

شاخص های مرکزی

۱. مد یا نما M_o

مد شاخصی مبنی بر تکرار بوده و به طبقه ای اطلاق می شود که دارای بیشترین فراوانی باشد. این شاخص ضعیف ترین و بی ثبات ترین شاخص مرکزی به شمار می رود چرا که حتی ویژگی ترتیب برای آن مهم نبوده (مقیاس اسمی) و لزوماً در مرکز توزیع قرار نمی گیرد.

۲. میانه M_d

میانه نقطه وسط یا پنجاه درصدی داده ها است.

این شاخص مبتنی بر ترتیب داده و دارای مقیاس ترتیبی است؛ پس برخلاف میانگین شاخصی کیفی بوده که تحت تاثیر ارزش عددی یا حجم داده ها قرار نمی گیرد. لذا میانه به داده های پرت کوچک یا بزرگ حساس نیست؛ بنابراین می توان گفت اگر توزیع داده ها چولگی یا داده پرت داشته باشد، بهترین و مناسب ترین شاخص مرکزی میانه است.

۳. میانگین M یا $\bar{x}$

ویژگی ها

۱. میانگین دارای مقیاس فاصله ای بوده و در مقیاس های فاصله ای و نسبی می توان از آن استفاده کرد.

۲. اگر توزیع داده ها نرمال باشد بهترین و مناسب ترین شاخص مرکزی میانگین است.

۳. همواره و در هر توزیعی مجموع انحراف داده ها از میانگین برابر صفر است به همین دلیل است که میانگین را مرکز ثقل یا نقطه تعادل توزیع می نامند:

$$\sum(x - \bar{x}) = 0$$

شاخص های پراکندگی (فاصله ای)

۱. دامنه تغییرات $R = \max - \min + 1$

نه تنها دامنه تغییرات بلکه تمامی شاخص های پراکندگی به دنبال فاصله میان داده ها بوده و دارای مقیاس فاصله ای هستند. دامنه تغییرات ضعیف ترین و بی ثبات ترین شاخص پراکندگی به شمار می رود چرا که در محاسبه ی آن فقط از دو عدد ابتدایی و انتهایی استفاده شده و با تغییر یکی از این اعداد مقدار R نیز شدیداً تغییر می کند (شدیدا حساس به داده پرت).

۲. انحراف چارک $Q = \frac{Q_3 - Q_1}{2}$ (چارک متوسط)

انحراف چارکی نشان می دهد داده ها بطور متوسط چقدر از میانه یا مرکز توزیع انحراف دارند. در واقع انحراف چارکی میزان پراکندگی را در ۵۰٪ وسط توزیع بررسی کرده لذا به داده های پرت کوچک یا بزرگ حساس نیست. پس می توان گفت زمانی که توزیع داده ها چولگی قابل ملاحظه ای داشته باشد، بهترین و مناسب ترین شاخص پراکندگی انحراف چارکی خواهد بود.

۳. انحراف متوسط MD یا $AD = \frac{\sum|x - \bar{x}|}{n}$

انحراف متوسط عبارت است از متوسط قدر مطلق انحراف داده ها از میانگین.

این شاخص نشان می دهد که داده ها به طور متوسط چقدر از میانگین انحراف دارند.

۴. واریانس $S^2 = \frac{\sum(x - \bar{x})^2}{n}$

واریانس عبارت است از متوسط مجذور انحراف داده ها از میانگین.

۵. انحراف استاندارد (معیار) $s = \sqrt{s^2}$

انحراف معیار برابر است با جذر یا رادیکال واریانس.

زمانی که توزیع داده ها نرمال باشد یعنی داده پرت یا چولگی شدیدی نداشته باشیم، بهترین و مناسب ترین شاخص پراکندگی انحراف استاندارد خواهد بود.

$$\text{ضریب تغییرات (ضریب پراکندگی)} = \frac{S}{\bar{X}} \text{ یا } c.v$$

برای مقایسه پراکندگی یک ویژگی در گروه های مختلف، می توان از S^2 یا S استفاده کرد؛ اما برای مقایسه پراکندگی دو یا چند ویژگی مختلف با واحدهای متفاوت می بایست از شاخصی به نام ضریب تغییرات که مستقل از واحد اندازه گیری است استفاده کرد.

تاثیر اعمال ریاضی بر شاخص های مرکزی و پراکندگی:

به طور کلی شاخص های مرکزی به تمامی اعمال ریاضی حساس بوده درحالی که شاخص های پراکندگی فقط به ضرب و تقسیم حساس هستند. چنانچه تمامی داده های یک توزیع را در عدد ثابتی ضرب یا تقسیم کنیم، شاخص های پراکندگی نیز در همان عدد ضرب یا تقسیم خواهند شد به جز واریانس که در مجذور آن عدد ضرب یا تقسیم می شود.

چولگی (کجی)

شاخص چولگی تحت تاثیر داده پرت بوده و جهت بررسی تقارن توزیع استفاده می شود. اگر داده پرت بزرگ داشته باشیم (مانند توزیع درآمد تهرانیها)، میانگین به سمت راست رفته در نتیجه چولگی راست یا مثبت خواهیم داشت (در واقع در چولگی مثبت، میانگین بزرگترین شاخص مرکزی است) در چولگی منفی یا چپ، میانگین کوچکترین شاخص مرکزی است چرا که توسط داده پرت کوچک، به سمت چپ کشیده شده است.

کشیدگی (برآمدگی)

شاخص کشیدگی تحت تاثیر مقدار انحراف استاندارد (پراکندگی) بوده و جهت بررسی تراکم توزیع استفاده می شود. مقدار کشیدگی با مقدار پراکندگی داده ها رابطه عکس دارد بطوریکه هر چه پراکندگی بیشتر شود، توزیع داده ها پخ تر خواهد شد.

نمرات استاندارد و منحنی طبیعی

نمرات استاندارد (نمرات معیار، نمرات تراز)

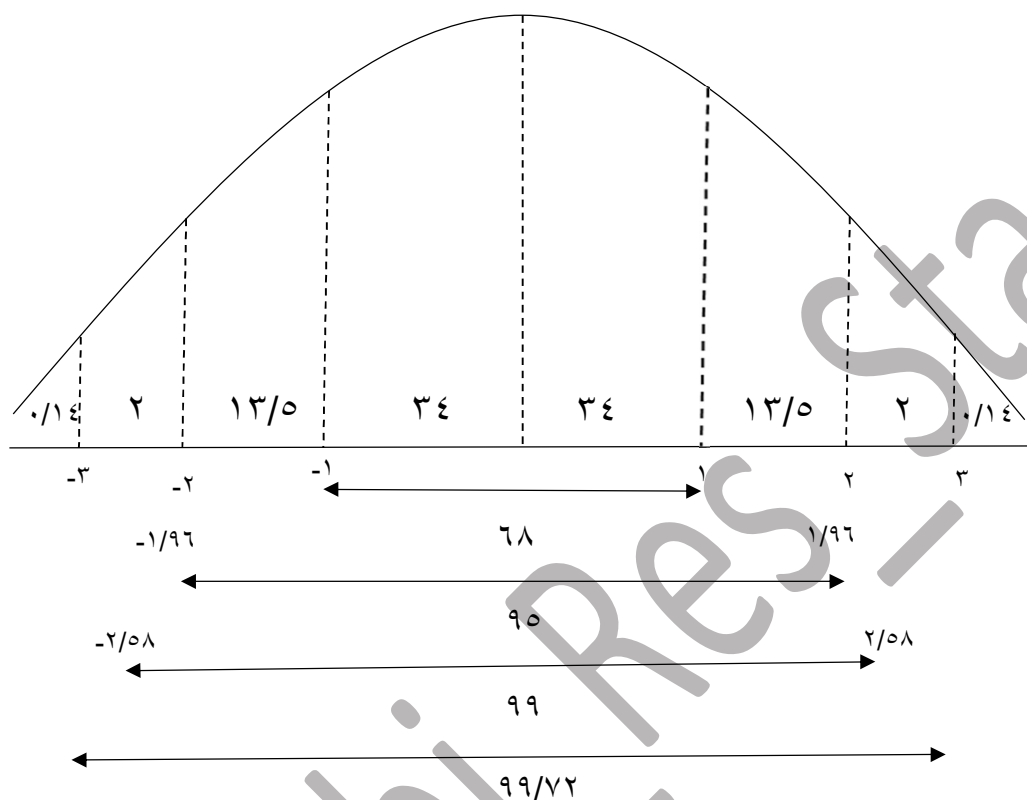
برای تفسیر پذیر و مقایسه پذیرکردن نمرات خام، می توان به کمک نمرات استاندارد، فاصله ی نمرات خام را از میانگین براساس واحد انحراف استاندارد مشخص کرد:

$$Z = \frac{X - \bar{X}}{S}$$

ویژگی های نمره Z

۱. نمرات استاندارد دارای مقیاس فاصله ای هستند.
۲. نمرات استاندارد دارای واحد اندازه گیری به نام انحراف استاندارد هستند.
۳. دامنه نمرات Z نسبت به نمرات خام بسیار محدود تر است تا جایی که در حالت نرمال نمرات Z بین -۳ تا ۳ قرار می گیرند.
۴. همواره و در هر توزیعی، زمانی که نمرات خام به نمرات Z تبدیل می شوند توزیعی به دست می آید که میانگین آن همواره صفر و انحراف معیار آن همواره برابر یک خواهد شد.
۵. تبدیل نمرات خام به نمرات Z یک تبدیل خطی است، بدین معنی که شکل توزیع تغییر نمی کند.
۶. نمرات استاندارد Z دو محدودیت دارند، منفی و اعشاری شدن؛ برای رفع این دو محدودیت است که نمره استاندارد T ساخته شده است.

$$T = 10Z + 50$$



همبستگی و رگرسیون

همبستگی (هم تغییری): توصیفی دو متغیره

ضرایب همبستگی شدت و جهت رابطه را نشان می دهند:

مقدار عددی $r = \pm$: که علامت نشان دهنده جهت و مقدار عددی نشان از شدت رابطه دارد
شدت همبستگی، مقداری بین صفر تا یک را نشان می دهد که با در نظر گرفتن علامت می توان دامنه همبستگی را بین -۱ تا +۱ در نظر گرفت.

ویژگی های همبستگی

- ۱- شدت همبستگی بین دو متغیره، تحت تأثیر هیچ کدام از اعمال ریاضی قرار نمی گیرد؛ اما در مورد جهت همبستگی می توان گفت چنان چه یکی از متغیرها قرینه شود، جهت همبستگی نیز قرینه خواهد شد.
- ۲- ضرایب همبستگی فاقد واحد اندازه گیری بوده و تحت تأثیر اختلاف واحد اندازه گیری نیز قرار نمی گیرند.
- ۳- معناداری همبستگی شدیداً به حجم نمونه یا تعداد آزمودنی ها حساس است؛ به همین دلیل است که حداقل حجم نمونه در پژوهش های همبستگی، $n = 30$ در نظر گرفته می شود.

۴- در یک تحقیق همبستگی، اگر آزمودنی ها همگن یا شبیه هم انتخاب شده باشند، به دلیل کاهش دامنه تغییرات، مقدار هم تغییری یا همبستگی میان متغیرها نیز کاهش می یابد. به عنوان مثال، همبستگی میان هوش و پیشرفت تحصیلی در دانش آموزان تیزهوش، کمتر از همبستگی همین متغیرها در یک نمونه تصادفی از دانش آموزان است.

۵- ضرابی مانند پیرسون، جهت بررسی روابط خطی استفاده می شوند؛ در حالی که متغیرهایی یافت می شود که میانشان رابطه یا همبستگی غیر خطی یا منحنی وجود دارد. برای بررسی همبستگی های غیر خطی، از ضریب همبستگی نسبی یا نسبت همبستگی (تا η) استفاده می شود.

همبستگی های تفکیکی و نیمه تفکیکی

به مطالعه همبستگی میان دو متغیر پس از کنترل یا حذف متغیر سوم، همبستگی تفکیکی (سه می یا جزئی) گفته می شود. در واقع اگر در ارتباط میان ۳ متغیر (مانند X_1 ، X_2 و X_3)، اثر متغیر X_3 را از روی هر دو متغیر حذف کنیم، همبستگی تفکیکی بوده و اگر اثر متغیر X_3 را فقط از روی یکی از متغیرها، کنترل یا حذف کنیم، همبستگی نیمه تفکیکی خواهد بود.

نمایش همبستگی تفکیکی: $r_{12.3}$

نمایش همبستگی نیمه تفکیکی: $r_{1(2.3)}$

ضریب همبستگی پیرسون

$$r_{xy} = \frac{Cov_{xy}}{S_x S_y} \quad \text{و} \quad Cov_{xy} = \frac{\sum (x - \bar{x})(y - \bar{y})}{n}$$

در آمار های کاربردی، به دو دلیل استفاده از پیرسون بر کواریانس ترجیح دارد:

- ۱- پیرسون برخلاف کواریانس، به واحد اندازه گیری حساس نیست.
- ۲- پیرسون برخلاف کواریانس، دامنه ثابتی دارد ($-1 \leq r \leq +1$). به همین دلیل به کواریانس، پیرسون بدون ابعاد نیز گفته می شود.

ضرایب همبستگی	سطح سنجش متغیرها	توضیحات	مثال
پیرسون	هر دو متغیر کمی	قوی ترین ضریب همبستگی	همبستگی بین میزان انگیزش و میزان عزت نفس دانشجویان
اسپیرمن	هر دو متغیر رتبه ای (ترتیبی)	جایگزین پیرسون در شرایط غیرنرمال یا حجم نمونه کمتر از ۳۰ یا روابط غیرخطی	همبستگی بین رتبه درس فیزیک و رتبه درس شیمی دانش آموزان
تاو کندال	هر دو متغیر رتبه ای (ترتیبی)	مناسب برای حجم نمونه کمتر از ۱۰	همانند اسپیرمن
توافقی (C)	هر دو متغیر اسمی	جدول مربعی ۳×۳ یا بیشتر	همبستگی بین نوع شغل و نوع رشته تحصیلی
دورشته ای نقطه ای	یک متغیر کمی و متغیر دیگر دو ارزشی واقعی	دو ارزشی های واقعی مانند جنسیت، ماهیتا کیفی هستند.	همبستگی بین معدل لیسانس و تغییر یا عدم تغییر رشته تحصیلی
دو رشته ای	یک متغیر کمی و متغیر دیگر دو ارزشی ساختگی	دو ارزشی های ساختگی مانند قبول/رد، از روی متغیرهای کمی ساخته می شوند.	همبستگی بین هوش و درآمد بالا و پایین

فی (فای ϕ)	هر دو متغیر دو ارزشی واقعی	جدول مربعی 2×2 واقعی	همبستگی بین جنسیت و شرکت یا عدم شرکت در انتخابات
تتراکوریک	هر دو متغیر دو ارزشی ساختگی	جدول مربعی 2×2 ساختگی	همبستگی بین بلند قدی/کوتاه قدی و باهوشی/کم هوشی
کرامر	هر دو متغیر اسمی	جدول مستطیلی 2×3 یا بیشتر	همبستگی بین جنسیت و انتخاب رشته تحصیلی

پیش بینی

زمانی که بین دو متغیر همبستگی وجود داشته باشد، می توان به کمک معادله رگرسیون و در میان اعضای جامعه، دست به پیش بینی زد؛ به طوری که می توان نمره شخصی را در یک متغیر به کمک نمره همان شخص در متغیر دیگر پیش بینی کرد. دقت این پیش بینی: به شدت همبستگی دو متغیر بستگی دارد، به طوری که: در همبستگی کامل ($r = \pm 1$)، پیش بینی نیز دقیق و کامل خواهد بود. در همبستگی ناقص ($0 < r < \pm 1$)، پیش بینی نیز ناقص خواهد بود (خطای پیش بینی خواهیم داشت). در همبستگی صفر ($r = 0$)، هیچ گونه پیش بینی نمی توان داشت.

ضریب تبیین (ضریب تعیین، مقدار واریانس تبیین شده، واریانس مشترک، مقدار پیش بینی): r^2

اگر بین دو متغیر به میزان r همبستگی وجود داشته باشد، می توان به میزان r^2 ، میان آنها تبیین یا پیش بینی انجام داد. به عنوان مثال: اگر همبستگی دو متغیر 0.7 - باشد، بدین معناست که می توان 49% از تغییرات متغیر ملاک (y) را به کمک متغیر پیش بین (x)، تبیین یا پیش بینی کرد.

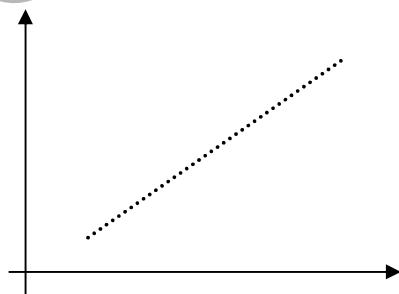
$$r_{xy}^2 = \frac{Cov_{xy}^2}{S_x^2 S_y^2} = \frac{\text{واریانس مشترک}}{\text{واریانس کل}}$$

ضریب عدم تبیین (مقدار واریانس تبیین نشده، مقدار خطا یا عدم پیش بینی): $1 - r_{xy}^2$

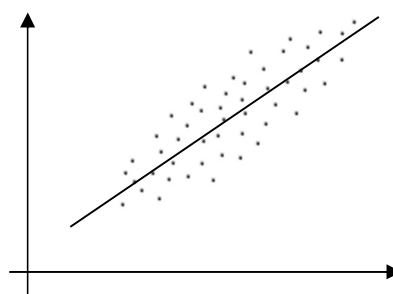
خط رگرسیون

در همبستگی های کامل، به دلیل اینکه تمامی نقاط بر روی یک خط قرار دارند، با داشتن معادله این خط، می توان بدون هیچ گونه خطایی پیش بینی انجام داد.

اما در همبستگی های ناقص، به دلیل پراکنش نقاط، از خطی فرضی به نام خط رگرسیون که از میانگین نقاط عبور کرده، برای پیش بینی استفاده می کنیم. این خط کمک می کند که محقق پیش بینی انجام دهد با حداقل خطا. به همین دلیل، خط رگرسیون را خط حداقل مجذورها و هم چنین خط برازش نیز می نامند.



همبستگی کامل



همبستگی ناقص

معادله خط رگرسیون $y' = bx + a$

a: عرض از مبدأ؛ محل تلاقی خط رگرسیون با محور Y یا عرض ها؛ مقدار پایه یا ثابت؛ مقدار Y به ازای $x = 0$
b: شیب یا ضریب زاویه خط رگرسیون؛ میزان تغییرات Y به ازای یک واحد تغییر در X

نحوه تشکیل معادله رگرسیون:

$$\left. \begin{aligned} b &= r_{xy} \frac{S_y}{S_x} \\ a &= \bar{y} - b\bar{x} \end{aligned} \right\} y' = bx + a$$

پیش بینی نمرات استاندارد

وقتی نمرات X و Y را به نمرات استاندارد تبدیل کنیم، میانگین ها برابر صفر و انحراف معیارها برابر یک خواهد شد؛ لذا خواهیم داشت:

$$\left. \begin{aligned} b &= r \\ a &= 0 \end{aligned} \right\} y = bx + a \quad Z_y = r_{xy} Z_x$$

$$(S_{xy}) \text{ خطای استاندارد پیش بینی} = S_y \sqrt{1 - r^2}$$

برآورد

خطای نمونه گیری:

اگر از جامعه ای بزرگ، نمونه ای انتخاب کرده و ویژگی های آن را محاسبه کنیم، بین ویژگی های جامعه و نمونه، تفاوت وجود خواهد داشت؛ علت این امر وجود نوعی خطا به نام خطای نمونه گیری است. این خطا الزاماً نتیجه اشتباه در نمونه گیری نبوده، بلکه به دلیل وجود پراکندگی و تفاوت میان اعضای جامعه است.

قضیه حد مرکزی

این قضیه بیان می کند که اگر از جامعه ای تعداد زیادی نمونه تصادفی و با حجم برابر انتخاب کرده و میانگین آن ها را محاسبه کنیم، توزیع این میانگین ها تشکیل یک توزیع طبیعی یا نرمال داده؛ میانگین کل این نمونه ها به میانگین واقعی جامعه نزدیک شده و مقدار پراکندگی یا انحراف معیار بین این میانگین ها برابر $\frac{\sigma}{\sqrt{n}}$ خواهد بود.

خطای استاندارد نمونه گیری (خطای استاندارد برآورد، خطای استاندارد میانگین)

به پراکندگی یا انحراف معیار میان میانگین های نمونه های تصادفی انتخاب شده از جامعه، خطای استاندارد برآورد میانگین می گویند. مقدار این خطا با انحراف معیار یا پراکندگی جامعه، رابطه مستقیم و با حجم نمونه رابطه عکس دارد:

$$S_{\bar{x}} = \frac{\sigma}{\sqrt{n}} \approx \frac{S}{\sqrt{n}}$$

$$S_{\bar{x}} = \frac{S}{\sqrt{n}} \rightarrow S_{\bar{x}}^2 = \frac{S^2}{n}$$

واریانس میانگین ها انحراف معیار میانگین ها
واریانس نمونه گیری انحراف معیار نمونه گیری

فاصله اطمینان جامعه : $\mu = \bar{X} \pm ZS_{\bar{x}}$

پارامتر نامرکزیت $ZS =$ → $n = \frac{Z^2 S^2}{d^2} = \left(\frac{ZS}{d}\right)^2$ فرمول محاسبه حجم نمونه

$d = ZS_{\bar{x}}$: نصف طول فاصله اطمینان- میزان خطای حاشیه ای- تفاوت بین میانگین ها- حداکثر خطای قابل قبول- اندازه یا حجم اثر

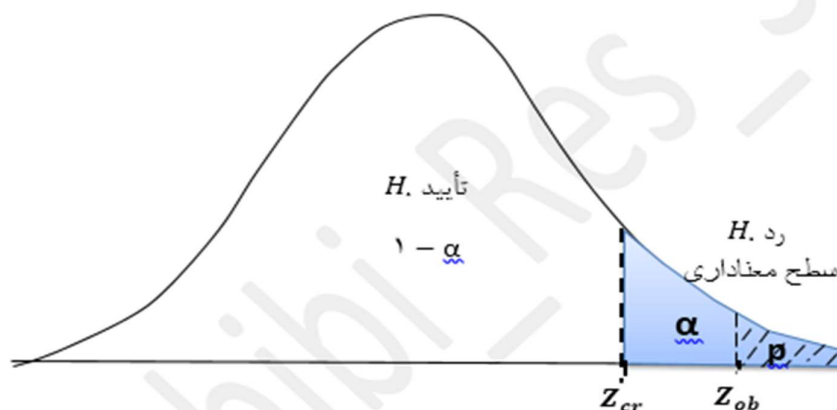
فرضیه	واقعیت	H. صحیح	H. غلط
رد H.	α	$1 - \beta$	β
تأیید H.	$1 - \alpha$	$1 - \beta$	β

آزمون فرضیه

نکته: خطاهای α و β با حجم نمونه رابطه عکس دارند.

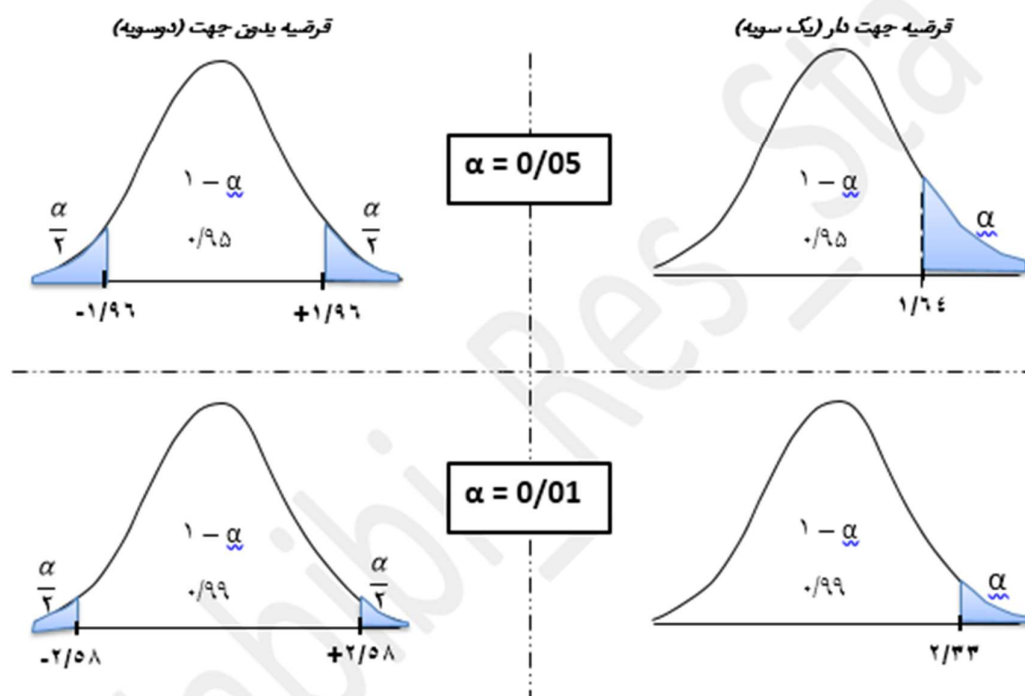
نکته: خطاهای α و β عکس یکدیگرند؛ بدین معنی که این دو خطا همانند دو روی یک سکه بوده و همزمان اتفاق نمی افتند.

مکانیسم تشخیص معناداری (مکانیسم تأیید یا رد فرضیه)



$$|Z_{ob}| \geq |Z_{cr}| \rightarrow p \leq \alpha$$

آزمون های تک دالته و دو دالته



همانطور که $\alpha = 0.05$ شانس معناداری بیشتری از $\alpha = 0.01$ دارد، در آزمون های تک دالته نیز احتمال معناداری با رد H_0 بیشتر از آزمون های دودالته است؛ چرا که با افزایش سطح معناداری، مقدار بحرانی کوچک تر شده و عبور از آن راحت تر می شود (تدازه سطح معناداری (α) با مقدار بحرانی رابطه عکس دارد).

آزمون های آماری پارامتریک

مقایسه رتبه	آزمون U من ویتنی	مقایسه دو گروه مستقل در مقیاس ترتیبی
	آزمون ویلکاکسون	مقایسه دو گروه وابسته در مقیاس ترتیبی
	آزمون کراسکال والیس	مقایسه بیش از دو گروه مستقل در مقیاس ترتیبی
	آزمون فریدمن	مقایسه بیش از دو گروه وابسته در مقیاس ترتیبی
مقایسه فراوانی، نسبت یا درصد	آزمون خی دو	مقایسه دو یا چند گروه مستقل در مقیاس اسمی

پارامتریک	نپارامتریک
t مستقل	U من ویتنی
t وابسته	ویلکاکسون
تحلیل واریانس (F مستقل)	کراسکال والیس
تحلیل واریانس مکرر (F وابسته)	فریدمن

آزمون خی دو تک متغیره (خی دو انطباق)

آزمون خی دو (χ^2 یا کای دو یا کای اسکوتر) فراوانی، درصد یا نسبتهای مشاهده شده را با فراوانی، درصد یا نسبتهای مورد انتظار مقایسه کرده تا مشخص شود که بین آنچه انتظار می رود و آنچه اتفاق افتاده، مطابقت وجود دارد یا خیر. در واقع این آزمون، انطباق میان اصول نظری (E) یعنی داده های بدست آمده از پیشینه یا نظریه ها) و واقعیت (O) یعنی داده های حاصل از آزمودنی ها) را بررسی می کند. به همین دلیل، آزمون خی دو تک متغیره، جزء آزمون های برازش یا نیکویی انطباق به شمار می رود (بررسی انطباق میان مدل حاصل از نظریه با داده های حاصل از آزمودنیها).

آزمون خی دو دومتغیره (خی دو استقلال)

به کمک خی دو دومتغیره می توان رابطه یا همبستگی میان دو متغیر اسمی و گسسته را بررسی کرد؛ بدین ترتیب که اگر خی دو دومتغیره معنادار شود (مانند جدول الف)، نشان از همبستگی یا رابطه میان متغیرها داشته و چنانچه این آزمون عدم معناداری را نشان دهد (مانند جدول ب)، مستقل بودن دو متغیر را نتیجه خواهد داد. به همین دلیل است که آزمون خی دو را در حالت دومتغیره، خی دو استقلال می نامند.

جدول الف

	روانشناسی	علوم تربیتی	جامعه شناسی
زن	۸۷	۶۴	۱۵
مرد	۱۲	۳۵	۷۴

جدول ب

	روانشناسی	علوم تربیتی	جامعه شناسی
زن	۶۴	۶۴	۶۵
مرد	۶۷	۶۲	۶۶

کاربرد درجه آزادی در آمار

بر خلاف توزیع Z که شکل ثابتی دارد، سایر توزیع های آماری مانند t، F و χ^2 ثابت نبوده و با توجه به حجم نمونه یا تعداد گروه ها تغییر می کنند. لذا به غیر از آزمون Z، مقادیر بحرانی سایر آزمون ها ثابت نبوده و در توزیع های نمونه گیری محاسبه و در جداولی تنظیم شده اند. به همین دلیل محقق هنگام استفاده از این آزمون ها لازم است جهت تشخیص معناداری، با توجه به مقادیر α و df، به جدول توزیع مربوطه مراجعه کرده و مقدار بحرانی را استخراج کند.

درجه آزادی خی دو تک متغیره: $df = K - 1$

K: تعداد گروه ها، طبقات یا مقوله ها

$$df = (r-1)(c-1)$$

درجه آزادی خی دو دومتغیره (درجه آزادی جداول توافقی $r \times c$):

r: تعداد سطرها
c: تعداد ستون ها

آزمون های آماری پارامتریک

مقایسه تک گروه (یک نمونه ای) ← آزمون های z یا t تک گروه

مقایسه دو گروه ← مستقل: آزمون های z یا t مستقل
وابسته: آزمون های z یا t وابسته

مقایسه چند گروه ← مستقل: آزمون تحلیل واریانس (F مستقل)
وابسته: آزمون تحلیل واریانس با اندازه گیری مکرر (F وابسته)

تفاوت دو توزیع z و t

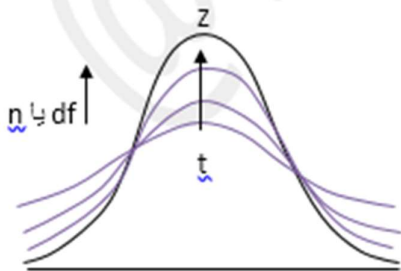
برخلاف توزیع z که شکل ثابتی دارد، شکل توزیع t ثابت نبوده و بستگی به مقدار درجه آزادی دارد؛ به طوری که با افزایش حجم نمونه یا درجه آزادی، S_t به عدد یک نزدیک شده و در نتیجه توزیع t شبیه به z می شود.

طبق قرارداد علم آمار، اگر حجم نمونه کمتر از ۳۰ باشد، از توزیع t استفاده کرده و در غیر اینصورت ($n \geq 30$) از همان توزیع z یا نرمال هم می توان استفاده کرد.

از آنجایی که مقدار انحراف معیار t بزرگتر از انحراف معیار z است، لذا توزیع t پخ تر یا هموارتر (دارای کشیدگی منفی) بوده و به همین دلیل، مقادیر بحرانی آزمون t از مقادیر بحرانی z بزرگتر بوده و می توان گفت آزمون t آزمونی دقیق تر یا سخت گیرانه تر به شمار می رود.

$$\left\{ \begin{array}{l} X_z = 0 \\ S_z = 1 \end{array} \right. \quad \text{توزیع} \quad \left\{ \begin{array}{l} X_t = 0 \\ S_t > 1 \end{array} \right. \quad \text{توزیع} \quad \left\{ \begin{array}{l} X_t = 0 \\ S_t > 1 \end{array} \right. \quad \text{استودنت}$$

$$S_t = \sqrt{\frac{df}{df-2}} \quad \left\{ \begin{array}{l} df=0 : S_t = 1/3 \\ df=1 : S_t = 1/0.1 \end{array} \right.$$



$df = 1$ (حجم نمونه خیلی کم): بیشترین تفاوت بین دو توزیع
 $df = \infty$ (حجم نمونه خیلی زیاد): بیشترین شباهت بین دو توزیع
(در درجه آزادی بینهایت، توزیع t بر z منطبق می شود)

آزمون مقایسه تک گروه

زمانی که بخواهیم میانگین یک نمونه را با میانگین جامعه یا با یک عدد ثابت ادعا شده مقایسه کنیم، از آزمون های z یا t تک گروه استفاده می کنیم:

$$z \text{ یا } t \text{ تک گروه} = \frac{\bar{x} - \mu}{S_{\bar{x}}}$$

و درجه آزادی t تک نمونه ای برابر است با : $df = n - 1$

آزمون مقایسه دو گروه مستقل

زمانی که بخواهیم میانگین دو گروه مستقل را با هم مقایسه کنیم (مانند مقایسه میانگین درآمد مردان و زنان)، از آزمون های z یا t مستقل استفاده می کنیم:

$$z \text{ یا } t \text{ مستقل} = \frac{\bar{X}_1 - \bar{X}_2}{S_{\bar{X}_1, \bar{X}_2}} = \frac{\text{تفاضل دو میانگین}}{\text{خطای استاندارد تفاضل دو میانگین}}$$

$df = n_1 + n_2 - 2$ یا $n - 2$: درجه آزادی t مستقل

آزمون مقایسه دو گروه وابسته

زمانی که قصد مقایسه دو گروه وابسته یا همبسته را در یک متغیر کمی داشته باشیم (مقایسه نمرات میان ترم و پایان ترم دانش آموزان یک کلاس)، از آزمون های Z یا t وابسته استفاده می کنیم:

$$z \text{ یا } t = \frac{\bar{D}}{S_{\bar{D}}} = \frac{\text{میانگین تفاوت نمرات}}{\text{خطای استاندارد میانگین تفاوت نمرات}}$$

و درجه آزادی t وابسته نیز همانند t تک گروه برابر است با: $df = n - 1$

نکته: مقدار خطای استاندارد (خطای برآورد) با توان آزمون رابطه عکس دارد؛ بدین معنی که با افزایش خطا، مقدار آماره آزمون و در نتیجه شانس معناداری کاهش می یابد.

$$\frac{\dots}{\uparrow \text{خطا}} = t \text{ یا } z \downarrow$$

نکته: آزمون های وابسته مانند t وابسته یا F وابسته (اندازه گیری مکرر)، توان بالاتری نسبت به آزمون های مستقل دارند؛ چرا که در مقایسه های وابسته یا درون گروهی نسبت به مقایسه های مستقل یا بین گروهی، تفاوت های فردی و در نتیجه خطا یا پراکندگی کمتر بوده و این امر باعث افزایش توان یا شانس معناداری خواهد شد.

آزمون **خی دو**: مقایسه دو نسبت مستقل - مقایسه چند نسبت مستقل

مقایسه نسبت یا درصد

آزمون **Z**: مقایسه نسبت تک گروه - مقایسه دو نسبت مستقل - مقایسه دو نسبت وابسته

در مقایسه دو نسبت (درصد) مستقل می توان از هر دو آزمون استفاده کرد که در این حالت اولویت با آزمون Z است. ضمناً اگر در مقایسه دو نسبت مستقل از هر دو روش استفاده کنیم، تساوی مقابل برقرار خواهد بود:

$$\chi^2 = Z^2$$

بررسی معناداری همبستگی:

$$t = r \sqrt{\frac{n-2}{1-r^2}}$$

به کمک آزمون t می توان معناداری ضریب همبستگی را بررسی کرد:

درجه آزادی معناداری یک ضریب همبستگی: $df = n - 2$

آزمون مقایسه چند گروه (مقایسه چند میانگین)

بهترین آزمون برای مقایسه بیش از دو گروه، آزمون F است. اگر در این حالت از آزمون t استفاده کنیم، به دلیل تعدد استفاده ها، مقدار خطا افزایش می یابد (تورم آلفا رخ می دهد).

در مقایسه دو گروه نیز، دقیق ترین روش آماری آزمون t است؛ اما اگر در این حالت (مقایسه دو میانگین) از هر دو آزمون F و t استفاده کنیم،

رابطه مقابل برقرار خواهد بود: $F_{ob} = t_{ob}^2$

فرمول محاسبه تحلیل واریانس:

$$F = \frac{\text{میانگین مجذورات بین گروهی}}{\text{میانگین مجذورات درون گروهی}} = \frac{\text{واریانس بین گروهی}}{\text{واریانس درون گروهی}} = \frac{\text{پراکندگی بین گروهی}}{\text{پراکندگی درون گروهی}} = \frac{\text{تفاوت بین گروه ها}}{\text{تفاوت درون گروه ها}} = \frac{\text{Mean Squares of Between}}{\text{Mean Squares of Within}}$$

$$F = \frac{S_b^2}{S_w^2} = \frac{MS_b}{MS_w} = \frac{MS_{tr}}{MS_e} = \frac{\text{واریانس آزمایشی}}{\text{واریانس خطا}} = \frac{\text{اثر متغیر مستقل}}{\text{اثر متغیر مزاحم}} = \frac{\text{درجه آزادی بین گروهی}}{\text{مجموع مجذورات درون گروهی}} = \frac{\text{مجموع مجذورات بین گروهی}}{\text{درجه آزادی درون گروهی}} = \frac{SS_b}{df_b} \div \frac{SS_w}{df_w}$$

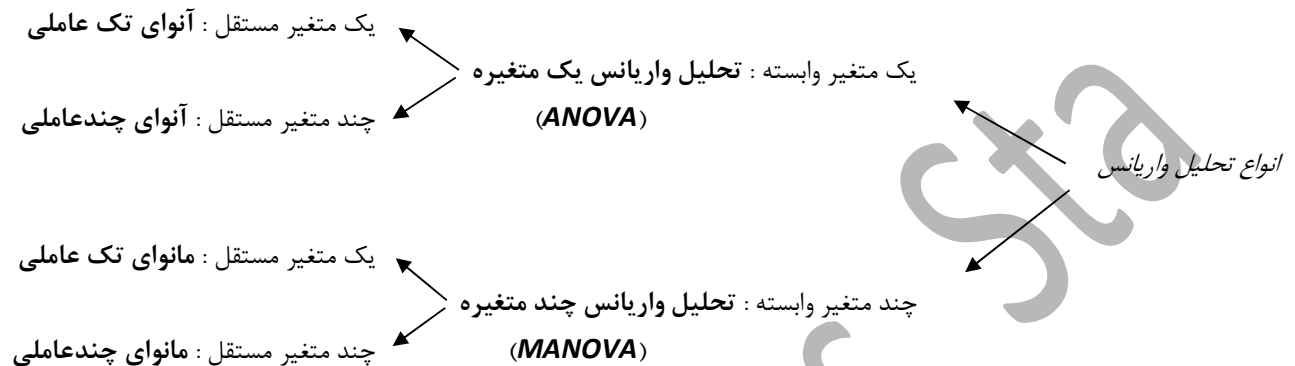
$$df_b = K - 1$$

درجه آزادی بین گروهی یا آزمایشی (صورت):

$$df_w = n_{\text{کل}} - K$$

درجه آزادی درون گروهی یا خطا (مخرج):

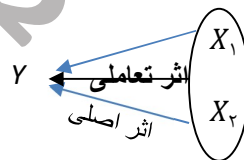
K: تعداد گروه ها یا مقوله ها و n: تعداد کل آزمودنی ها ($n_1 + n_2 + \dots$)



زمانی که بیش از یک متغیر مستقل داریم، تعداد گروه ها برابر با حاصلضرب تعداد سطوح متغیرهای مستقل خواهد بود.

به مطالعه جداگانه تک تک متغیرهای مستقل، اثر اصلی گفته می شود؛ به همین دلیل تعداد اثرهای اصلی با تعداد متغیرهای مستقل برابر است.

به مطالعه همزمان حداقل دو متغیر مستقل، اثر تعاملی یا متقابل گفته می شود. می توان گفت در اثر تعاملی، اثر متغیرهای مستقل از یکدیگر مستقل نبوده و همزمان مورد مطالعه قرار می گیرند:



تعداد اثرهای اصلی و تعاملی در طرح سه عاملی

۳ اثر اصلی (سه متغیر مستقل یا سه عامل)

تعداد اثرهای تعاملی: ۴ اثر تعاملی (۳ اثر تعاملی درجه یک و ۱ اثر تعاملی درجه دو)

تعاملی درجه یک یا مرتبه اول، به تعامل دوتایی متغیرها گفته شده؛ تعاملی درجه دو یعنی تعامل سه تایی و به همین ترتیب

آزمون تحلیل کواریانس (آنکوا - مانکوا)

در روابط علی، محقق سعی در کنترل متغیر مزاحم دارد. یکی از روش هایی که می توان به کمک آن متغیر مزاحم را حذف یا کنترل کرد، روش تحلیل کواریانس است.

روش تحلیل کواریانس، روش ارتقا یافته یا پیشرفته ی تحلیل واریانس است چرا که قابلیت اضافه ی آن نسبت به تحلیل واریانس، کنترل عوامل مزاحم است. در تحلیل کواریانس، متغیر مزاحم (متغیر همپراش، کمکی یا مشتبه کننده) اندازه گیری شده و سپس به کمک رگرسیون، اثر آن از روی نتایج حذف می شود.

در واقع هرگاه در تست ها به کنترل مزاحم یا کنترل پیش آزمون اشاره شد، مناسب ترین روش آماری، تحلیل کواریانس خواهد بود.

تحلیل کواریانس = تحلیل واریانس + رگرسیون خطی

نکته: تحلیل کواریانس علاوه بر مفروضه های تحلیل واریانس (نرمال بودن و یکسانی واریانس ها)، دو مفروضه خاص نیز دارد:

۱- رابطه متغیر مزاحم و متغیر وابسته از نوع خطی باشد.

۲- همگنی شیب رگرسیون: بدین معنا که بین متغیر مزاحم (همپراش) و متغیر مستقل، رابطه یا تعامل معناداری وجود نداشته باشد.

طرح بلوکی

اگر رابطه بین متغیر مزاحم و متغیر وابسته قوی تر از ۰/۶ باشد، بهترین روش کنترل استفاده از تحلیل کواریانس است؛ اما چنانچه این رابطه ضعیف تر از ۰/۴ باشد، برتری با طرح بلوکی خواهد بود (بین ۰/۴ تا ۰/۶ این دو روش برتری به هم ندارند).

در روش بلوکی، متغیر مزاحم به عنوان یک متغیر مستقل (متغیر مقوله ای یا طبقه ای) وارد تحقیق می شود؛ بدین صورت که آزمودنی ها بر اساس متغیر مزاحم، به بلوک ها یا گروه هایی همگن تقسیم شده، سپس در هر کدام از گروه ها رابطه اصلی (رابطه مستقل وابسته) بررسی می شود. در واقع به کمک طرح بلوکی می توان واریانس درون گروهی یا خطا (تفاوت های فردی) را کاهش داده و به افزایش واریانس بین گروهی نیز کمک کرد.

به عنوان مثال اگر در مطالعه اثرگذاری روش تدریس بر یادگیری، متغیر هوش به عنوان مزاحم شناخته شود، می توان آزمودنی ها را بصورت زیر، به گروه هایی همگن تقسیم کرد:

$$\begin{array}{l} IQ < 90 \\ 90 < IQ < 110 \\ IQ > 110 \end{array} \quad \left. \begin{array}{l} \\ \\ \end{array} \right\} \text{گروه بندی آزمودنی ها بر اساس هوش}$$

اندازه (حجم) اثر

آزمون های آماری، دو محدودیت اساسی دارند:

۱- وابستگی شدید به حجم نمونه ۲- عدم محاسبه میزان اثرگذاری متغیرها

برای برطرف کردن این دو محدودیت، لازم است از شاخصی تکمیلی به عنوان اندازه اثر استفاده کنیم؛ این شاخص قادر است مستقل از حجم نمونه، میزان اثر گذاری متغیر مستقل را محاسبه کند.

روش های مختلفی برای محاسبه اندازه اثر وجود دارد که رایج ترین آن، محاسبه مجذور اتا است:

$$\eta^2 = \frac{SS_b}{SS_{total}}$$

مجموع مجذورات کل، حاصل جمع مجموع مجذورات بین و درون گروهی است: $SS_{total} = SS_b + SS_w$

اگر مقدار مجذور اتا بیشتر از ۰/۵۰ باشد، نشان می دهد مقدار مجموع مجذورات بین گروهی بیشتر از مجموع مجذورات درون گروهی است (متغیر مستقل بیشتر از عوامل دیگر یا مزاحم اثرگذار بوده) و اگر اندازه اثر کمتر از ۰/۵۰ باشد، نشان از اثرگذاری زیاد متغیرهای مزاحم دارد.