2} & x_{nk} \end{bmatrix}_{n \times k}$$

k: تعداد متغیرهای مستقل و n: تعداد داده‌ها

لازم به توضیح است که در ماتریس x اگر مقدار ثابت B₀ در ماتریس B در نظر گرفته نشود عدد یک را در ستون اول نمی‌نویسیم که پارامترها به روش کمترین مربعات به صورت زیر محاسبه می‌شود.

$$\hat{b} = (x'x)^{-1}x'y$$

$$\text{cov}(b) = E\left([b - E(b)]*[b - E(b)]'\right) \Rightarrow \text{cov}(b) = \sigma^2 (x'x)^{-1}$$

$$SS_E = S_{yy} - SS_R = y'y - b'x'y$$

$$\sigma^2 = \text{MSE} = \frac{SS_E}{n-p}, p = k+1$$

ماتریس $(x'x)^{-1}$ یک ماتریس متقارن است اگر اعضای این ماتریس را به صورت c_{ij} نمایش دهیم آنگاه واریانس $\hat{B}_i$ از رابطه

$$(x'x)^{-1} = \begin{bmatrix} c_{..} & c_{.1} & c_{.2} & c_{.k} \\ c_{1.} & c_{11} & c_{12} & c_{1k} \\ \vdots & \vdots & \vdots & \vdots \\ c_{k.} & c_{k1} & c_{k2} & c_{kk} \end{bmatrix}$$

$$\hat{y} = xb = x(x'x)^{-1}xy = Hy$$

ماتریس $H = x(x'x)^{-1}x$ یک ماتریس خود توان است.

آزمون فرض ضرایب انفرادی رگرسیون

$$\begin{cases} H_0: \beta_j = \beta_j \\ H_1: \beta_j \neq \beta_j \end{cases} \quad |t| > t_{\frac{\alpha}{2}, df} \Rightarrow H_0 \text{ رد می شود}$$

$$\begin{cases} H_0: \beta_j \geq \beta_j \\ H_1: \beta_j < \beta_j \end{cases} \quad t < -t_{\alpha, df} \Rightarrow H_0 \text{ رد می شود} \quad \begin{aligned} df &= n - 2 \\ t &= \frac{b_j - \beta_j}{\sqrt{c_{jj} \times \hat{\sigma}^2}} \end{aligned}$$

$$\begin{cases} H_0: \beta_j \leq \beta_j \\ H_1: \beta_j > \beta_j \end{cases} \quad t > t_{\alpha, df} \Rightarrow H_0 \text{ رد می شود}$$

بر آورد فاصله ضرایب رگرسیون:

$$B_j: b_j \pm t_{\frac{\alpha}{2}, df} \sqrt{\hat{\sigma}^2 c_{jj}}$$

بر آورد میانگین پاسخ و بر آورد فاصله‌ای میانگین پاسخ:

$$\begin{aligned} \tilde{y}^* &= x^{*'}b \\ \Rightarrow E(y^*): \tilde{y}^* &\pm t_{\frac{\alpha}{2}, df} \sqrt{\hat{\sigma}^2 x^{*'}(x'x)^{-1}x^*} \end{aligned}$$

پیش بینی مشاهدات جدید و فاصله اطمینان پیش بینی جدید:

$$\begin{aligned} \tilde{y}^* &= x^{*'}b \\ \Rightarrow y^*: \tilde{y}^* &\pm t_{\frac{\alpha}{2}, df} \sqrt{\hat{\sigma}^2 (1 + x^{*'}(x'x)^{-1}x^*)} \end{aligned}$$

آزمون معنادار بودن رگرسیون: با استفاده از تحلیل واریانس این آزمون را انجام می‌دهیم

منبع تغییرپذیری SOV	مجموع مربعات SS	درجات آزادی df	میانگین مربعات $MS = \frac{SS}{df}$	آماره آزمون F
رگرسیون	SS_R	K	$MS_R = \frac{SS_R}{k}$	$F = \frac{MS_R}{MS_E}$
خطا	SS_E	$n - k - 1 = n - P$	$MS_E = \frac{SS_E}{n - P} = \hat{\sigma}^2$	
کل	$SS_T = S_{yy}$	$n - 1$		

در صورتی که $F > F_{\alpha, k-1, n-P}$ باشد H_0 رد می‌شود و حداقل یکی از متغیرها نقش قابل توجهی دارد که در آن:

$$S_{yy} = y'y - \frac{\left(\sum_{i=1}^n y_i\right)^2}{n} = \sum_{i=1}^n y_i^2 - n\bar{y}^2$$

$$SS_R = b'x'y - \frac{\left(\sum_{i=1}^n y_i\right)^2}{n} = b'x'y - n\bar{y}^2$$

$$SS_E = S_{yy} - SS_R = y'y - b'x'y$$

$$R^2 = \frac{SS_R}{S_{yy}} = \frac{S_{yy} - SS_E}{S_{yy}} = 1 - \frac{SS_E}{S_{yy}}$$

رابطه F با ضریب تعیین

$$F = \frac{MS_R}{MS_E} = \frac{\frac{SS_R}{k}}{\frac{SS_E}{n-k-1}} = \frac{(n-k-1) \frac{SS_R}{S_{yy}}}{k \cdot \frac{SS_E}{S_{yy}}} = \frac{n-k-1}{k} = \frac{R^2}{1-R^2}$$

$$\Rightarrow F = \frac{n-k-1}{k} \cdot \frac{R^2}{1-R^2} = \frac{n-p}{k} \cdot \frac{R^2}{1-R^2}$$

نکته: MS_E و MS_R یک برآوردگر نااریب برای σ^2 هستند

ضریب تعیین اصلاح شده $(\bar{R}^2)$:

$$\bar{R}^2 = 1 - \frac{n-1}{n-P} (1 - R^2)$$

$$\bar{R}^2 = 1 - \frac{n-1}{S_{yy}} MS_E, \quad MS_E = \frac{SS_E}{n-P}, \quad \bar{R}^2 = 1 - \frac{\sum \frac{(y - \hat{y})^2}{n-P}}{\sum \frac{(y - \hat{y})^2}{n-1}}$$

مثال: در مدل خطی $y = B_1 X_1 + B_2 X_2 + B_3 X_3 + \varepsilon$ تعیین کنید:

(۱) ضرایب خط رگرسیون (۲) واریانس خطا (۳) آزمون فرض $B_1 = 0$ در سطح خطای ۵٪

(۴) جدول آنالیز واریانس

y_i	X_{i1}	X_{i2}	X_{i3}
۱	۰	۰	۱
۲	۱	۱	۰
-۱	۱	۰	۰
۰	۱	-۱	۰
۰	۰	۰	۰

(۱) جواب: چون B_1 در مدل وجود ندارد در ستون اول X عدد یک وجود ندارد.

$$X = \begin{bmatrix} 0 & 0 & 1 \\ 1 & 1 & 0 \\ 1 & 0 & 0 \\ 1 & 1 & 0 \\ 0 & 1 & 0 \end{bmatrix} \Rightarrow X'X = \begin{bmatrix} 3 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 1 \end{bmatrix} \Rightarrow (X'X)^{-1} = \begin{bmatrix} 1/3 & 0 & 0 \\ 0 & 1/2 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

$$x'y = \begin{bmatrix} 0 & 1 & 1 & 1 & 0 \\ 1 & 1 & 0 & -1 & 0 \\ 1 & 0 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 1 \\ 2 \\ -1 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 2 \\ 1 \end{bmatrix}$$

$$\hat{B} = (x'x)^{-1} x'y = \begin{bmatrix} 1/3 & 0 & 0 \\ 0 & 1/2 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 2 \\ 1 \end{bmatrix} = \begin{bmatrix} 1/3 \\ 1 \\ 1 \end{bmatrix} = \begin{bmatrix} B_1 \\ B_2 \\ B_3 \end{bmatrix}$$

(۲)

$$y'y = \sum y_i^2 = 6$$

$$b'x'y = \begin{bmatrix} 1/3 & 1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 2 \\ 1 \end{bmatrix} = \frac{10}{3}$$

$$SSE = y'y - b'x'y = 6 - \frac{10}{3} = \frac{8}{3}$$

$$\sigma^2 = \frac{SSE}{n-p} = \frac{8/3}{5-(3+1)} = \frac{2}{3}$$

(۳)

$$t = \frac{b_1}{\sqrt{\sigma^2 C_{jj}}} = \frac{1/3}{\sqrt{2/3 \times 1/3}} = 0.35$$

چون $t = 0.35 > t_{0.25,1} = 12/70.6$ ، بنابراین H_0 پذیرفته می‌شود و B_1 با احتمال ۹۵٪ برابر صفر است.

(۴)

منبع تغییرات	ss	df	$Ms = \frac{SS}{df}$	F
رگرسیون	$SS_R = \frac{7}{3}$	۳	$\frac{7}{9}$ $\frac{8}{3}$	$F = \frac{\frac{7/9}{1/3}}{2/3} = \frac{7}{24}$
خطا	$SS_R = \frac{8}{3}$	۱		
کل	$SS = 5$	۴		

آزمون معنادار بودن رگرسیون: با استفاده از جدول ANOVA رگرسیون این آزمون را انجام می‌دهیم

منبع تغییر پذیری SOV	مجموع مربعات SS	درجات آزادی df	میانگین مربعات $MS = \frac{SS}{df}$	آماره آزمون F
رگرسیون	SS_{reg}	k	$MS_{reg} = \frac{SS_{reg}}{k}$	$F_{\alpha} = \frac{MS_{reg}}{MS_{res}}$
خطا	SS_{res}	$n - k - 1$	$MS_{res} = \frac{SS_{res}}{n - k - 1} = \hat{\sigma}^2$	
کل	$S_{yy} = SS_T$	$n - 1$		

اگر $F_{\alpha, k, N-K-1} > F_{\alpha}$ فرض صفر رد میشود و حداقل یکی از شیب خط‌ها معنی دار است

نکته ۱۲: مجذور اتا (η^2) شدت دو متغیر وقتی متغیر وابسته کمی و متغیر مستقل کیفی باشد را نشان میدهد

نکته ۱۳: میزان تغییرات Y نسبت به X را ضریب تعیین ($V = R^2$) مشخص میکند.

$$R^2 = \frac{SS_{reg}}{SS_T} = 1 - \frac{SS_{res}}{SS_T} \quad \text{و} \quad 1 - R^2 = \text{ضریب بیگانگی}$$

نکته ۱۴: به منظور بررسی مسئله چند گانگی خطی از دو آماره رواداری یا تحمل و عامل تورم واریانس

$$(VIF) \text{ استفاده میشود و رابطه } VIF = \frac{1}{1 - R^2} \text{ بدست می آید}$$

مثال: نتایج حاصل از سه متغیر پیش بین با تعداد ۱۰ عدد در جدول آنوا زیر ثبت شده. جدول را کامل کنید و توضیح دهید فرض صفر تایید میشود؟ ضریب تعیین ضریب بیگانگی را بدست آورید.

منبع تغییر پذیری SOV	مجموع مجذورات SS	درجات آزادی df	میانگین مربعات $MS = \frac{SS}{df}$	آماره آزمون F
رگرسیون	۲۰			
خطا				
کل	۵۰			

جواب:

منبع تغییر پذیری SOV	مجموع مربعات SS	درجات آزادی df	میانگین مربعات $MS = \frac{SS}{df}$	آماره آزمون F
رگرسیون	۲۰	$K = 3$	$MS_{reg} = \frac{20}{3} = 6.67$	$F_{\alpha} = \frac{MS_{reg}}{MS_{res}}$ $F_{\alpha} = \frac{6.67}{5} = 1.33$
خطا	$30 = 50 - 20$	۶	$MS_{res} = \frac{30}{6} = 5$	
کل	۵۰	$n - 1 = 10 - 1 = 9$		

چون $F_{\alpha, 3, 6} = 4.75 > F_M = 1.34$ فرض صفر تایید میشود و از شیب خط ها معنی دار نیست با احتمال ۹۵٪

$$R^2 = \frac{SS_{reg}}{SS_T} = \frac{20}{50} = 0.4 \quad \text{ضریب بیگانگی} = 1 - R^2 = 1 - 0.4 = 0.6$$

فرض های رگرسیون چند متغیری :

۱- نمره ها مستقل باشند ۲- نمره های موجود در هر خانه از جامعه ای انتخاب شده باشند

۳- خطی باشند برای این کار تمام متغیرهای مستقل دیگر را ثابت نگه میدارند

محاسبات رگرسیون چند متغیری:

$$1) b_1 = \frac{r_{y1} - r_{y2}r_{12}}{1 - r_{12}^2} \left(\frac{S_y}{S_1} \right)$$

$$b_2 = \frac{r_{y2} - r_{y1}r_{12}}{1 - r_{12}^2} \left(\frac{S_y}{S_2} \right)$$

$$2) a = \bar{y} - b_1 \bar{x}_1 - b_2 \bar{x}_2$$

$$3) y = a + b_1 x_1 + b_2 x_2 + \dots$$

ضرایب استاندارد شده:

$$\beta_1 = b_1 \frac{S_y}{S_1} \quad , \quad \beta_2 = b_2 \frac{S_y}{S_2}$$

$$R^2 = \frac{r_{y1}^2 + r_{y2}^2 - 2r_{y1}r_{y2}r_{12}}{1 - r_{12}^2} \quad \text{مجذور همبستگی (ضریب تعیین):}$$

ضریب تعیین تعدیل شده : در نتایج پژوهش ها درج می شود چون تاثیر حجم نمونه در آن نشان داده شده است

$$R_{adj}^2 = 1 - \left[(1 - R^2) \cdot \left(\frac{N-1}{N-K-1} \right) \right]$$

مثال : اگر در نتایج حاصل از سه متغیر پیش بین با تعداد ۱۰ عدد ضریب تعیین ۰/۸ باشد ضریب تعیین تعدیل شده را بدست آورید.

$$N = 10, K = 3, R^2 = 0.8$$

$$R_{adj}^2 = 1 - \left[(1 - R^2) \cdot \left(\frac{N-1}{N-K-1} \right) \right] = 1 - \left[(1 - 0.8) \cdot \left(\frac{10-1}{10-3-1} \right) \right] = 1 - 0.3 = 0.7$$

$$F = \frac{\frac{R^2}{K}}{\frac{1-R^2}{N-K-1}} \quad \text{یا} \quad F = \frac{n-k-1}{k} \cdot \frac{R^2}{1-R^2} = \frac{n-p}{k} \cdot \frac{R^2}{1-R^2} : F \text{ آماره آزمون}$$

مثال : در سوال قبل آماره آزمون را بدست آورید

$$F = \frac{\frac{R^2}{K}}{\frac{1-R^2}{N-K-1}} = \frac{\frac{.4}{3}}{\frac{1-.4}{10-3-1}} = \frac{\frac{.4}{3}}{\frac{.6}{6}} = \frac{.4/3}{.1} = 4$$

چون $F_{.05,3,6} = 4.75$ و $F_{\alpha} = 4 < F_{.05,3,6}$ فرض صفر رد میشود و حداقل یکی از شیب خط ها معنی دار است با احتمال ۹۵٪

نکته : به منظور بررسی مسئله چند گانگی خطی از دو آماره رواداری یا تحمل و عامل تورم واریانس (VIF) استفاده میشود

انواع روش های ورود داده ها در رگرسیون :

۱-همزمان (Enter) ۲- گام به گام (Stepwise) ۳- روش حذف (Remove)

۴-روش پس رونده (Backward) ۵-پیش رونده (Forward)

رگرسیون سلسله مراتبی : رگرسیون سلسله مراتبی راهی است برای نشان دادن اینکه آیا کمیت های مورد علاقه شما پس از قرار گرفتن سایر کمیت ها در مدل، مقدار قابل توجهی از واریانس کمیت وابسته (وابسته) منظور همان ضریب تعیین R^2 است (را توضیح می دهند یا خیر. در این روش، چندین مدل رگرسیون را با اضافه کردن کمیت هایی به مدل قبلی در هر مرحله می سازید. رگرسیون سلسله مراتبی (Hierarchical regression) نشان می دهد متغیرهای پیش بین پس از ورود متغیرهای دیگر به صورت معنادار تغییرات متغیر وابسته را توجیه می کنند. این چارچوب بیشتر برای مدل های مقایسه ای مناسب است تا روش های آماری. یکی از کاربردهای اساسی این روش در حل مسائل متغیر تعدیل کننده است. به منظور بررسی نقش تعدیل کنندگی یک متغیر مثلا Z در رابطه بین دو متغیر X و Y از رگرسیون سلسله مراتبی استفاده می شود. رگرسیون سلسله مراتبی یا ترتیبی این امکان را فراهم می آورد که تاثیر چند متغیر مستقل بر یک متغیر وابسته طی چند مرحله مشخص شود.

تحلیل مسیر (Path analysis): تحلیل مسیر تکنیکی آماری است که با استفاده از معادلات رگرسیون خطی استاندارد (معادلات رگرسیون خطی بر حسب ضرایب رگرسیون استاندارد) به میزان انطباق یک مدل علی نظری با واقعیت (داده ها) می پردازد. در این روش از ضریب بتای استاندارد رگرسیون جهت تعیین جهت و شدت روابط میان متغیرها استفاده می شود. مقدار آماره تی نیز معناداری روابط را نشان می دهد. به عبارت دیگر تحلیل مسیر تکنیکی برای آزمون تجربی مدل علی نظری است. علاوه بر این اگر مدل علی نظری با داده های جمعیتی معین انطباق داشته باشد می توان انواع اثرات (اثر مستقیم و غیرمستقیم و کاذب و خالص) تک تک متغیرهای مستقل بر متغیر وابسته مدل علی نظری را نیز حساب

کرد. هدف تحلیل مسیر به دست آوردن برآوردهای کمی روابط علی (همکنشی یکجانبه یا کواریته) بین مجموعه ای از متغیرهاست. ساختن یک مدل علی لزوماً به معنای وجود روابط علی در بین متغیرهای مدل نیست بلکه این علیت بر اساس مفروضات همبستگی و نظر و پیشینه تحقیق استوار است. اغلب برای انجام محاسبات مربوط به تحلیل مسیر می توان از نرم افزارهای Amos یا Lisrel یا SPSS استفاده کرد. تحلیل مسیر جهت و شدت روابط متغیرهای تحقیق را نشان می دهد. مقادیری که جهت و میزان تاثیر میان متغیرها را نشان می دهند ضریب مسیر نامیده می شوند ضرایب مسیر همان ضریب استاندارد شده (β) رگرسیون هستند. بنابراین برای تحلیل مسیر باید از رگرسیون خطی ساده استفاده شود. تحلیل مسیر تنها بر روی متغیرهای قابل مشاهده انجام پذیر است و اگر بخواهید بین ابعاد تحلیل مسیر را اجرا کنید باید میانگین سوالات هر بعد را حساب کنید تا متغیر پنهان به یک متغیر قابل مشاهده تبدیل شود.

فصل هفتم

تحلیل واریانس

تحلیل واریانس: هرگاه میانگین بیش از دو جامعه را بخواهیم مقایسه کنیم به جای آزمون t از تحلیل

واریانس استفاده می‌کنیم زیرا برای مقایسه K جامعه با یکدیگر باید $\binom{k}{2} = \frac{k(k-1)}{2}$ آزمون t انجام

دهیم.

تحلیل واریانس یک عامله: در این تحلیل مدل در نظر گرفته شده به صورت $x_{ij} = \mu + \alpha_i + e_{ij}$ یامی باشد

که در آن μ میانگین کل α_i اثر تیمار و e_{ij} خطا می‌باشد. که جدول آنالیز واریانس تحلیل یک عاملی به صورت زیر است

$$\left\{ \begin{array}{l} H_0: \mu_1 = \mu_2 = \dots = \mu_k \\ H_1: \mu_i \text{ ها متفاوت است} \end{array} \right. \quad \text{یا حداقل یکی از} \quad \left\{ \begin{array}{l} H_0: \alpha_1 = \alpha_2 = \dots = \alpha_k = 0 \\ H_1: \alpha_i \text{ ها متفاوت است} \end{array} \right.$$

جدول آنوا ANOVA

متغیرات	SS	df	$MS = \frac{SS}{df}$	F
بین گروهها (تیمار)	SSb	$k-1$	$MS_b = \frac{SS_b}{k-1}$	$F = \frac{MS_b}{MS_w}$
درون گروهها (خطا)	$SS_w = SS_{res}$	$k(n-1) = N-k$	$MS_w = \frac{SS_w}{k(n-1)}$	
کل	$SS = SS_T$	$N-1$		

اگر $F_{\alpha, k-1, k(n-1)} > F_{\alpha}$ باشد H_0 رد می‌شود و حداقل یکی از μ_i ها متفاوت است

$$t_{df}^2 = F_{\alpha, df}$$

نکته ۱:

K : تعداد تیمار n_i : تعداد هر تیمار N : تعداد کل i : جمع هر تیمار T : جمع کل

$$\bar{x}_{..} = \frac{\sum_{i=1}^k \sum_{j=1}^n x_{ij}}{N}, \bar{x}_{i.} = \frac{\sum_{j=1}^n x_{ij}}{n_i} \Rightarrow \bar{x}_{..} = \frac{\sum_{i=1}^k \bar{x}_{i.}}{k}$$

$\bar{x}_{..}$: میانگین کل

$$S_{Tr} = SS_B = n \sum_{i=1}^k (\bar{x}_{i.} - \bar{x}_{..})^2 = \sum_{i=1}^k \frac{T_i^2}{n_i} - \frac{T^2}{N}$$

$$SS_W = \sum_{i=1}^k \sum_{j=1}^n (x_{ij} - \bar{x}_{i.})^2 = \sum_{i=1}^k (n_i - 1) s_i^2, SS_t = \sum_{i=1}^k \sum_{j=1}^n (x_{ij} - \bar{x}_{..})^2 = \sum_{i=1}^k \sum_{j=1}^n x_{ij}^2 - \frac{T^2}{N}$$

$$\eta^2 = \frac{SS_B}{SS_T}$$

نکته ۱: در صورتیکه $\eta^2 \leq 0.5$ باشد مجموع مجذورات درون گروهی بیشتر از مجموع مجذورات بین گروهی است و بالعکس

نکته ۲: تفاوت های فردی در طرح اندازه گیری مکرر یک عاملی کنترل میشود (تکرار سنجش)

توجه: محدود η^2 به صورت زیر است

۱- η^2 حدود ۰/۰۱ یعنی اثر کم $\eta^2 - 2$ حدود ۰/۰۶ یعنی اثر متوسط $\eta^2 - 3$ حدود ۰/۱۴ یعنی اثر زیاد

مثال: در سطح معنی دار ۵٪ آزمون کنید که آیا اختلاف معنی داری بین میانگینها سه جامعه زیر وجود دارد.

A	۳	۴	۲	۵	۱۴
B	۲	۱	۶	۳	۱۲
C	۴	۹	۷	۶	۲۶

$$T = \sum_{i=1}^k T_i = 14 + 12 + 26 = 52$$

$$SS_{Tr} = \sum_{i=1}^k \frac{T_i^2}{n_i} - \frac{T_r^2}{N} = \frac{14^2}{4} + \frac{12^2}{4} + \frac{26^2}{4} - \frac{52^2}{12} = 28/67$$

$$SS_T = \sum_{i=1}^k \sum_{j=1}^n x_{ij}^2 - \frac{T_r^2}{N} = 3^2 + 4^2 + 2^2 + 5^2 + \dots + 7^2 + 6^2 - \frac{52^2}{12} = 60/67$$

$$SS_E = SS_T - SS_{Tr} = 60/67 - 28/67 = 32$$

منبع	SS	df	$MS = \frac{SS}{df}$	F
تیمار	۶/۲۸	۳-۱=۲	$\frac{۲۸/۶۷}{۲} = ۱۴/۳۳۵$	$F = \frac{MS_{Tr}}{MS_E} = \frac{۱۴/۳۳۵}{۳/۲} = ۴/۴۷۹$ چون $F = ۴/۴۷۹ > F_{\cdot, ۵, ۲, ۱۰} = ۴/۱۰۲۸$
خطا	۳۲	۱۲-۳=۱۰	$\frac{۳۲}{۱۰} = ۳/۲$	
کل	۶۷/۶۰	۱۲-۱=۱۱		

چون $F = 4/479 > F_{0.05, 2, 10} = 4/1028$ است بنابراین H_0 رد می شود و یکی از اینها با احتمال ۹۴٪ متفاوت است

مثال: در آزمون تحلیل واریانس برای مقایسه سه گروه مستقل در صورتیکه $SS_T = 670$ و $SS_B = 100$ و تعداد هر گروه مساوی 20 نفر باشد اندازه F چند است؟

متغیرات	SS	df	$MS = \frac{SS}{df}$	F
بین گروهها (تیمار)	100	$3-1=2$	$MS_B = \frac{100}{2} = 50$	$F = \frac{MS_B}{MS_W} = \frac{50}{10} = 5$ گزینه 3 صحیح است
درون گروهها (خطا)	$670-100=570$	$60-3=57$	$MS_W = \frac{570}{57} = 10$	
کل	670	$60-1=59$		

تحلیل واریانس دو عامله یک مشاهده به ازای هر ترکیب (بدون تاثیر متقابل): مدل در نظر گرفته شده به صورت $x_{ij} = \mu + \alpha_i + \beta_j + e_{ij}$ یا می باشد که در آن رابطه خطی زیر برقرار است

K: تعداد تیمار (ستون) n: تعداد بلوکها (سطر)

$\begin{cases} H_0: \beta_1 = \beta_2 = \dots = \beta_n = 0 \\ H_1: \text{حداقل یکی از } \beta \text{ ها مخالف صفر است} \end{cases}$
 $\begin{cases} H_0: \alpha_1 = \alpha_2 = \dots = \alpha_k = 0 \\ H_1: \text{حداقل یکی از } \alpha \text{ ها مخالف صفر است} \end{cases}$

وسیله آموزش			نوع درس	نوع درس
تلویزیون	کامپیوتر	رادیو		
5	3	4	ریاضی	
4	7	3	فیزیک	
5	6	5	ادبیات	
6	4	5	عربی	

جدول تحلیل واریانس رده بندی دو طرفه بدون تأثیر متقابل

متغیرات	SS	df	$MS = \frac{SS}{df}$	F
تیمار (ستون)	SS_A	$k - ۱$	$MS_A = \frac{SS_A}{k - ۱}$	$F_A = \frac{MS_A}{MS_w}$ $F_B = \frac{MS_B}{MS_w}$
بلوک (سطر)	SS_B	$n - ۱$	$MS_B = \frac{SS_B}{n - ۱}$	
خطا	$SS_{res} = SS_w$	$(n - ۱)(k - ۱)$	$MS_w = \frac{SS_w}{(n - ۱)(k - ۱)}$	
کل	$SS = SS_T$	$N - ۱$		

اگر $F_A > F_{\alpha, k-1, (n-1)(k-1)}$ باشد H_0 رد می شود و حداقل یکی از α_i ها مخالف صفر است

اگر $F_B > F_{\alpha, n-1, (n-1)(k-1)}$ باشد H_0 رد می شود و حداقل یکی از β_j ها مخالف صفر است

جدول تحلیل واریانس رده بندی دو طرفه بدون تأثیر متقابل

متغیرات	SS	Df	$SS^* = Ms = \frac{SS}{df}$	F
تیمار	$SS_{\tau} = SS_{Tr}$	$k - ۱$	$MS_{Tr} = \frac{SS_{Tr}}{k - ۱}$	$F_{Tr} = \frac{MS_{Tr}}{MS_E}$ $F_B = \frac{MS_B}{MS_E}$
بلوک	$SS_{\beta} = SS_B$	$n - ۱$	$MS_B = \frac{SS_B}{n - ۱}$	
خطا	$SS_{\gamma} = SS_E$	$(n - ۱)(k - ۱)$	$MS_E = \frac{SS_E}{(n - ۱)(k - ۱)}$	
کل	$SS = SS_T$	$N - ۱$		

اگر $F_B > F_{\alpha, k-1, (n-1)(k-1)}$ باشد H_0 رد می شود و حداقل یکی از α_i ها مخالف صفر استاگر $F_B > F_{\alpha, n-1, (n-1)(k-1)}$ باشد H_0 رد می شود و حداقل یکی از B_j ها مخالف صفر است

$$SS_{Tr} = n \sum_{i=1}^k (\bar{x}_{i.} - \bar{x}_{..})^2 = \sum_{i=1}^k \frac{T_{i.}^2}{n_i} - \frac{T_{..}^2}{N}$$

$$SS_B = k \sum_{j=1}^n (\bar{x}_{.j} - \bar{x}_{..})^2 = \sum_{j=1}^n \frac{T_{.j}^2}{k} - \frac{T_{..}^2}{N}$$

$$SS_T = \sum_{i=1}^k \sum_{j=1}^n (x_{ij} - \bar{x}_{..})^2 = \sum_{i=1}^k \sum_{j=1}^n x_{ij}^2 - \frac{T_{..}^2}{N}$$

مثال: در سطح خطای ۵٪ آزمون کنید در تعداد فروش کالای خاص، با توجه فصول مختلف و نوع وسیله تبلیغاتی مختلف تفاوت معنی داری وجود دارد؟

T_j	وسیله تبلیغ			فصل
	تلویزیون	روزنامه	رادیو	
۱۲	۵	۳	۴	بهار
۱۴	۴	۷	۳	تابستان
۱۶	۵	۶	۵	پاییز
۱۵	۶	۴	۵	زمستان
	۲۰	۲۰	۱۷	$T_{i.}$

$$T = \sum_{i=1}^k T_{i.} = 20 + 20 + 17 = 57$$

$$SS_{Tr} = \sum_{i=1}^k \frac{T_i^2}{n_i} - \frac{T^2}{N} = \frac{17^2}{4} + \frac{20^2}{4} + \frac{20^2}{4} - \frac{57^2}{4} = 1/5$$

$$SS_B = \sum_{j=1}^k \frac{T_j^2}{n_j} - \frac{T^2}{N} = \frac{12^2}{3} + \frac{14^2}{3} + \frac{16^2}{3} + \frac{15^2}{3} - \frac{57^2}{12} = 2/91$$

$$SS_T = \sum \sum x_{ij}^2 - \frac{T^2}{N} = 4^2 + 3^2 + \dots - \frac{57^2}{12} = 16/25$$

$$SS_E = SS_T - SS_A - SS_{Tr} = 16/25 - 1/5 - 2/91 = 11/84$$

منبع تغییرات (تیمار)	SS	Df	MS = $\frac{SS}{df}$	MS
فصل (تیمار)	1/5	2	0/75	$F_{Tr} = \frac{0/75}{1/97} = 0/25$ $F_B = \frac{0/97}{1/97} = 0/49$
بلوک	2/91	3	0/97	
خطا	11/84	6	1/97	
کل	16/25	11		

چون $H_0, F_{Tr} = 0/25 < F_{0.05, 2, 6} = 5/1433$ پذیرفته می‌شود و وسیله تبلیغ با احتمال ۹۵٪ مؤثر نیست

چون $H_0, F_B = 0/49 < F_{0.05, 3, 6} = 4/7571$ پذیرفته می‌شود و فصل با احتمال ۹۵٪ مؤثر نیست

تحلیل واریانس دو عامله یا m مشاهده به ازای هر ترکیب (طرح فاکتوریل): مدل در نظر گرفته شده به صورت $x_{ijr} = \mu + \alpha_i + \beta_j + (\alpha\beta)_{ij} + \varepsilon_{ijr}$ می‌باشد که رابطه‌های زیر برقرار است. طرح فاکتوریل کامل، که به عنوان طرح کاملاً متقاطع نیز شناخته می‌شود، به یک طرح آزمایشی اطلاق می‌شود که از دو یا چند عامل تشکیل شده است که هر عامل دارای مقادیر یا «سطوح» ممکن مجزای متعددی است. با استفاده از این طرح می‌توان تمامی ترکیبات ممکن سطوح عامل را در هر تکرار بررسی کرد. اثر متقابل موقعی اتفاق می‌افتد که اثر یک متغیر مستقل بر متغیر وابسته بسته به سطح یک متغیر مستقل دیگر تغییر کند. این اصطلاح بیشتر برای روش‌های طراحی آزمایش به کار می‌رود. اثر متقابل بر حسب تعداد متغیرهای اصلی و برهم کنش‌ها به دو سطحی، سه سطحی و.....

نوع درس	وسیله آموزش		
	رادیو	کامپیوتر	تلویزیون
ریاضی	۴	۳	۵
	۳	۵	۳
	۵	۹	۲
فیزیک	۳	۷	۴
	۶	۴	۲
	۶	۵	۳
ادبیات	۵	۶	۵
	۳	۸	۶
	۷	۹	۹
عربی	۵	۴	۶
	۶	۷	۷
	۸	۶	۹

جدول تحلیل واریانس رده بندی دو طرفه با تاثیر متقابل ANOVA جدول آنوا

منبع متغیرات	SS	df	$SS^* = MS = \frac{SS}{df}$	F
تیمار	SS_A	$k - 1$	$MS_A = \frac{SS_A}{k - 1}$	$F_A = \frac{MS_A}{MS_E}$ $F_B = \frac{MS_B}{MS_E}$ $F_{A.B} = \frac{MS_{A.B}}{MS_W}$
بلوک	SS_B	$n - 1$	$MS_B = \frac{SS_B}{n - 1}$	
تاثیر متقابل	$SS_{A.B}$	$(n - 1)(k - 1)$	$MS_{A.B} = \frac{SS_{A.B}}{(k - 1)(n - 1)}$	
خطا	SS_W	$kn(m - 1)$	$MS_W = \frac{SS_W}{kn(m - 1)}$	
کل	$SS_T = SS_{yy}$			

اگر $F_A > F_{\alpha, k-1, kn(m-1)}$ باشد فرض H_0 (یعنی $\alpha_i = 0$) رد می‌شود و حداقل یکی از تیمارها متفاوت است اگر $F_B > F_{\alpha, n-1, kn(m-1)}$ باشد فرض H_0 (یعنی $\beta_j = 0$) رد می‌شود و حداقل یکی از بلوک‌ها متفاوت است

اگر $F_{A.B} > F_{\alpha, (k-1)(n-1), kn(m-1)}$ باشد فرض H_0 (یعنی $(\alpha\beta)_{ij} = 0$) رد می‌شود و تاثیر متقابل وجود دارد

$$SS_{Tr} = nm \sum_{i=1}^k (\bar{x}_{i..} - \bar{x}_{...})^2 = \sum_{i=1}^k \frac{T_{i..}^2}{n_i} - \frac{T^2}{N}$$

$$SS_B = km \sum_{j=1}^m (\bar{x}_{.j} - \bar{x}_{...})^2 = \sum_{j=1}^n \frac{T_{.j}^2}{k} - \frac{T^2}{N}$$

$$SS_{Tr.B} = m \sum_{i=1}^k \sum_{j=1}^n (\bar{x}_{ij.} - \bar{x}_{i..} - \bar{x}_{.j} + \bar{x}_{...})^2$$

$$SS_E = \sum_{i=1}^k \sum_{j=1}^n \sum_{r=1}^m (\bar{x}_{ijr} - \bar{x}_{ij.})^2 = \sum_{i=1}^k \sum_{j=1}^n \sum_{r=1}^m x_{ijr}^2 - \frac{\sum_{i=1}^k \sum_{j=1}^n x_{ij.}^2}{m}$$

$$SS_T = \sum_{i=1}^k \sum_{j=1}^n \sum_{r=1}^m (x_{ijr} - \bar{x}_{...})^2 = \sum_{i=1}^k \sum_{j=1}^n \sum_{r=1}^m x_{ijr}^2 - \frac{T^2}{N}$$

$$SS_{Tr.B} = SS_T - SS_{Tr} - SS_B - SS_E$$

مثال: در سطح خطای ۵٪ آزمون کنید

الف) تبلیغات در میانگین فروش مؤثر است

ب) تخفیف در میانگین فروش مؤثر است

ج) تاثیر متقابلی بین قیمت گذاری و سیاست فروش وجود دارد

سیاست فروش		قیمت گذاری
بدون تبلیغ	با تبلیغ	
۶	۱۰	با تخفیف
۵	۱۱	
۵	۶	بدون تخفیف
۴	۷	

ابتدا جمع در هر سلول را محاسبه می کنیم و مانند بدون تکرار مشاهدات SS_B, SS_{Tr} را محاسبه کنید.

$T_{.j}$	بدون تبلیغ	با تبلیغ	
۳۲	۱۱	۲۱	با تخفیف
۲۲	۹	۱۳	بدون تخفیف
	۲۰	۳۴	$T_{i.}$

$$SS_{Tr} = \sum_{i=1}^k \frac{T_i^2}{n_i} - \frac{T^2}{N} = \frac{34^2}{4} + \frac{20^2}{4} - \frac{54^2}{8} = 24/5$$

$$SS_B = \sum_{j=1}^n \frac{T_j^2}{k} - \frac{T^2}{N} = \frac{32^2}{4} + \frac{22^2}{4} - \frac{54^2}{8} = 12/5$$

$$SS_T = \sum_{i=1}^k \sum_{j=1}^n \sum_{r=1}^m x_{ijr}^2 - \frac{T^2}{N} = 10^2 + 11^2 + \dots + 5^2 + 4^2 - \frac{54^2}{8} = 43/5$$

$$SS_E = \sum_{i=1}^k \sum_{j=1}^n \sum_{r=1}^m x_{ijr}^2 - \frac{\sum_{i=1}^k \sum_{j=1}^n x_{ij}^2}{m} = 10^2 + 11^2 + \dots + 5^2 + 4^2 - \frac{11^2 + 21^2 + 9^2 + 13^2}{2} = 2$$

$$SS_{Tr.B} = SS_T - SS_{Tr} - SS_B - SS_E = 43/5 - 24/5 - 12/5 - 2 = 4/5$$

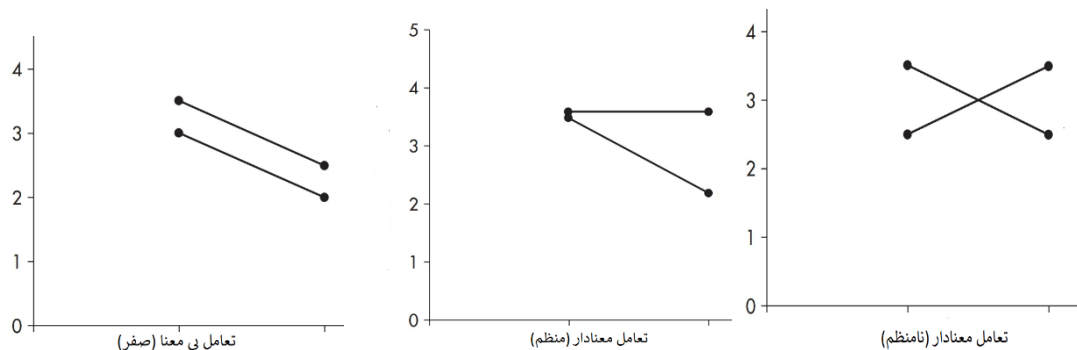
منبع تغییرات (تیمار)	SS	df	MS = $\frac{SS}{df}$	MS
تیمار	24/5	1	24/5	$F_{Tr} = \frac{24/5}{0/5} = 49$
بلوک	12/5	1	12/5	
تأثیر متقابل	4/5	1	4/5	$F_B = \frac{12/5}{0/5} = 25$
خطا	2	4	0/5	$F_{Tr.B} = \frac{4/5}{0/5} = 9$
کل	43/5	7		

چون $F_{Tr} = 49 > F_{0.05,1,4} = 7/7086$ می‌باشد بنابراین H_0 رد می‌شود و با اطمینان ۹۵٪ تبلیغ مؤثر است.

چون $F_B = 25 > F_{0.05,1,4} = 7/7086$ می‌باشد بنابراین H_0 رد می‌شود و با اطمینان ۹۵٪ شیوه قیمت گذاری مؤثر است.

چون $F_{Tr.B} = 9 > F_{0.05,1,4} = 70/7086$ می‌باشد بنابراین H_0 رد می‌شود و با اطمینان ۹۵٪ تأثیر متقابل وجود دارد.

نکته: اگر دو خط یکدیگر را قطع کنند تعامل (تأثیر متقابل) معنا دار است و اگر دو خط موازی یکدیگر باشند تعامل بی معنا است. و در صورتی که اثر تعاملی نامنظم وجود داشته باشد باید تفسیر متغیرهای مستقل با احتیاط صورت گیرد.



متغیر کنترل: به متغیر مزاحمی گفته می شود که محقق تصمیم می گیرد اثر آن را خنثی یا حذف کند.

محقق به دو روش می تواند اثر متغیر مزاحم را کنترل کند:

۱- کنترل پژوهشی: با ثابت نگه داشتن متغیر در یک سطح، اثر آن خنثی می شود. مثلاً فقط از بین دانش آموزان دختر نمونه گیری انجام می شود.

۲- کنترل آماری: به کمک روش آماری تحلیل کواریانس، اثر متغیر مزاحم مطالعه و از روی نتایج حذف می شود.

فصل هشتم

ناپارامتری

آزمون‌های ناپارامتری:

هر گاه توزیع جامعه مشخص نباشد از آزمون‌های ناپارامتری (آمار توزیع آزاد) استفاده می‌شود که بر اساس تعداد می‌باشد. که دقت آن نسبت به آمار پارامتری کمتر است.

آزمون علامت یک نمونه‌ای:

از این آزمون برای آزمون میانه یا چارک‌ها و اگر توزیع جامعه متقارن باشد آزمون میانگین هم می‌توان استفاده کرد و سپس $(X_i - \theta_0)$ را محاسبه می‌کنیم اگر مثبت باشد علامت + و اگر منفی باشد

علامت - قرار می‌دهیم که پارامترهای توزیع دوجمله‌ای می‌باشد که برای میانه $P = \frac{1}{2}$ و برای چارک‌ها

$\frac{1}{4}$ یا $\frac{3}{4}$ قرار می‌دهیم که احتمال موردنظر است و اگر $nP > 5$, $nq > 5$ باشد می‌توان از تقریب توزیع

نرمال استفاده کرد.

$$z = \frac{x - np}{\sqrt{npq}}$$

مثال: در یک جامعه تاکنون میانه برابر ۴۰ بوده است در سال جاری به علت تغییرات احساس می‌شود

میانه تغییر کرده است به این منظور نمونه‌ای ۱۴ تایی انتخاب کرده و اعداد

۴۰-۴۵-۴۷-۴۰-۳۳-۴۲-۵۰-۳۹-۴۰-۳۷-۴۹-۴۳-۴۰-۴۶

بدست آمده است در سطح معنی‌داری $\alpha = 0.1$ این اعداد را آزمون کنیم.

جواب: ابتدا $(X_i - \theta_0)$ را بدست می‌آوریم و علامت را تعیین می‌کنیم و اگر حاصل صفر شود علامت

نمی‌گذاریم.

++-++-++

بنابراین $n = 10$, $p = \frac{1}{2}$ یعنی $B(10, \frac{1}{2})$ بنابراین چون تعداد مثبت‌ها بیش از $np = 10 \times \frac{1}{2} = 5$ است

بنابراین:

$$P(X \geq 8) = 0.1 = 1 - P(X \leq 7) = 1 - 0.945 = 0.055 < 0.1$$

چون تعداد علامت مثبت در نمونه یعنی آماره آزمون برابر ۷ و در ناحیه بحرانی قرار نمی‌گیرد و H_0

پذیرفته می‌شود یا می‌توان بر اساس منفی بحث کرد یعنی اولین جایی که احتمال در $n = 10$ و $p = 0.5$ کمتر از ۰/۱ باشد.

$$P(X \leq C) = 0.1 = C_{0.1, 10, 0.5} = 2$$

چون تعداد منفی‌ها بیش از ۲ است بنابراین در ناحیه رد قرار نمی‌گیرد و H_0 پذیرفته می‌شود.

کاربردهای آزمون مربع کای دو (خی‌دو): همانگونه که می‌دانیم توزیع خی دو مجموع مربعات نرمال

استاندارد است $(X^2 = \sum_{i=1}^{df} Z_i^2)$ که میانگین آن برابر درجه آزادی و واریانس دو برابر درجه آزادی است

توزیع خی دو دارای منحنی چوله به راست که در آزمون فرض‌های زیر استفاده می‌شود.

۱- جدولهای توافقی :

که با نام‌های آزمونهای استقلال و آزمونهای همگونی (نسبت‌ها) مطرح می‌شود برای مواقعی که r سطر و c

ستون داشته باشیم که فراوانی مشاهده شده O_{ij} و فراوانی مورد انتظار e_{ij} می باشد که از رابطه

$$e_{ij} = \frac{O_{i.} \times O_{.j}}{n} \quad \text{بدست می آید و آماره آزمون از رابطه } X^2 = \sum_{i=1}^r \sum_{j=1}^c \frac{(O_{ij} - e_{ij})^2}{e_{ij}} \quad \text{بدست می آید و با}$$

عدد بحرانی $\chi^2_{\alpha, df}$ با درجه آزادی $df = (r-1)(c-1)$ مقایسه می شود. در صورتی که $X^2 > \chi^2_{\alpha, df}$ باشد H_0 رد می شود و در نتیجه داده مستقل نیستند یا می توان گفت نسبتها برابر نیست لازم به توضیح است این آزمون در مواقعی بکار می رود که هیچ یک از e_{ij} ها کمتر از ۵ نباشد در این مواقع برخی خانه ها را با یکدیگر ادغام می کنیم.

نکته: در صورتی که درجه آزادی برابر یک شود از تصحیح یتس استفاده می شود و آماره آزمون از رابطه

$$X^2 = \sum \frac{(|F_{oi} - F_{ei}| - 0.5)^2}{F_{ei}} \quad \text{بدست می آید.}$$

مثال: در سطح خطای ۵٪ آزمون کنید میزان تحصیلات و جنسیت مستقل از یکدیگر هستند.

تحصیلات جنسیت	A	B	C	D
خانم	۵۵	۲۰	۱۵	۲۰
آقا	۳۵	۲۰	۲۵	۱۰

تحصیلات جنسیت	A	B	C	D	جمع
خانم	۵۵	۲۰	۱۵	۲۰	۱۱۰
آقا	۳۵	۲۰	۲۵	۱۰	۹۰
جمع	۹۰	۴۰	۴۰	۳۰	

تحصیلات جنسیت	A	B	C	D
خانم	$\frac{110 \times 90}{200} = 49.5$	$\frac{110 \times 40}{200} = 22$	$\frac{110 \times 40}{200} = 22$	$\frac{110 \times 30}{200} = 16.5$
آقا	$\frac{90 \times 90}{200} = 40.5$	$\frac{90 \times 40}{200} = 18$	$\frac{90 \times 40}{200} = 18$	$\frac{90 \times 30}{200} = 13.5$

$$X^2 = \sum_{i=1}^r \sum_{j=1}^c \frac{(O_{ij} - e_{ij})^2}{e_{ij}} = \frac{(55 - 49/5)^2}{49/5} + \frac{(20 - 22)^2}{22} + \frac{(15 - 22)^2}{22} + \frac{(20 - 16/5)^2}{16/5} \\ + \frac{(35 - 40/5)^2}{40/5} + \frac{(20 - 18)^2}{20} + \frac{(25 - 18)^2}{18} + \frac{(10 - 13/5)^2}{13/5} = 8/0.92 \\ df = (r-1)(C-1) = (2-1)(4-1) = 3$$

چون $X^2 = 8/0.92 > \chi^2_{0.05,3} = 7/815$ است بنابراین H_0 رد می‌شود و با احتمال ۹۵٪ با توجه به نمونه جنسیت و میزان تحصیلات مستقل از یکدیگر نیستند.

مثال: در سطح خطای ۵٪ آزمون کنید بازمانده‌های لامپهای خلاء و بلوک بالا و بلوک پائین مستقل از یکدیگر هستند.

نوع بازمانی O_{ij}

	A	B	C	D
بلوک بالا	۵۵	۲۰	۱۵	۲۰
بلوک پائین	۳۵	۲۰	۲۵	۱۰

	A	B	C	D	جمع
بلوک بالا	۵۵	۲۰	۱۵	۲۰	۱۱۰
بلوک پائین	۳۵	۲۰	۲۵	۱۰	۹۰
جمع	۹۰	۴۰	۴۰	۳۰	

	A	B	C	D
بلوک بالا	$\frac{110 \times 90}{200} = 49/5$	$\frac{110 \times 40}{200} = 22$	$\frac{110 \times 40}{200} = 22$	$\frac{110 \times 30}{200} = 16/5$
بلوک پائین	$\frac{90 \times 90}{200} = 40/5$	$\frac{90 \times 40}{200} = 18$	$\frac{90 \times 40}{200} = 18$	$\frac{90 \times 30}{200} = 13/5$

$$X^2 = \sum_{i=1}^r \sum_{j=1}^c \frac{(O_{ij} - e_{ij})^2}{e_{ij}} = \frac{(55 - 49/5)^2}{49/5} + \frac{(20 - 22)^2}{22} + \frac{(15 - 22)^2}{22} + \frac{(20 - 16/5)^2}{16/5} \\ + \frac{(35 - 40/5)^2}{40/5} + \frac{(20 - 18)^2}{20} + \frac{(25 - 18)^2}{18} + \frac{(10 - 13/5)^2}{13/5} = 8/0.92 \\ df = (r-1)(C-1) = (2-1)(4-1) = 3$$

چون $X^2 = 8/0.92 > \chi^2_{0.05,3} = 7/815$ است بنابراین H_0 رد می‌شود و با احتمال ۹۵٪ بازمانده‌های لامپ‌های خلاء و بلوک مستقل از یکدیگر نیستند.

ضریب توافقی C: اگر در آزمون استقلال دو صفت وابسته باشند به کمک ضریب توافقی C می‌توان شدت

$$C = \sqrt{\frac{X^2}{X^2 + n}}$$

این بستگی را مشخص کرد.

$$\Phi = \sqrt{\frac{\chi^2}{n}}$$

ضریب فی :

$$V = \sqrt{\frac{X^2}{N \cdot \min[(K-1), (L-1)]}}$$

ضریب کریمر: که تعمیم ضریب فی است

آزمون نیکویی برازش:

از این آزمون در مواقعی استفاده می‌شود که بخواهیم بدانیم داده‌ها از توزیع خاصی تبعیت می‌کنند یا خیر. که فراوانی مشاهده شده و $F_{ei} = Np$ فراوانی مورد انتظار است و آماره آزمون از رابطه

$$X^2 = \sum_{i=1}^k \frac{(F_{oi} - F_{ei})^2}{F_{ei}}$$

و درجه آزادی از رابطه $df = k - p - 1$ بدست می‌آید که در آن k تعداد طبقات و p

تعداد پارامتری که باید برآورد شود. در صورتی که $X^2 > X^2_{\alpha, df}$ باشد H_0 رد می‌شود و فراوانی نظری و تجربی برابر نیست و داده‌ها از آن توزیع تبعیت نمی‌کنند.

مثال: تعداد کالای تولیدی مشاهده توسط چهار دستگاه به صورت زیر نسبت شده است

A	B	C	D
۲۵۰	۱۱۰	۹۰	۵۰

مدیر کارخانه مدعی است که تولید کالا به نسبت اعداد زیر می‌باشد

A	B	C	D
۸	۷	۶	۴

در سطح اطمینان ۵٪ آزمون کنید تردیدی برای ادعای فوق وجود دارد.

دستگاه	F_{oi}	$F_{ei} = n.p$
A	۲۵۰	$۵۰۰ \times \frac{۱۳}{۲۵} = ۲۶۰$
B	۱۱۰	$۵۰۰ \times \frac{۷}{۲۵} = ۱۴۰$
C	۹۰	$۵۰۰ \times \frac{۶}{۲۵} = ۱۲۰$
D	۵۰	$۵۰۰ \times \frac{۴}{۲۵} = ۸۰$
	۵۰۰	۵۰۰

$$X^2 = \sum_{i=1}^k \frac{(F_{oi} - F_{ei})^2}{F_{ei}} = \frac{(۲۵۰ - ۲۶۰)^2}{۲۶۰} + \frac{(۱۱۰ - ۱۴۰)^2}{۱۴۰} + \frac{(۹۰ - ۱۲۰)^2}{۱۲۰} + \frac{(۵۰ - ۸۰)^2}{۸۰} = ۲۵/۵۶$$

$$df = k - p - 1 = ۴ - ۰ - 1 = ۳$$

چون $X^2 = ۲۵/۵۶ > X^2_{۰.۰۵, ۳} = ۷/۸۱۵$ است بنابراین H_0 رد می‌شود و با احتمال ۹۵٪ نظر مدیر صحیح نیست.

نکته: در صورتی که درجه آزادی برابر یک شود از تصحیح یتس استفاده می‌شود و آماره آزمون از رابطه

$$X^2 = \sum \frac{(|F_{oi} - F_{ei}| - ۰.۵)^2}{F_{ei}}$$
 بدست می‌آید.

مثال: از ۶۰ کالای تولید شده یک کارخانه در یک ماه، ۱۵ کالا معیوب است در سطح خطای ۵٪ آزمون کنید کالاهای معیوب از توزیع برنولی با پارامتر ۰/۲ تبعیت می‌کند.

کیفیت	f_{oi}	P	$F_{ei} = NP$
معیوب	۱۵	۰/۲	۱۲
سالم	۴۵	۰/۸	۴۸
	۶۰		

$$df = k - p - 1 = ۲ - ۰ - 1 = ۱ \Rightarrow$$

$$X^2 = \sum_{i=1}^k \frac{(|F_{oi} - F_{ei}| - 0.5)^2}{F_{ei}} = \frac{(|15 - 12| - 0.5)^2}{12} + \frac{(|45 - 48| - 0.5)^2}{48} = 0.65$$

چون $X^2 = 0.65 < X^2_{0.05,1} = 3.841$ بنابراین H_0 پذیرفته می‌شود و با احتمال ۹۵٪ تعداد کالاهای معیوب از توزیع برنولی با پارامتر $p=0.2$ تبعیت می‌کند.

مثال: نمونه‌ای شامل ۲۰۰ صفحه تایپ شده را بررسی کرده‌ایم و تعداد اشتباهات در صفحات ثبت شده است در سطح خطای ۰/۰۱ آزمون کنید آیا تعداد اشتباهات از توزیع پواسون تبعیت می‌کند

تعداد اشتباهات X	۰	۱	۲	۳	۴	۵
F_{oi} تعداد صفحات	۵۰	۶۰	۵۰	۲۴	۱۲	۴

همانگونه که می‌دانیم در توزیع پواسون پارامتر توزیع برابر میانگین توزیع است.

$$\mu = \frac{\sum F_i X_i}{N} = \frac{0 \times 50 + 1 \times 60 + 2 \times 50 + 3 \times 24 + 4 \times 12 + 5 \times 4}{200} = 1/5$$

با استفاده از رابطه پواسون $P(X=x) = \frac{e^{-\mu} \mu^x}{x!}$ احتمالها را بدست می‌آوریم

X	F_{oi}	$P = \frac{e^{-\mu} \mu^x}{x!}$	$F_{ei} = NP$
۰	۵۰	۰/۲۲۳۱	۴۴/۶۲
۱	۶۰	۰/۳۳۴۷	۶۶/۹۴
۲	۵۰	۰/۲۵۱۰	۵۰/۲
۳	۲۴	۰/۱۲۵۵	۲۵/۱
۴	۱۲	۰/۰۴۷۱	۹/۴۲
۵	۴	۰/۰۱۷۶	۳/۵۲

چون فراوانی نظری رده آخر کمتر از ۵ است بهتر است دو رده آخر را با هم ادغام کنیم

X	F _{oi}	F _{ei}
۰	۵۰	۴۴/۶۲
۱	۶۰	۶۶/۹۴
۲	۵۰	۵۰/۲
۳	۲۴	۲۵/۱
۴≥	۱۶	۹/۴۲

$$X^2 = \sum_{i=1}^k \frac{(F_{oi} - F_{ei})^2}{F_{ei}} = \frac{(50 - 44/62)^2}{44/62} + \frac{(60 - 66/94)^2}{66/94} + \frac{(50 - 50/2)^2}{50/2} + \frac{(24 - 25/1)^2}{25/1} + \frac{(16 - 12/94)^2}{12/94} = 2/14$$

$$df = k - p - 1 = 5 - 1 - 1 = 3$$

چون $X^2 = 2/14 \neq \chi^2_{0.1,3} = 11/345$ بنابراین فرض H_0 پذیرفته می شود و با احتمال ۹۹٪ تعداد اشتباهات از توزیع پواسون تبعیت می کند.

آزمون استقلال در جدول 2×2 :

در جدول 2×2 می توان به صورت زیر آماره آزمون را بدست آورد:

حالت اول

حالت دوم	a	b	a+b
	c	d	c+d

$$X^2 = \frac{(ad - cb)^2 (a + b + c + d)}{(a + b)(a + c)(b + d)(c + d)}$$

اگر $X^2 > X^2_{\alpha,1}$ باشد فرض H_0 رد می شود و مستقل نیستند.

همچنین برای این آزمون می‌توان از آماره Z استفاده کرد.

$$Z = \frac{(ad - cb)\sqrt{a+b+c+d}}{\sqrt{(a+b)(a+c)(b+d)(c+d)}}$$

اگر $|Z| > Z_{\alpha/2}$ باشد فرض H_0 رد می‌شود و مستقل نیستند.

آزمون رتبه علامت‌دار: در صورتی که داده به صورت زوجی باشد استفاده می‌شود که ابتدا اختلاف دو حالت را به دست می‌آوریم سپس قدر مطلق اختلافها را حساب می‌کنیم و آنها را از کوچک به بزرگ رتبه می‌دهیم و در صورتی که اعداد یکسان داشته باشیم رتبه متوسط می‌دهیم و سپس مجموع رتبه‌های

مثبت $\left(T^+\right)$ یا مجموع رتبه‌های منفی $\left(T^-\right)$ را بدست می‌آوریم اگر $n \geq 15$ باشد توزیع T^+ یا T^-

تقریباً نرمال است با میانگین و واریانس $E(T) = \frac{n(n+1)}{4}$ ، $Var(T) = \frac{n(n+1)(2n+1)}{24}$ که در

این صورت آماره آزمون به صورت $Z = \frac{T - E(T)}{\sqrt{Var(T)}}$ بدست می‌آید و بسته به شرایط آزمون، آزمون صورت می‌گیرد.

مثال: ورقه‌های ۱۵ دانشجو توسط دو استاد به طور مستقل شده و نمرات آنها در جدول زیر ثبت شده است در سطح ۵٪ آزمون کنید آیا نمرات دو استاد تفاوت معنی‌دار دارد.

استاد اول (x)	۶	۱۲	۱۹	۱۴	۷	۱۵	۱۰	۱۳	۱۱	۵	۱۸	۲۰	۱۲	۱۷	۱۴
استاد دوم (y)	۸	۱۳	۱۷	۱۳	۶	۱۸	۱۲	۱۱	۱۵	۸	۲۰	۱۹	۱۴	۱۵	۱۶
d=x-y	-۲	-۱	۲	۱	۱	-۳	-۲	۲	-۴	-۳	-۲	-۱	-۲	۲	-۲
قدر مطلق مرتب شده	۱	۱	۱	۱	۲	۲	۲	۲	۲	۲	۲	۲	۳	۳	۴
رتبه فرض	۱	۲	۳	۴	۵	۶	۷	۸	۹	۱۰	۱۱	۱۲	۱۳	۱۴	۱۵
رتبه اصلی	۲/۵	۲/۵	۲/۵	۲/۵	۲/۵	۸/۵	۸/۵	۸/۵	۸/۵	۸/۵	۸/۵	۸/۵	۱۳/۵	۱۳/۵	۱۵

$$T^+ = 2/5 + 2/5 + 8/5 + 8/5 + 8/5 = 30/5$$

$$E(T) = \frac{n(n+1)}{4} = \frac{15 \times 16}{4} = 60$$

$$\text{Var}(T) = \frac{n(n+1)(2n+1)}{24} = \frac{15 \times 16 \times 31}{24} = 310$$

$$Z = \frac{T - E(T)}{\sqrt{\text{Var}(T)}} = \frac{30/5}{\sqrt{310}} = \frac{-29/5}{17/6} = -1/67$$

$$\begin{cases} H_0: \mu_1 = \mu_2 \\ H_1: \mu_1 \neq \mu_2 \end{cases} \quad \begin{cases} H_0: \bar{d} = 0 \\ H_1: \bar{d} \neq 0 \end{cases}$$

چون $|z| = 1/67 \neq z_{0.25} = 1/96$ بنابراین H_0 پذیرفته می‌شود و نمرات در استاد در سطح اطمینان ۹۵٪ با یکدیگر تفاوت معنی‌دار ندارد.

آزمون مجموع رتبه‌ها:

آزمون U من — وتینی یا آزمون ویلکاکسون جایگزین آزمون t در آزمون پارامتری برای مقایسه یکسان بودن میانگین دو جامعه استفاده می‌شود. برای انجام این آزمون ابتدا دو جامعه را با یکدیگر ادغام کرد. و از کوچک به بزرگ رتبه می‌دهیم و در صورت یکسان بودن اعداد رتبه متوسط می‌دهیم سپس مجموع رتبه‌ها هریک از جوامع را بدست می‌آوریم (R_1, R_2) سپس آماره‌های U_1, U_2 را بدست آورده و $\min$ آنها را انتخاب می‌کنیم. اگر $n_1 > 8, n_2 > 8$ باشد در این صورت U تقریباً دارای توزیع نرمال می‌باشد.

$$E(U) = \frac{n_1 n_2}{2}$$

$$\text{Var}(U) = \frac{n_1 n_2 (n_1 + n_2 + 1)}{12}$$

و آماره آزمون از رابطه $Z = \frac{U - E(U)}{\sqrt{\text{Var}(U)}}$ بدست می‌آید.

آزمون مجموع رتبه‌ها آزمون H :

آزمون H یا آزمون کروسکال — والیس تعمیم یافته آزمون U من ویتنی است و برای مقایسه میانگین k جامعه به کار می‌رود و جایگزین تحلیل واریانس در آمار پارامتری است برای انجام این جوامع را با یکدیگر ادغام می‌کنیم و از کوچک به بزرگ مرتب می‌کنیم و رتبه می‌دهیم مانند حالت‌های قبل سپس مجموع رتبه‌های هر جامعه را تعیین می‌کنیم (R_i) که در این صورت آماره آزمون از رابطه

$$H = \frac{12}{n(n+1)} \sum_{i=1}^k \frac{R_i^2}{n_i} - 3(n+1)$$

بدست می‌آید اگر حجم نمونه‌ها حداقل برابر ۵ باشد آماره H دارای

توزیع خی دو χ^2 با درجه آزادی $df = k-1$ است اگر $H > \chi^2_{d,df}$ باشد. H_0 رد می‌شود و میانگین حداقل یک جامعه متفاوت است.

تجزیه و تحلیل فریدمن: در این آزمون داده‌ها حداقل جلید دارای مقیاس رتبه‌ای باشند و تعمیم‌یافته آزمون رتبه علامت‌دار است یعنی روی افراد ثلثت k آزمایش انجام شود برای انجام این آزمون به هر جامعه رتبه می‌دهیم و انتظار داریم که فراوانی رتبه‌ها در k جامعه برابر باشد. بنابراین مجموع رتبه‌های

$$X^2 = \frac{12}{nk(k+1)} \sum_{i=1}^k R_i^2 - 3n(k+1) \quad \text{رابطه ۳}$$

هر جامعه را بدست می‌آوریم (R_i) سپس آماره آزمون را از رابطه ۳ بدست می‌آوریم که n هر جامعه و k تعداد جوامع است اگر $X^2 > x^2_{\alpha, K-1}$ باشد H_0 رد می‌شود و حداقل یکی از جوامع متفاوت است.

مثال: تحقیق کنید آیا نتیجه ارزشیابی ۳ استاد یکسان است در سطح خطای ۵٪.

نمرات

دانشجو	I	II	III
۱	۴۸	۷۵	۱۰۰
۲	۳۰	۱۰۰	۱۰۰
۳	۱۰۰	۸۶	۹۶
۴	۳۳	۵۸	۶۰
۵	۶۰	۶۵	۹۰
۶	۲۰	۶۰	۴۰

رتبه‌ها

دانشجو	I	II	III
۱	۱	۲	۳
۲	۱	۲/۵	۲/۵
۳	۱	۲	۳
۴	۱	۲	۳
۵	۲	۱	۳
۶	۱	۳	۲
Ri	۷	۱۲/۵	۱۶/۵

$$X^2 = \frac{12}{nk(k+1)} \sum_{i=1}^n R_i^2 - 3n(k+1)$$

$$n=6 \quad k=3$$

$$X^2 = \left[\frac{12}{6 \times 3(3+1)} \left[7^2 + 12/5^2 + 16/5^2 \right] \right] - 3 \times 6 \times (3+1) = 7/583$$

$$df = k-1 = 3-1 = 2$$

چون $X^2 = 7/583 > X^2_{0.05,2} = 5.991$ بنابراین H_0 رد می‌شود و ارزشیابی اساتید در سطح خطای ۵٪ یکسان نیست.

آزمون ردیف (گردش):

از این آزمون می‌توان برای آزمون استقلال یا تصادفی بودن ردیف استفاده می‌شود. در صورتی که $n_1 \geq 10, n_2 \geq 10$ باشد می‌توان از تقریب نرمال استفاده کرد:

$$E(R) = \frac{2n_1n_2}{n_1+n_2} + 1$$

$$Var(R) = \frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1+n_2)^2(n_1+n_2-1)}$$

$$Z = \frac{R - E(R)}{\sqrt{Var(R)}}$$

که در آن R تعداد ردیفها یا تغییرات است n_1, n_2 تعداد هریک اعضا است.

نکته: اگر n_1, n_2 کمتر از ۱۰ باشند از جدول مخصوص آزمون ردیف استفاده می‌شود.

مثال: مراجعات بر جنسیت به یک بانک به صورت زیر است در سطح خطای ۵٪ آزمون کنید آیا ورود افراد از لحاظ جنسیت تصادفی است.

$$R = 8$$

$$E(R) = \frac{2n_1n_2}{n_1+n_2} + 1 = \frac{2 \times 11 \times 12}{11+12} + 1 = 12/47$$

$$Var(R) = \frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1+n_2)^2(n_1+n_2-1)} = \frac{2 \times 11 \times 12(2 \times 11 \times 12 - 11 - 12)}{(11+12)^2(11+12-1)} = 5/467$$

$$Z = \frac{R - E(R)}{\sqrt{Var(R)}} = \frac{8 - 12/47}{\sqrt{5/467}} = -1/911$$

چون $|Z| = 1/911 > Z_{0.025} = 1/96$ بنابراین H_0 پذیرفته می‌شود و ورود افراد تصادفی است با احتمال ۹۵٪.

ضریب همبستگی رتبه‌ای اسپیرمن:

در محاسبه ضریب همبستگی پواسون فرض اصلی این است که نمونه‌ها از جامعه نرمال دوبعدی باشد در صورتی که متغیرها کیفی باشند یا کمی و به متعلق به جامعه نرمال نباشند از ضریب همبستگی رتبه‌ای اسپیرمن استفاده می‌شود که برای محاسبه این ضریب به مقادیر y, x از کوچک به بزرگ رتبه می‌دهیم و در صورت یکسان بودن رتبه متوسط می‌دهیم سپس اختلاف رتبه‌های زوج y, x محاسبه می‌کنیم.

$$r_s = 1 - \frac{\sum_{i=1}^k d_i^2}{n(n^2 - 1)} \quad (d_i = r_{xi} - r_{yi})$$

باشد می‌توان از توزیع نرمال تقریب زد و آزمون فرض ضریب همبستگی را انجام می‌دهیم.

$$E(r_s) = 0$$

$$\text{Var}(r_s) = \frac{1}{n-1} \Rightarrow z = \frac{r_s - 0}{\sqrt{\frac{1}{n-1}}} \Rightarrow z = r_s \sqrt{n-1}$$

لازم به توضیح است اگر n کوچک باشد می‌توان از توزیع t برای آزمون فرض استفاده کرد که کمتر استفاده می‌شود.

$$t = \frac{r_s \sqrt{n-2}}{\sqrt{1-r_s^2}}$$

مثال: ضریب همبستگی اسپیرمن برای دو متغیر y, x را محاسبه کنید.

x	۶	۹	۷	۵	۴
y	-۲	-۵	۴	۶	۷
r_x	۳	۵	۴	۲	۱
r_y	۲	۱	۳	۴	۵
$d = r_x - r_y$	۱	۴	۱	-۲	-۴
d^2	۱	۱۶	۱	۴	۱۶
					$\sum d_i^2 = ۳۸$

$$r_s = 1 - \frac{\sum_{i=1}^k d_i^2}{n(n^2 - 1)} = 1 - \frac{۶ \times ۳۸}{۵(۲۵ - ۱)} = -۰/۹$$

آزمون کولموگروف اسمیرنوف (ks):

این آزمون برای همگونی یا نیکویی برازش بین دو فراوانی تجمعی نسبی در آمار ناپارامتری به کار می‌رود. برای انجام این آزمون فراوانی تجمعی نسبی مشاهده شده (F_0) و فراوانی تجمعی نسبی نظر

(F_e) را بدست می‌آوریم سپس اختلاف آنها را محاسبه می‌کنیم که آماره آزمون $D_n = \max |F_o - F_e|$ می‌باشد و اگر $D_n > D_{\alpha,n}$ باشد. H_0 رد می‌شود و از آن توزیع تبعیت نمی‌کند.

حدود تلرانس طبیعی آزاد - توزیع:

در صورتی که توزیع جامعه مشخص نباشد برای حدود تلرانس دوطرفه، تعداد مشاهدات لازم به طوری که با احتمال $1 - \alpha$ درصد از توزیع بین کوچکتر بین مشاهده و بزرگترین مشاهده قرار گیرد و

$$n \approx \left(\frac{2-\alpha}{\alpha} \right) \left(\frac{X_{1-\gamma,4}^2}{4} \right) + \frac{1}{\gamma} \quad \text{تقریباً از رابطه}$$

و برای حدود تلرانس یک طرفه تعداد مشاهدات لازم به طوری که با احتمال γ ، $1 - \alpha$ درصد از توزیع بین کوچکترین مشاهده و بزرگترین مشاهده قرار گیرد. تقریباً از رابطه $n = \frac{\log(1-\gamma)}{\log(1-\alpha)}$ بدست می‌آید.

مثال: تعداد مشاهدات لازم به منظور حدود تلرانس آزاد — توزیع دو طرفه که با احتمال ۹۰٪ دست کم ۹۵٪ درصد توزیع را در برگیرد را بدست آورید.

$$1 - \alpha = 0.95 \Rightarrow \alpha = 0.05, \quad \gamma = 0.90 \quad X_{0.1,4}^2 = 7.77944$$

$$n \approx \left(\frac{2-\alpha}{\alpha} \right) \left(\frac{X_{1-\gamma,4}^2}{4} \right) + \frac{1}{\gamma} = \left(\frac{2-0.05}{0.05} \right) \left(\frac{X_{0.1,4}^2}{4} \right) + \frac{1}{0.9} = 76$$

مثال: تعداد مشاهدات لازم به منظور حدود تلرانس آزاد — توزیع یک طرفه که با احتمال ۹۸٪ دست کم ۹۰٪ درصد توزیع را در برگیرد را بدست آورید.

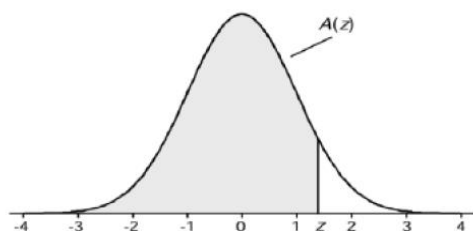
$$1 - \alpha = 0.9 \Rightarrow \alpha = 0.1, \quad \gamma = 0.98$$

$$n = \frac{\log(1-\gamma)}{\log(1-\alpha)} = \frac{\log(0.02)}{\log(0.1)} = \frac{-1.69}{-0.4} \approx 42$$

جدول توزیع نرمال استاندارد

TABLE A.1

Cumulative Standardized Normal Distribution



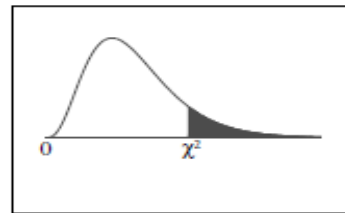
$A(z)$ is the integral of the standardized normal distribution from $-\infty$ to z (in other words, the area under the curve to the left of z). It gives the probability of a normal random variable not being more than z standard deviations above its mean. Values of z of particular importance:

z	$A(z)$	
1.645	0.9500	Lower limit of right 5% tail
1.960	0.9750	Lower limit of right 2.5% tail
2.326	0.9900	Lower limit of right 1% tail
2.576	0.9950	Lower limit of right 0.5% tail
3.090	0.9990	Lower limit of right 0.1% tail
3.291	0.9995	Lower limit of right 0.05% tail

z	0.00	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
0.0	0.5000	0.5040	0.5080	0.5120	0.5160	0.5199	0.5239	0.5279	0.5319	0.5359
0.1	0.5398	0.5438	0.5478	0.5517	0.5557	0.5596	0.5636	0.5675	0.5714	0.5753
0.2	0.5793	0.5832	0.5871	0.5910	0.5948	0.5987	0.6026	0.6064	0.6103	0.6141
0.3	0.6179	0.6217	0.6255	0.6293	0.6331	0.6368	0.6406	0.6443	0.6480	0.6517
0.4	0.6554	0.6591	0.6628	0.6664	0.6700	0.6736	0.6772	0.6808	0.6844	0.6879
0.5	0.6915	0.6950	0.6985	0.7019	0.7054	0.7088	0.7123	0.7157	0.7190	0.7224
0.6	0.7257	0.7291	0.7324	0.7357	0.7389	0.7422	0.7454	0.7486	0.7517	0.7549
0.7	0.7580	0.7611	0.7642	0.7673	0.7704	0.7734	0.7764	0.7794	0.7823	0.7852
0.8	0.7881	0.7910	0.7939	0.7967	0.7995	0.8023	0.8051	0.8078	0.8106	0.8133
0.9	0.8159	0.8186	0.8212	0.8238	0.8264	0.8289	0.8315	0.8340	0.8365	0.8389
1.0	0.8413	0.8438	0.8461	0.8485	0.8508	0.8531	0.8554	0.8577	0.8599	0.8621
1.1	0.8643	0.8665	0.8686	0.8708	0.8729	0.8749	0.8770	0.8790	0.8810	0.8830
1.2	0.8849	0.8869	0.8888	0.8907	0.8925	0.8944	0.8962	0.8980	0.8997	0.9015
1.3	0.9032	0.9049	0.9066	0.9082	0.9099	0.9115	0.9131	0.9147	0.9162	0.9177
1.4	0.9192	0.9207	0.9222	0.9236	0.9251	0.9265	0.9279	0.9292	0.9306	0.9319
1.5	0.9332	0.9345	0.9357	0.9370	0.9382	0.9394	0.9406	0.9418	0.9429	0.9441
1.6	0.9452	0.9463	0.9474	0.9484	0.9495	0.9505	0.9515	0.9525	0.9535	0.9545
1.7	0.9554	0.9564	0.9573	0.9582	0.9591	0.9599	0.9608	0.9616	0.9625	0.9633
1.8	0.9641	0.9649	0.9656	0.9664	0.9671	0.9678	0.9686	0.9693	0.9699	0.9706
1.9	0.9713	0.9719	0.9726	0.9732	0.9738	0.9744	0.9750	0.9756	0.9761	0.9767
2.0	0.9772	0.9778	0.9783	0.9788	0.9793	0.9798	0.9803	0.9808	0.9812	0.9817
2.1	0.9821	0.9826	0.9830	0.9834	0.9838	0.9842	0.9846	0.9850	0.9854	0.9857
2.2	0.9861	0.9864	0.9868	0.9871	0.9875	0.9878	0.9881	0.9884	0.9887	0.9890
2.3	0.9893	0.9896	0.9898	0.9901	0.9904	0.9906	0.9909	0.9911	0.9913	0.9916
2.4	0.9918	0.9920	0.9922	0.9925	0.9927	0.9929	0.9931	0.9932	0.9934	0.9936
2.5	0.9938	0.9940	0.9941	0.9943	0.9945	0.9946	0.9948	0.9949	0.9951	0.9952
2.6	0.9953	0.9955	0.9956	0.9957	0.9959	0.9960	0.9961	0.9962	0.9963	0.9964
2.7	0.9965	0.9966	0.9967	0.9968	0.9969	0.9970	0.9971	0.9972	0.9973	0.9974
2.8	0.9974	0.9975	0.9976	0.9977	0.9977	0.9978	0.9979	0.9979	0.9980	0.9981
2.9	0.9981	0.9982	0.9982	0.9983	0.9984	0.9984	0.9985	0.9985	0.9986	0.9986
3.0	0.9987	0.9987	0.9987	0.9988	0.9988	0.9989	0.9989	0.9989	0.9990	0.9990
3.1	0.9990	0.9991	0.9991	0.9991	0.9992	0.9992	0.9992	0.9992	0.9993	0.9993
3.2	0.9993	0.9993	0.9994	0.9994	0.9994	0.9994	0.9994	0.9995	0.9995	0.9995
3.3	0.9995	0.9995	0.9995	0.9996	0.9996	0.9996	0.9996	0.9996	0.9996	0.9997
3.4	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9998
3.5	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998
3.6	0.9998	0.9998	0.9999							

cum. prob one-tail two-tails	$t_{.50}$	$t_{.75}$	$t_{.80}$	$t_{.85}$	$t_{.90}$	$t_{.95}$	$t_{.975}$	$t_{.99}$	$t_{.995}$	$t_{.999}$	$t_{.9995}$
	0.50	0.25	0.20	0.15	0.10	0.05	0.025	0.01	0.005	0.001	0.0005
	1.00	0.50	0.40	0.30	0.20	0.10	0.05	0.02	0.01	0.002	0.001
df											
1	0.000	1.000	1.376	1.963	3.078	6.314	12.71	31.82	63.66	318.31	636.62
2	0.000	0.816	1.061	1.386	1.886	2.920	4.303	6.965	9.925	22.327	31.599
3	0.000	0.765	0.978	1.250	1.638	2.353	3.182	4.541	5.841	10.215	12.924
4	0.000	0.741	0.941	1.190	1.533	2.132	2.776	3.747	4.604	7.173	8.610
5	0.000	0.727	0.920	1.156	1.476	2.015	2.571	3.365	4.032	5.893	6.869
6	0.000	0.718	0.906	1.134	1.440	1.943	2.447	3.143	3.707	5.208	5.959
7	0.000	0.711	0.896	1.119	1.415	1.895	2.365	2.998	3.499	4.785	5.408
8	0.000	0.706	0.889	1.108	1.397	1.860	2.306	2.896	3.355	4.501	5.041
9	0.000	0.703	0.883	1.100	1.383	1.833	2.262	2.821	3.250	4.297	4.781
10	0.000	0.700	0.879	1.093	1.372	1.812	2.228	2.764	3.169	4.144	4.587
11	0.000	0.697	0.876	1.088	1.363	1.796	2.201	2.718	3.106	4.025	4.437
12	0.000	0.695	0.873	1.083	1.356	1.782	2.179	2.681	3.055	3.930	4.318
13	0.000	0.694	0.870	1.079	1.350	1.771	2.160	2.650	3.012	3.852	4.221
14	0.000	0.692	0.868	1.076	1.345	1.761	2.145	2.624	2.977	3.787	4.140
15	0.000	0.691	0.866	1.074	1.341	1.753	2.131	2.602	2.947	3.733	4.073
16	0.000	0.690	0.865	1.071	1.337	1.746	2.120	2.583	2.921	3.686	4.015
17	0.000	0.689	0.863	1.069	1.333	1.740	2.110	2.567	2.898	3.646	3.965
18	0.000	0.688	0.862	1.067	1.330	1.734	2.101	2.552	2.878	3.610	3.922
19	0.000	0.688	0.861	1.066	1.328	1.729	2.093	2.539	2.861	3.579	3.883
20	0.000	0.687	0.860	1.064	1.325	1.725	2.086	2.528	2.845	3.552	3.850
21	0.000	0.686	0.859	1.063	1.323	1.721	2.080	2.518	2.831	3.527	3.819
22	0.000	0.686	0.858	1.061	1.321	1.717	2.074	2.508	2.819	3.505	3.792
23	0.000	0.685	0.858	1.060	1.319	1.714	2.069	2.500	2.807	3.485	3.768
24	0.000	0.685	0.857	1.059	1.318	1.711	2.064	2.492	2.797	3.467	3.745
25	0.000	0.684	0.856	1.058	1.316	1.708	2.060	2.485	2.787	3.450	3.725
26	0.000	0.684	0.856	1.058	1.315	1.706	2.056	2.479	2.779	3.435	3.707
27	0.000	0.684	0.855	1.057	1.314	1.703	2.052	2.473	2.771	3.421	3.690
28	0.000	0.683	0.855	1.056	1.313	1.701	2.048	2.467	2.763	3.408	3.674
29	0.000	0.683	0.854	1.055	1.311	1.699	2.045	2.462	2.756	3.396	3.659
30	0.000	0.683	0.854	1.055	1.310	1.697	2.042	2.457	2.750	3.385	3.646
40	0.000	0.681</									

Chi-Square Distribution Table



The shaded area is equal to α for $\chi^2 = \chi^2_{\alpha}$.

df	$\chi^2_{.995}$	$\chi^2_{.990}$	$\chi^2_{.975}$	$\chi^2_{.950}$	$\chi^2_{.900}$	$\chi^2_{.100}$	$\chi^2_{.050}$	$\chi^2_{.025}$	$\chi^2_{.010}$	$\chi^2_{.005}$
1	0.000	0.000	0.001	0.004	0.016	2.706	3.841	5.024	6.635	7.879
2	0.010	0.020	0.051	0.103	0.211	4.605	5.991	7.378	9.210	10.597
3	0.072	0.115	0.216	0.352	0.584	6.251	7.815	9.348	11.345	12.838
4	0.207	0.297	0.484	0.711	1.064	7.779	9.488	11.143	13.277	14.860
5	0.412	0.554	0.831	1.145	1.610	9.236	11.070	12.833	15.086	16.750
6	0.676	0.872	1.237	1.635	2.204	10.645	12.592	14.449	16.812	18.548
7	0.989	1.239	1.690	2.167	2.833	12.017	14.067	16.013	18.475	20.278
8	1.344	1.646	2.180	2.733	3.490	13.362	15.507	17.535	20.090	21.955
9	1.735	2.088	2.700	3.325	4.168	14.684	16.919	19.023	21.666	23.589
10	2.156	2.558	3.247	3.940	4.865	15.987	18.307	20.483	23.209	25.188
11	2.603	3.053	3.816	4.575	5.578	17.275	19.675	21.920	24.725	26.757
12	3.074	3.571	4.404	5.226	6.304	18.549	21.026	23.337	26.217	28.300
13	3.565	4.107	5.009	5.892	7.042	19.812	22.362	24.736	27.688	29.819
14	4.075	4.660	5.629	6.571	7.790	21.064	23.685	26.119	29.141	31.319
15	4.601	5.229	6.262	7.261	8.547	22.307	24.996	27.488	30.578	32.801
16	5.142	5.812	6.908	7.962	9.312	23.542	26.296	28.845	32.000	34.267
17	5.697	6.408	7.564	8.672	10.085	24.769	27.587	30.191	33.409	35.718
18	6.265	7.015	8.231	9.390	10.865	25.989	28.869	31.526	34.805	37.156
19	6.844	7.633	8.907	10.117	11.651	27.204	30.144	32.852	36.191	38.582
20	7.434	8.260	9.591	10.851	12.443	28.412	31.410	34.170	37.566	39.997
21	8.034	8.897	10.283	11.591	13.240	29.615	32.671	35.479	38.932	41.401
22	8.643	9.542	10.982	12.338	14.041	30.813	33.924	36.781	40.289	42.796
23	9.260	10.196	11.689	13.091	14.848	32.007	35.172	38.076	41.638	44.181
24	9.886	10.856	12.401	13.848	15.659	33.196	36.415	39.364	42.980	45.559
25	10.520	11.524	13.120	14.611	16.473	34.382	37.652	40.646	44.314	46.928
26	11.160	12.198	13.844	15.379	17.292	35.563	38.885	41.923	45.642	48.290
27	11.808	12.879	14.573	16.151	18.114	36.741	40.113	43.195	46.963	49.645
28	12.461	13.565	15.308	16.928	18.939	37.916	41.337	44.461	48.278	50.993
29	13.121	14.256	16.047	17.708	19.768	39.087	42.557	45.722	49.588	52.336
30	13.787	14.953	16.791	18.493	20.599	40.256	43.773	46.979	50.892	53.672
40	20.707	22.164	24.433	26.509	29.051	51.805	55.758	59.342	63.691	66.766
50	27.991	29.707	32.357	34.764	37.689	63.167	67.505	71.420	76.154	79.490
60	35.534	37.485	40.482	43.188	46.459	74.397	79.082	83.298	88.379	91.952
70	43.275	45.442	48.758	51.739	55.329	85.527	90.531	95.023	100.425	104.215
80	51.172	53.540	57.153	60.391	64.278	96.578	101.879	106.629	112.329	116.321
90	59.196	61.754	65.647	69.126	73.291	107.565	113.145	118.136	124.116	128.299
100	67.328	70.065	74.222	77.929	82.358	118.498	124.342	129.561	135.807	140.169

TABLE A.3

F Distribution: Critical Values of *F* (5% significance level)

ν_1	1	2	3	4	5	6	7	8	9	10	12	14	16	18	20
ν_2															
1	161.45	199.50	215.71	224.58	230.16	233.99	236.77	238.88	240.54	241.88	243.91	245.36	246.46	247.32	248.01
2	18.51	19.00	19.16	19.25	19.30	19.33	19.35	19.37	19.38	19.40	19.41	19.42	19.43	19.44	19.45
3	10.13	9.55	9.28	9.12	9.01	8.94	8.89	8.85	8.81	8.79	8.74	8.71	8.69	8.67	8.66
4	7.71	6.94	6.59	6.39	6.26	6.16	6.09	6.04	6.00	5.96	5.91	5.87	5.84	5.82	5.80
5	6.61	5.79	5.41	5.19	5.05	4.95	4.88	4.82	4.77	4.74	4.68	4.64	4.60	4.58	4.56
6	5.99	5.14	4.76	4.53	4.39	4.28	4.21	4.15	4.10	4.06	4.00	3.96	3.92	3.90	3.87
7	5.59	4.74	4.35	4.12	3.97	3.87	3.79	3.73	3.68	3.64	3.57	3.53	3.49	3.47	3.44
8	5.32	4.46	4.07	3.84	3.69	3.58	3.50	3.44	3.39	3.35	3.28	3.24	3.20	3.17	3.15
9	5.12	4.26	3.86	3.63	3.48	3.37	3.29	3.23	3.18	3.14	3.07	3.03	2.99	2.96	2.94
10	4.96	4.10	3.71	3.48	3.33	3.22	3.14	3.07	3.02	2.98	2.91	2.86	2.83	2.80	2.77
11	4.84	3.98	3.59	3.36	3.20	3.09	3.01	2.95	2.90	2.85	2.79	2.74	2.70	2.67	2.65
12	4.75	3.89	3.49	3.26	3.11	3.00	2.91	2.85	2.80	2.75	2.69	2.64	2.60	2.57	2.54
13	4.67	3.81	3.41	3.18	3.03	2.92	2.83	2.77	2.71	2.67	2.60	2.55	2.51	2.48	2.46
14	4.60	3.74	3.34	3.11	2.96	2.85	2.76	2.70	2.65	2.60	2.53	2.48	2.44	2.41	2.39
15	4.54	3.68	3.29	3.06	2.90	2.79	2.71	2.64	2.59	2.54	2.48	2.42	2.38	2.35	2.33
16	4.49	3.63	3.24	3.01	2.85	2.74	2.66	2.59	2.54	2.49	2.42	2.37	2.33	2.30	2.28
17	4.45	3.59	3.20	2.96	2.81	2.70	2.61	2.55	2.49	2.45	2.38	2.33	2.29	2.26	2.23
18	4.41	3.55	3.16	2.93	2.77	2.66	2.58	2.51	2.46	2.41	2.34	2.29	2.25	2.22	2.19
19	4.38	3.52	3.13	2.90	2.74	2.63	2.54	2.48	2.42	2.38	2.31	2.26	2.21	2.18	2.16
20	4.35	3.49	3.10	2.87	2.71	2.60	2.51	2.45	2.39	2.35	2.28	2.22	2.18	2.15	2.12
21	4.32	3.47	3.07	2.84	2.68	2.57	2.49	2.42	2.37	2.32	2.25	2.20	2.16	2.12	2.10
22	4.30	3.44	3.05	2.82	2.66	2.55	2.46	2.40	2.34	2.30	2.23	2.17	2.13	2.10	2.07
23	4.28	3.42	3.03	2.80	2.64	2.53	2.44	2.37	2.32	2.27	2.20	2.15	2.11	2.08	2.05
24	4.26	3.40	3.01	2.78	2.62	2.51	2.42	2.36	2.30	2.25	2.18	2.13	2.09	2.05	2.03
25	4.24	3.39	2.99	2.76	2.60	2.49	2.40	2.34	2.28	2.24	2.16	2.11	2.07	2.04	2.01
26	4.22	3.37	2.98	2.74	2.59	2.47	2.39	2.32	2.27	2.22	2.15	2.09	2.05	2.02	1.99
27	4.21	3.35	2.96	2.73	2.57	2.46	2.37	2.31	2.25	2.20	2.13	2.08	2.04	2.00	1.97
28	4.20	3.34	2.95	2.71	2.56	2.45	2.36	2.29	2.24	2.19	2.12	2.06	2.02	1.99	1.96
29	4.18	3.33	2.93	2.70	2.55	2.43	2.35	2.28	2.22	2.18	2.10	2.05	2.01	1.97	1.94
30	4.17	3.32	2.92	2.69	2.53	2.42	2.33	2.27	2.21	2.16	2.09	2.04	1.99	1.96	1.93
35	4.12	3.27	2.87	2.64	2.49	2.37	2.29	2.22	2.16	2.11	2.04	1.99	1.94	1.91	1.88
40	4.08	3.23	2.84	2.61	2.45	2.34	2.25	2.18	2.12	2.08	2.00	1.95	1.90	1.87	1.84
50	4.03	3.18	2.79	2.56	2.40	2.29	2.20	2.13	2.07	2.03	1.95	1.89	1.85	1.81	1.78
60	4.00	3.15	2.76	2.53	2.37	2.25	2.17	2.10	2.04	1.99	1.92	1.86	1.82	1.78	1.75
70	3.98	3.13	2.74	2.50	2.35	2.23	2.14	2.07	2.02	1.97	1.89	1.84	1.79	1.75	1.72
80	3.96	3.11	2.72	2.49	2.33	2.21	2.13	2.06	2.00	1.95	1.88	1.82	1.77	1.73	1.70
90	3.95	3.10	2.71	2.47	2.32	2.20	2.11	2.04	1.99	1.94	1.86	1.80	1.76	1.72	1.69
100	3.94	3.09	2.70	2.46	2.31	2.19	2.10	2.03	1.97	1.93	1.85	1.79	1.75	1.71	1.68
120	3.92	3.07	2.68	2.45	2.29	2.18	2.09	2.02	1.96	1.91	1.83	1.78	1.73	1.69	1.66
150	3.90	3.06	2.66	2.43	2.27	2.16	2.07	2.00	1.94	1.89	1.82	1.76	1.71	1.67	1.64
200	3.89	3.04	2.65	2.42	2.26	2.14	2.06	1.98	1.93	1.88	1.80	1.74	1.69	1.66	1.62
250	3.88	3.03	2.64	2.41	2.25	2.13	2.05	1.98	1.92	1.87	1.79	1.73	1.68	1.65	1.61
300	3.87	3.03	2.63	2.40	2.24	2.13	2.04	1.97	1.91	1.86	1.78	1.72	1.68	1.64	1.61
400	3.86	3.02	2.63	2.39	2.24	2.12	2.03	1.96	1.90	1.85	1.78	1.72	1.67	1.63	1.60
500	3.86	3.01	2.62	2.39	2.23	2.12	2.03	1.96	1.90	1.85	1.77	1.71	1.66	1.62	1.59
600	3.86	3.01	2.62	2.39	2.23	2.11	2.02	1.95	1.90	1.85	1.77	1.71	1.66	1.62	1.59
750	3.85	3.01	2.62	2.38	2.23	2.11	2.02	1.95	1.89	1.84	1.77	1.70	1.66	1.62	1.58
1000	3.85	3.00	2.61	2.38	2.22	2.11	2.02	1.95	1.89	1.84	1.76	1.70	1.65	1.61	1.58

منابع

آمار و کاربرد آن در مدیریت جلد اول	دکتر عادل آذر	انتشارات سمت
آمار و کاربرد آن در مدیریت جلد دوم	دکتر عادل آذر	انتشارات سمت
آمار و کاربرد آن در مدیریت جلد اول	دکتر خدیجه جمشیدی	دانشگاه پیام نور
آمار و کاربرد آن در مدیریت جلد دوم	دکتر خدیجه جمشیدی	دانشگاه پیام نور
آمار و احتمال	شلدون راس	دانشگاه فردوسی مشهد
احتمال و کاربرد آن	دکتر علی عمیدی	دانشگاه پیام نور
آمار مهندسی لیبرمن	دکتر هاشم محلوچی	مرکز نشر دانشگاهی