

آمار و احتمالات (استاد فیلین خوشنود و مهریاش) (شروع خودی بود)

برگردد

۱- مقدماتی بر آمار و احتمالات

۲- احتمال

۳- آمار توصیفی

۴- آمار استنباطی

۵- مدل‌های احتمالی و توزیع‌های آماری

۶- رگرسیون و همبستگی

۷- آمار چند متغیره

۱- آمار چیست؟ آمار علمی است که بدین جهت آورده می‌شود، تجزیه، تحلیل و تفسیر داده‌ها می‌پردازد. آمار این امکان را می‌دهد که از داده‌های موجود به نتایج قابل اعتماد برسیم و الگوهای عمومی را شناسایی کنیم. آمار به دو دسته توصیفی و استنباطی تقسیم می‌شود.

توصیفی

استنباطی

آمار توصیفی داده‌ها به شکلی ساده و مختصر نمایش داده می‌شود، شاخص‌های اصلی آمار توصیفی شامل میانگین، میانه، واریانس و انحراف معیار است.

آمار استنباطی از داده‌های نمونه‌ای برای استنباط و نتیجه‌گیری دارد.



به کمک جمعیت بزرگتر (استگاه) می‌شود، هدف از آمار استنباطی

آن در مورد شرفیه ها و تخمین یا امارت های جمعیت می باشد

احتمالات بر بزرگی پدیده های تصادفی و پیوسته بین نتایج آنها می پردازد

احتمال : مقدار عددی می باشد که نشان دهنده احتمال وقوع یک رویداد است و مقداری بین ۰ و ۱ دارد.

احتمال صفر : رویدادی که غیر ممکن است

احتمال یک : رویدادی که قطعاً رخ می دهد.

احتمال های بین صفر و یک ، احتمال وقوع رویداد هایی است که به صورت تصادفی یا

احتمالی ممکن است اتفاقاً بیفتند.

دایره آمار و احتمال : در حالی که آمار به بررسی داده ها و استنباط اطلاعات

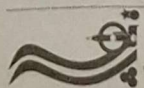
از آنها می پردازد احتمال در واقع به مدل سازی پدیده های تصادفی کمک می کند ، احتمال

بایه و اساس بسیاری از روش های استنباطی در آمار می باشد.

مفاهیم پایه در احتمال : فضای نمونه : مجموعه ای از تمام نتایج ممکن یک آزمایش

تصادفی . رویداد : هر زیر مجموعه ای از فضای نمونه که نتایج خاص وایندی

می کند . احتمال شرطی : احتمال وقوع یک رویداد با فرض وقوع رویداد دیگری



قوانین احتمال : قانون جمع اگر دو رویداد A و B مستقل باشند، احتمال وقوع

$$P(A \cup B) = P(A) + P(B) - P(A \cap B) \quad \text{هر یک از آنها برابر است با}$$

تأویض ضرب : برای رویدادهای مستقل، احتمال وقوع هر دو برابر است با

$$P(A \cap B) = P(A) \times P(B)$$

آغاز تحلیل داده ها : هنگامی که داده ها جمع آوری می شوند، اولین قدم در تحلیل آنها

استفاده از آمار توصیفی می باشد. تا بتوانیم کلی از داده ها بدست آوریم؛ سپس می توان از آمار

استنباطی برای انجام آزمون های فرضیه یا تخمین پارامترهای جمعیت استفاده کرد.

نمونه ای در یک بازی یک ناس شش وجهی انداخته می شود. احتمال اینکه عدد ۴ ظاهر شود

$$A = \{1, 2, 3, 4, 5, 6\} \quad p = \frac{1}{6} \quad \text{چقدر است؟}$$

تقریباً در یک کلاس ۹ درصد از دانش آموزان دختر هستند و ۴ درصد پسر هستند.

از میان دخترها ۷ درصد نمرات خوب دارند، از میان پسرها ۵ درصد نمرات خوب دارند اگر

یک دانش آموزی انتخاب شود که نمرات خوبی دارد احتمال دختر بودن چقدر است.

D	دختر	$P(G/D) = 0.7$	$\left((0.7) \times (0.7) \right) + \left((0.5) \times (0.4) \right) = 0.47$
P	پسر	$P(G/P) = 0.5$	
G	نمره خوب	$P(D/G) = \frac{P(G/D) \cdot P(D)}{P(G)}$	
$P(D) = 0.4$		$P(G) = P(G/D) \cdot P(D) + (P(G/P) \cdot P(P))$	$P(D/G) = \frac{0.7 \times 0.4}{0.47}$
$P(P) = 0.5$			۰.۴۷۷

نمونه آزمون آزمون فردی مثبت باشد احتمال اینکه آن فرد واقعا بیمار باشد

B	فرد بیمار	$P(B) \Rightarrow$ بیمار	چقدر است.
H	فرد سالم	$P(H) \Rightarrow$ فرد سالم	
+	نتیجه آزمون	$P(+)$ آزمون +	
			$\frac{P(+ B) P(B)}{P(+ B) P(B) + P(+ H) P(H)}$
			$\Rightarrow P(B +)$

~~10, 8, 4~~

حالت دوم

10, 8, 4

متغیر تصادفی: یک تابع است که به هر نتیجه از یک آزمون تصادفی یک عدد

اختصاص می دهد. به عبارت دیگر متغیر تصادفی، مقدار عددی یک رویداد

تصادفی را نشان می دهد. متغیرهای تصادفی به دو دسته تقسیم می شوند.

۱- متغیر تصادفی نسبتی ۲- متغیر تصادفی پیوسته

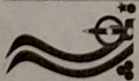
در متغیر تصادفی نسبتی مقادیر این نوع متغیر محدود و شمارش پذیر می باشد

مثلا تعداد ناس های که در یک بازی می افتد یا تعداد صندلی های که دارد یک فروشگاه

می شوند، در متغیر تصادفی پیوسته این مقادیر بین دو عدد پیوسته

است و می تواند هر مقداری را در یک بازه مشخص بگیرد. اینها در معمولاً

با اعداد حقیقی نشان داده می شوند، مثلاً حد یک شخص یا



درمائی هوا در یک روز خاص

توزیع های احتمالی

توزیع های احتمالی گسسته: برای متغیرهای تصادفی گسسته توزیع های مختلفی وجود دارند که هر کدام ویژگی های خاص خود را دارند. مهم ترین توزیع های احتمالی گسسته عبارتند از الف) توزیع برنولی: توزیع برنولی برای آزمایش های مستقل که تنها دو نتیجه ممکن دارند. مثلاً پرتاب سکه، نتایج تست با خط می باشد. (این توزیع در واقع یک متغیر تصادفی گسسته است که تنها دو مقدار منفی یا یکد را می پذیرد.) احتمال وقوع نتیجه یک (مثلاً سر) را با p نشان می دهیم و احتمال وقوع نتیجه صفر (مثلاً خط) را با $1-p$ نشان می دهیم.

$$P(X=1) = p$$

نتیجه صفر (مثلاً خط) را با $1-p$ نشان می دهیم

$$P(X=0) = 1-p$$

ب) توزیع دوجمله ای: زمانی که یک آزمایش برنولی چندین بار تکرار شود، توزیع در جمله ای بدست می آید. در این توزیع تعداد موفقیت ها (مثلاً تعداد سرها در چندین پرتاب سکه) به صورت تصادفی و براساس تعداد آزمایش ها و احتمال موفقیت در هر آزمایش مشخص می شود.

تعداد آزمایش ها (نمونه)

$$P(X=k) = \binom{n}{k} p^k (1-p)^{n-k}$$

احتمال
موفقیت

تعداد موفقیت ها
(نتایج موفقیت آمیز)

تیمین: فرض کنید سکه داریم که احتمال شیر (نتیجه موفقیت) برابر با $p = 0,4$ و احتمال خط (نتیجه شکست) برابر با $1 - p = 0,6$ می باشد. حال محاسبه کنید که احتمال وقوع هر کدام از این دو نتیجه چقدر است؟

$$P(X=1) = p = 0,4 \rightarrow 40\%$$

$$P(X=0) = 1 - p = 0,6 \rightarrow 60\%$$

تیمین: فرض می کنیم که یک سکه داریم که احتمال شیر (موفقیت) در هر پرتاب $p = 0,4$ و احتمال خط (شکست) در هر پرتاب $1 - p = 0,6$ حالا ۵ بار این سکه را پرتاب می کنیم و می خواهیم احتمال وقوع دقیقاً سه شیر (موفقیت) را در پنج پرتاب محاسبه کنیم.

$$P(X=k) = \binom{n}{k} p^k (1-p)^{n-k}$$

$10 \times 0,4^3 \times 0,6^2$

$$P(X=3) = \binom{5}{3} 0,4^3 (0,6)^{5-3} = 0,3456$$

$0,4$

$$\binom{5}{3} = \frac{5!}{3!(5-3)!} = \frac{5 \times 4 \times 3 \times 2 \times 1}{(3 \times 2 \times 1)(2!)} = 10$$

2×1

ج) توزیع پواسون این توزیع برای مدل سازی تعداد اتفاقاتی که در یک بازه زمانی یا فضای ثابت اتفاق می افتد (مثل تعداد تماس های دریافتی در یک مرکز خدمات یا تعداد دفعات در یک ساعت) این توزیع معمولاً برای اتفاقاتی استفاده می شود که احتمال وقوع آنها به طور متوسط در یک بازه زمانی مشخص ثابت باشد.

$$P(X=k) = \frac{\lambda^k e^{-\lambda}}{k!}$$

$$p(x=k) = \frac{d^k e^{-d}}{k!} \rightarrow \text{ثابت نمبر}$$

تعداد اشیاء هایی که احتمال وقوع دارند
را محاسبه می کنند

تمرین: فرض کنید یک مرکز تماس در یک روز خاص به طور متوسط ۴ تماس در هر ساعت دریافت می کند. احتمال اینکه در یک ساعت خاص این مرکز تماس دقیقاً ۳ تماس دریافت کند را محاسبه کنید.

$$p(x=k) = \frac{d^k e^{-d}}{k!}$$

$$p = P = 3 = \frac{4^3 e^{-4}}{3!} = p(x=3) = \frac{(4)^3 (2,718)^{-4}}{3!}$$

$$\frac{64 \times (0,135)}{6} = 0,145$$

تمرین: در یک کارخانه تولید به طور متوسط هر روز در نقص فنی در خط تولید رخ می دهد. محاسبه کنید که احتمال اینکه در یک روز خاص دقیقاً ۳ نقص فنی در خط تولید رخ می دهد را محاسبه کنید.

$$p(x=3) = \frac{2^3 2,718}{3!} = \frac{8 \times 0,135}{6} = 0,18$$

$$k=3$$

$$d=2$$

۱۸٪

احتمال وقوع

تیمین، یک فرد سگ، اطلاعاتی به طور متوسط ۵ ساله را به طور متوسط در بافت می کند. احتمال اینکه در یک ساعت دقیقه ۷ بار سگ در یک لحظه را مشاهده کنید.

$$p = \alpha = k$$

$$\alpha = \frac{5^7 (2,718)^{-5}}{7!} = 0,1038 \quad \boxed{10,38\%}$$

احتمال وقوع

$7! = 7 \times 6 \times 5 \times 4 \times 3 \times 2 \times 1$

جلسه سوم

قفسه جدید مرکزی

این قفسه بیان می کند که وقتی یک متغیر تصادفی به مقدار زیاد نمونه برداری شود، توزیع میانگین های نمونه ها تقریباً به توزیع نرمال نزدیک خواهد شد حتی اگر توزیع اصلی داده ها نرمال نباشد. فرض کنید شما داده های از یک جمعیت با توزیع دلخواه دارید (مثلاً ممکن است داده ها به صورت تصادفی و ناموزون توزیع شده باشند) حالا اگر از این داده ها، چندین نمونه تصادفی بگیرید و میانگین هر نمونه را محاسبه کنید، توزیع این میانگین ها تقریباً به شکل یک منحنی نرمال (گوسی) در خواهد آمد به ویژه زمانی که اندازه نمونه ها بزرگتر شود. فرض کنید یک جمعیت با توزیع دلخواه داریم و می خواهیم نمونه های از این جمعیت بگیریم اگر اندازه نمونه ها، بزرگ باشد توزیع میانگین های نمونه ها (که همان میانگین هر نمونه تصادفی است) به شکل زیر رفتار خواهد کرد. میانگین توزیع نمونه ها برابر با میانگین اصلی خواهد بود، انحراف معیار توزیع نمونه ها به نام خطای استاندارد شناخته می شود و به اندازه نمونه بستگی دارد.

هر چه اندازه نمونه بزرگتر باشد، انحراف معیار توزیع میانگین ها کوچکتر می شود، این ویژگی باعث می شود که بتوانیم بدراحتی با استفاده از توزیع نرمال تحلیل های آماری را انجام دهیم، حتی زمانی که داده های اصلی به طور طبیعی نرمال نیستند.

آری μ (میانگین جمعیت باشد)، میانگین توزیع میانگین ها نیز μ خواهد بود
 اگر σ (استیما) انحراف معیار جمعیت باشد، انحراف معیار توزیع میانگین ها
 به صورت زیر محاسبه می شود.

$$SE = \frac{\sigma}{\sqrt{n}}$$

اندازه نمونه

به کمک قضیه حد مرکزی، توزیع میانگین های نمونه ها با افزایش اندازه نمونه، بزرگ توزیع نرمال تبدیل می شود، بدین معنا که حتی اگر توزیع اصلی جمعیت نرمال نباشد، توزیع میانگین های نمونه ها، برای n بزرگ به نرمال نزدیک خواهد شد.

تمرین: فرض کنید در یک کارخانه تولیدی مقدار قطعات معیوب از میان ... با قطعه تصادفی اندازه گیری شود و میانگین آن ۵۰ درصد است، همچنین انحراف معیار مقدار قطعات معیوب در هر ۱۰۰ قطعه ۲ درصد است، حال فرض کنید که ما یک نمونه تصادفی از ۱۰۰ قطعه انتخاب می کنیم، حال احتمال اینکه میانگین قطعات معیوب از ۴ درصد بیشتر باشد را محاسبه کنید.

$$SE = \frac{\sigma}{\sqrt{n}} = \frac{2}{\sqrt{100}} = 0.2$$

تمرین ۱ فرض کنید در یک مرکز تماس، زمان پاسخگویی کارمندان به تماس مشتریان
 که به دقتی از یک توزیع غیر نرمال پیروی می کنند، میانگین کل زمان پاسخگویی ۴
 دقیقه می باشد، انحراف معیار کل ۳ دقیقه می باشد حال، نمونه هایی از ۳۴ تماس
 را تصادفی انتخاب کرده ایم، میانگین زمان پاسخگویی آنها را حساب کنید.
انحراف معیار توزیع میانگینها

$$SE = \frac{\sigma}{\sqrt{n}} = \frac{3}{\sqrt{34}} = 0.51$$

اندازه نمونه

$$\mu = 4$$

$$\sigma = 3$$

تمرین ۲ می دانیم که زمان لازم برای درست کردن یک قهوه در یک کافی شاپ
 توزیع نرمال ندارد، اما از داده های دانسته شد $\mu = 4$ ، $\sigma = 1.2$ فرض کنید
 از این کافی شاپ نمونه هایی با اندازه $n = 34$ می گیریم طبق قضیه حد مرکزی
 انحراف معیار توزیع میانگینها را برابر می دانیم.

$$SE = \frac{\sigma}{\sqrt{n}} = \frac{1.2}{\sqrt{34}} = 0.2$$

انحراف معیار توزیع
 میانگینها

اندازه نمونه

$$\mu = 4$$

$$\sigma = 1.2$$

$$n = 34$$

مفهوم میانگینها

میانگینها همان محل است جایی که داده ها تقسیم می شود.
مثال: فرض کنید نمرات دانشجو در یک آزمون ریاضی به صورت زیر است
 میانگین نمرات دانشجو را بدست آورید.

پاسخ

۱۴

۱۳

۱۲

۱۱

۱۰

$$\mu = \frac{14 + 13 + 12 + 11 + 10}{5} = 12$$

میانگین
X اقلین بار

توزیع می‌یابند

میان عددی است که در وسط داده‌های مرتب شده قرار دارد، یعنی نیمی از داده‌ها کوچکتر و نیمی بزرگتر از آن هستند، اگر تعداد داده‌ها زوج بود دو میانه برابر است با میانگین دو عدد وسط.

مثال: برای داده‌های زیر میانه را حساب کنید.

۲, ۴, ۶, ۸

$$\frac{4 + 6}{2} = 5$$

میانه

مثال: فرض کنید قد ۵ نفر به سادگی مرتب به صورت زیر می‌باشد.

۱۵۰

۱۶۰

۱۶۵

میانه

چند روشی میانه است
از آن‌ها به این روش

۱۷۰

۱۷۵



Subject :

Year :

Month :

Date :

تمرین: فرض کنید نمرات ۴ دانشجو به شکل زیر می باشد. مایلی

۲۰ ۱۹ ۱۸ ۱۵ ۱۲ ۱۰

$$\begin{array}{r} ۳۳ \\ ۱۵ + ۱۸ \\ \hline ۳ \end{array} = ۱۹,۵$$

۱۵

مد

عددی است که بیشترین تکرار را دارد، یعنی عددی که بیش از بقیه، در داده ها ظاهر شده است.

مثال: نمرات ۸ دانشجو در درس آمار و احتمالات به شکل زیر می باشد.

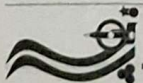
مد آن را بیابید.
۱۵, ۱۵, ۱۴, ۱۳, ۱۷, ۱۸, ۱۵, و ۷

مد آن ۱۵ است

تمرین: به ۷ دانشجو، پروژه ای داده ایم، داده های آنها به شرح زیر می باشد.
مایلی، مایه، و مد اینها را محاسبه کنید.

۷۵ ۱۰۰۰ ۷۰ ۹۵ ۶۰ ۵۵ ۵۰

صاف



Date:

Tab Back Main Run Stop Speed Exit

Subject:

میانگین

$$\bar{X} = \frac{50 + 55 + 40 + 45 + 70 + 75 + 100}{7} = \frac{435}{7} = 62,14$$

میانگین

$$X: 50, 55, 40, 45, 70, 75, 100 = 45$$

میانگین

میانگین ندارد

مثال: فرض کنید درآمد ماهانه ۷ نفر از یک گروه کاری به شرح زیر است

$$15.000 + 20.000 + 22.000 + 25.000 + 30.000 + 35.000 + 40.000 = 187.000$$

$$x = 25.000$$

میانگین و میانگین دوم

انحراف معیار

اندازه ای است که نشان می دهد داده ها به طور متوسط چقدر از میانگین فاصله دارند. انحراف معیار کم نشان دهنده این است که داده ها به طور کلی نزدیک به میانگین قرار دارند. انحراف معیار زیاد نشان دهنده این است که داده ها از میانگین پراکنده شده اند.

$$s = \sqrt{\frac{\sum_{i=1}^N (x_i - \bar{x})^2}{N}}$$

انحراف معیار جمعیت

تعداد داده ها

میانگین جمعیت

نمات ۵ دانشجو به شرح زیر است

۶۰ ۷۰ ۸۰ ۹۰ ۱۰۰

توضیح سوال: برای محاسبه انحراف معیار ابتدا میانگین داده‌ها را محاسبه می‌کنیم سپس کم کردن میانگین از هر داده و سپس مربع کردن آنها بعد از آن جمع کردن مربع‌های اختلافات بعد از آن تقسیم بر تعداد داده‌ها می‌کنیم و در نهایت جذر را از حاصل تقسیم می‌گیریم

$$\textcircled{1} \bar{X} = \frac{۶۰ + ۷۰ + ۸۰ + ۹۰ + ۱۰۰}{۵} = ۸۰$$

$$(۶۰ - ۸۰)^2 = ۴۰۰$$

$$(۷۰ - ۸۰)^2 = ۱۰۰$$

$$(۸۰ - ۸۰)^2 = ۰$$

$$(۹۰ - ۸۰)^2 = ۱۰۰$$

$$(۱۰۰ - ۸۰)^2 = ۴۰۰$$

$$\textcircled{2} ۴۰۰ + ۱۰۰ + ۰ + ۱۰۰ + ۴۰۰ = ۱۰۰۰$$

$$\textcircled{3} \frac{۱۰۰۰}{۵} = ۲۰۰ \quad \textcircled{4} \sqrt{۲۰۰} = ۱۴,۱۴$$

هر داده را از میانگین کم کردن و تقسیم بر تعداد کل داده‌ها می‌کنیم

مثال: داده‌های آمارگاهی دانشجو به شرح زیر می‌باشد. انحراف معیار جمعیت آن را حساب کنید.

۴۰, ۴۵, ۶۰, ۶۵, ۵۰

$$\frac{۴۰ + ۴۵ + ۶۰ + ۶۵ + ۵۰}{۵} = ۵۲$$

$$(۴۰ - ۵۲)^2 = ۱۴۴$$

$$(۴۵ - ۵۲)^2 = ۴۹$$

$$(۶۰ - ۵۲)^2 = ۶۴$$

$$(۶۵ - ۵۲)^2 = ۱۴۹$$

$$(۵۰ - ۵۲)^2 = ۴$$

جمع ۴۳۰

$$\frac{۴۳۰}{۵} = ۸۶$$

$$\sqrt{۸۶} = ۹,۲۷۴$$

تربیه داده های زیر میزان آلودگی در ۵ نمونه آب متفاوت را نشان می دهد واریانس و انحراف معیار داده را بدست آورید.

$$\begin{aligned} 50 & \quad \frac{50 + 55 + 40 + 45 + 70}{5} = 46 \\ 55 & \\ 40 & \quad (50 - 40)^2 = 100 \\ 45 & \quad (55 - 40)^2 = 225 \\ 70 & \quad (40 - 40)^2 = 0 \quad \text{و} \quad \frac{450}{5} = 90 \\ & \quad (45 - 40)^2 = 25 \\ & \quad (70 - 40)^2 = 900 \end{aligned}$$

$$\sqrt{90} = 9.47$$

$$100 + 225 + 100 + 225 = 650$$

انحراف معیار

آزمون های فرضی

این صحبت، ابزارهای آماریست که به ما کمک می کند که در جواب باراده ها تصمیمات علمی و معتبر بگیریم. آزمون فرضی به طور کلی یک روش آماری است که برای ارزیابی صحبت یا تدریسی یک فرضیه در مورد جمعیت بکار می رود. این فرایند شامل استفاده از داده های نمونه برای ارزیابی اینکه آیا تئوری خاصی برای ما پذیرفتنی خاص و مورد دارد یا نه می باشد. آزمون فرضی معمولاً در دو مرحله انجام می شود.

① تعریف فرضیه‌ها ② ارزیابی داده‌ها و تصمیم‌گیری بر اساس شواهد موجود

① در هر آزمون فرض، دو فرضیه اصلی وجود دارد

فرضیه صفر (H_0): این فرضیه معمولاً فرضی است که می‌خواهیم آن را رد کنیم، معمولاً فرضیه صفر بیان می‌کند که هیچ تفاوت یا اثر خاصی در داده وجود ندارد.

فرضیه جایگزین (H_1 یا H_a): این فرضیه، مخالف فرضیه صفری است و معمولاً بیان می‌کند که یک تفاوت یا اثر خاصی در داده‌ها وجود دارد. به عنوان مثال اگر بخواهیم آزمون کنیم که آیا میانگین قد افراد در یک جامعه برابر با ۱۷۰ سانتی متر است، فرضیه‌ها به صورت زیر خواهند بود.

(H_0): میانگین قد افراد برابر با ۱۷۰ سانتی متر است

(H_1 یا H_a): میانگین قد افراد متفاوت از ۱۷۰ سانتی متر است

② پس از تکمیل فرضیه‌ها باید یک آزمون آماری انتخاب کنیم که به کمک آن می‌توانیم شواهد را بررسی کنیم و تصمیم بگیریم که آیا فرضیه صفر را رد کنیم یا نه؟ مراحل کار شامل موارد زیر می‌باشد.

① محاسبه آماره آزمون

این مقدار معمولاً با استفاده از داده‌ها، نمونه و فرمول‌های خاص هر آزمون آماری محاسبه می‌شود.

② انتخاب سطح معنادار (α) (آلفا) این سطح احتمال است

که حای از احتمال اشتباه در رد فرضیه صفر می‌باشد، معمولاً $(\alpha = 0.05)$ انتخاب می‌شود که به معنای این است که ۵ درصد احتمال داریم که فرضیه صفر را اشتباه رد کنیم.

۳) مقایسه آماری آزمون با مقدار بحرانی یا استفاده از P-Value

احتمال مشاهده داده‌های نمونه با P داده‌های مشاهده‌شده تحت فرضیه صفر می‌باشد، اگر کوچکتر از سطح معنادار α باشد، فرضیه صفر رد می‌شود.

۴) مقدار بحرانی

در این روش باید آماره آزمون را (مثل توزیع) با مقدار بحرانی که از جدول توزیع مربوطه گرفته می‌شود مقایسه کنیم.

انواع آزمون‌های فرضی

آزمون‌های فرضی انواع مختلفی دارند که بر اساس نوع داده‌ها و نوع آزمون بکار می‌روند. برخی از مهم‌ترین آزمون‌ها را توضیح خواهیم داد.

۱) آزمون student - t

برای مقایسه میانگین‌های دو گروه از داده‌ها استفاده می‌شود. معمولاً زمانی که اندازه نمونه‌ها کوچک است، این آزمون می‌تواند، برای آزمون یک نمونه (برای مقایسه میانگین با یک مقدار ثابت) یا دو نمونه مستقل (برای مقایسه دو میانگین نمونه‌ای) استفاده می‌شود.

۲) آزمون Z

برای مقایسه میانگین‌ها یا نسبت‌ها در داده‌هایی که از توزیع نرمال پیروی می‌کنند و معمولاً برای نمونه‌های بزرگ ($n > 30$) استفاده می‌شود. آزمون Z همچنین برای مقایسه یک میانگین با مقدار مشخص در زمانی که انحراف معیار جامعه را می‌دانیم بکار می‌رود.

۳) آزمون chi - square

این آزمون برای بررسی ارتباط بین دو متغیر کیفی یا مقایسه توزیع‌های مشاهده شده با توزیع‌های مورد انتظار استفاده می‌شود. این آزمون به طور گسترده در مطالعات نظر سنجی و تحلیل جدول‌های توانایی نگار می‌رود.

⑤ آزمون ANOVA (تحلیل واریانس)

این آزمون برای مقایسه میانگین‌ها در سه یا بیشتر گروه‌های مستقل استفاده می‌شود. هدف آن بررسی این است که آیا اختلاف معناداری بین میانگین‌های گروه‌ها وجود دارد یا نه.

مفاهیم کلیدی در آزمون‌های فرضی

خطا نوع اول (α): احتمال رد فرضیه صفر، زمانی که فرضیه صفر درست است به عبارتی دیگر، خطا نوع اول یعنی اشتباه در رد فرضیه صحیح.

خطا نوع دوم (β): احتمال نپذیرفتن فرضیه صفر، زمانی که فرضیه صفر غلط می‌باشد، به عبارتی دیگر، فرضیه صفر باید رد شود ولی اشتباهاً نپذیرفته می‌شود.

مثال فرضی: می‌خواهیم بررسی کنیم که آیا میانگین قد دانش‌آموزان یک کلاس با ۱۷ سانی متر متفاوت است یا خیر. برای این کار مراحل زیر را طی می‌کنیم.

① فرضیه‌ها: میانگین قد دانش‌آموزان برابر با ۱۷۰ سانتی متر می‌باشد.
(H_0)

② میانگین قد دانش‌آموزان متفاوت با ۱۷۰ سانتی متر می‌باشد.
(H_1)

- ۱) انتخاب سطح معیار: فرض می‌کنیم که $\alpha = 0.05$
- ۲) جمع آوری داده‌ها: فرض می‌کنیم اندازه نمونه ۳۵ نفر باشد و میانگین قد نمونه ۱۷۲ سانتی متر و انحراف معیار ۱۰ سانتی متر باشد.
- ۳) آزمون Z: برای انجام آزمون Z از جدول زیر استفاده می‌کنیم.

$$Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} = Z = \frac{172 - 170}{\frac{10}{\sqrt{35}}} = 1.195$$

اثر از معیار
تعداد داده‌ها

۴) مقایسه با مقدار بحرانی یا استفاده از P-Value

با توجه به جدول توزیع Z و سطح معیار $\alpha = 0.05$ مقدار بحرانی برای آزمون دوطرفه حدود ۱.۹۶ می‌باشد چون $1.195 < 1.96$ فرضیه صفر را رد نمی‌کنیم، در این مثال چون مقدار Z می‌باشد کمتر از مقدار بحرانی می‌باشد یعنی نتوانیم فرضیه صفر را رد کنیم و نتیجه می‌گیریم که مشاهده کافی برای ادعا اینکه میانگین قد دانش آموزان با ۱۷۰ سانتی متر متفاوت است وجود ندارد.

کاربردهای آزمون‌های فرضی

در تحقیقات علمی برای تأیید یا رد فرضیه‌های علمی در آمار پیرایشی برای بررسی اثر بخشی دارو ها یا درمان‌ها در آمار اقتصادی برای تحلیل داده‌های اقتصادی و پیش‌بینی روندهای اقتصادی، در بازاریابی برای بررسی اثر تبلیغات یا تغییرات در قیمت و مقدار مصرف کننده‌ها

تمرین: میزان آفات ۶ درخت به صورت زیر ارائه شده است
مطلوب است واریانس؟ میانگین؟ انحراف معیار؟ مد؟

۲۲

۲۴

۲۱

۲۰

۲۵

۲۴

$$\bar{x} = \frac{22 + 24 + 21 + 20 + 25 + 24}{6} = 22,4$$

میانگین

$$mod = 24$$

بیشترین تکرار

$$\bar{x}_{med} = \frac{22 + 24}{2} = \frac{46}{2} = 23$$

مد

$$(x_1 - 22,4)^2 = 0,2$$

$$(x_2 - 22,4)^2 = 3,2$$

$$(x_3 - 22,4)^2 = 0,36$$

$$(x_4 - 22,4)^2 = 4,76$$

$$(x_5 - 22,4)^2 = 4,76$$

$$(x_6 - 22,4)^2 = 0,16$$

$$0,2 + 3,2 + 0,36 + 4,76 + 4,76 + 0,16 = 13,44$$

$$s^2 = \frac{13,44}{6} = 2,24$$

$$s = \sqrt{2,24} = 1,49$$

۱۷۵
مثال: در یک تحقیق میانگین قد مردان در یک کشور برابر با ۱۷۵ cm است. یک نمونه نمونه‌ای از ۱۷۸ مرد به طور تصادفی انتخاب شده است و میانگین قد آن‌ها ۱۷۸ سانتی متر می‌باشد. انحراف معیار جامعه برابر با ۸ سانتی متر است. مطلوب است محاسبه آمار آزمون Z؟

Date:



Subject:

$$Z = \frac{1V\Delta - 1V\Delta}{\frac{1}{\sqrt{\epsilon_0}}} = \frac{3}{1,295} = 2,3V$$

$$2,3V > 1,99 \quad \text{فرضیه درست است}$$

$$\alpha = 9,95$$

Date:

Sat Sun Mon Tue Wed Thu Fri

Subject:

$$t_1 = t_2 = t$$

از هر یک زمان طی روزها مذکور نام برابر است
در نتیجه می توان نوشت:

$$\bar{V} = \frac{x_1 + x_2}{t_1 + t_2}$$

$$\bar{V} = \frac{x_1 + x_2}{t} = \frac{1}{2} \left(\frac{x_1}{t} + \frac{x_2}{t} \right) = \frac{1}{2} (v_1 + v_2) = \frac{1}{2}$$

$$t_1 = t_2 = t$$

$$\frac{1}{2} (90 + 30) = 60 \text{ m/s}$$

چندک

با توجه به تعاریفی که برای چندک جامعه و چندک نمونه‌ای وجود دارد، چندک α ام کمی است که در اصطلاح α چندم داده‌ها کوچکتر یا مساوی آن هستند به عنوان مثال چندک هفتم مقدار است که $\frac{7}{100}$ مشاهدات کوچکتر یا مساوی آن هستند، برخی از چندک‌ها به دلیل کاربرد فراوان، اسامی خاصی دارند به عنوان مثال چارک α ام، چندکی است که $\frac{\alpha}{4}$ ام مشاهدات کوچکتر یا مساوی آن هستند و یا Q نامی داده می‌شود. Q_1 صدک α ام نقطه‌ای است که $\frac{\alpha}{100}$ و یا $\frac{\alpha}{100}$ مشاهدات کوچکتر یا مساوی آن می‌باشند. ما هم دیتای مانند دهک و پینچ‌هاک نیز به طور فامشابه قابل تقریب هستند، چندک‌ها لزوماً بیان کننده مکان تمرکز مشاهدات نیستند، بلکه مکانی از توزیع را نشان می‌دهند که کسری از مشاهدات کوچکتر یا مساوی آن هستند.

گونه محاسبه چندک‌ها

① مشاهدات را به صورت صعودی مرتب کرده و مشاهده مرتب شده α ام

را به صورت $X_{(\alpha)}$ نشان می‌دهیم

② اگر p نسبت چندک یا کتر مشاظر با آن و n اندازه نمونه باشد آنگاه کمیته C را از رابطه زیر محاسبه می‌کنیم:

$$C = (n+1)p$$

③ فرض کنید نسبت صحیح کمیته C را با r و نسبت اعشاری

آن را با w نشان دهیم در این صورت چندک به صورت زیر بدست می‌آید:

$$(1-w)X_{(r)} + wX_{(r+1)} = X_{(r)} + w[X_{(r+1)} - X_{(r)}]$$

تقریباً سن یازده نفر از مراجعین به لیسک اطفال به صورت

11, 10, 9, 13, 12, 11, 10, 11, 12, 13, 14, 15, 16

گزارش شده است، پنجمین هفتک ایستاده هارا حساب کنید

① ابتدا مساعداً را به صورت سعودی مرتب می‌کنیم

1, 11, 12, 13, 13, 14, 15, 14, 16, 16, 17

۲) کسر متناظر با یخچال هفت را به صورت $\rho = \frac{A}{V}$ می نویسیم و C را محاسبه می کنیم.

$$C = (n+1)p \rightarrow C = (n+1) \frac{\Delta}{V} = \Lambda + \frac{p}{V}$$

$$1^\circ \frac{\Delta}{V} = \frac{42}{V} = 1 + \frac{15}{V}$$

(۳) حال با توجه به عدد بدست آمده $\frac{4}{V}$ آن را به صورت $\frac{A+C}{V}$ می نویسیم چون

می خواهیم مقدار r و W ، ایدست آوریم

$$r = \lambda \quad w = \frac{\lambda}{v}$$

$$X + w \left[X_{(q)} - X_{(n)} \right] =$$

$$(14 + \frac{R}{V} [1V - 14]) = 14, \Delta V 1$$

معیارهای مقایسه‌ای

معیارهای مقایسه‌ای، معیارهایی هستند که برای براندازی حول پارامترهای مکانی یک توزیع را اندازه‌گیری می‌کنند. اینها معیارها و وابستگی، میانگین، و قدر مطلق انحرافات، دامنه بین جایی و فریب تغییرات، از معیارهای براندازی معروف می‌باشند.

دامنه

اگر $x_1, x_2, \dots, x_n$ یک نمونه n تایی بدست آمده از یک جامعه باشد، آنگاه دامنه نمونه اختلافات مقداری کوچکترین و بزرگترین مشاهده می‌باشد.

$$R = \max(x_i) - \min(x_i)$$

با جایگزینی اندازه جامعه N بجای n جایگزین اندازه n دامنه جامعه را نیز می‌توان بدست آورد.

دامنه به وجود نقاط دور افتاده به شدت حساس است و در صورت وجود نقطه دور افتاده، مقدار دامنه به شدت تحت تأثیر طبعی بزرگ می‌شود، همچنین دامنه معیار است که از اطلاعات موجود در تمام مشاهدات، استفاده نمی‌کند و فقط به اطلاعات کوچکترین و بزرگترین مشاهده نیازمند است. بنابراین می‌توان از پراکندگی مشاهده نیازمند، اما نزدیک آن به آنگاه استفاده از آن می‌باشد.

نکته: یکی از مهمترین مزایای اینها، نسبت به معیارهای مانند دامنه را می‌توان با شکل بررسی کرد. همانطور که از فرم تابعی اینها توزیع‌ها، مشخص است، اینها دو شکل پراکندگی متفاوت دارند. اما دامنه هر دو توزیع تقریباً با هم برابر است. در حالی که اینها معیار آنها متفاوت می‌باشد. دو شکل زیر توزیع‌هایی

با روش پیرامون آنرا و محاسبه متفاوت راندها می دهند.



۱۳۸۲، ۹، ۲۲ جلسه هفتم

تحلیل واریانس یکطرفه

فرض کنید k جامعه مختلف با توزیع نرمال و میانگین های به ترتیب $\mu_1, \mu_2, \dots, \mu_k$ و واریانس مشترک و مجهول σ^2 داریم. می خواهیم فرض را بر این بگذاریم که این جوامع را آزمودن کنیم. در این حالت، فرض صفر و فرض مقابل آن را می توان به صورت زیر نمایش داد.

$$H_0: \mu_1 = \mu_2 = \dots = \mu_k$$

$$H_1: \mu_i \text{ ها با یکدیگر متفاوت است}$$

برای آزمون این فرض k نمونه تصادفی n تایی از جوامع به صورت زیر در نظر می گیریم.

Date: / /

Sat Sun Mon Tue Wed Fri

Subject:

جامعه ۱	X_{11}	X_{12}	...	X_{1n}
جامعه ۲	X_{21}	X_{22}	...	X_{2n}
.	.	.	.	.
.	.	.	.	.
.	.	.	.	.
جامعه k	X_{k1}	X_{k2}	...	X_{kn}

X_{ij} مشاهده i ام از نمونه جامعه i ام می باشد که به طور مستقل دارای توزیع نرمال با میانگین μ_i و واریانس σ_i^2 می باشد به طور محاسباتی X_{ij} را می توان به صورت زیر نمایش داد

$$X_{ij} = \mu_i + \epsilon_{ij}$$

$$i = 1, 2, \dots, k$$

$$j = 1, 2, \dots, n$$

که ϵ_{ij} دارای توزیع نرمال با میانگین صفر و واریانس σ_i^2 می باشد.

$$\varepsilon_{ij} = X_{ij} - \mu_i$$

$$i = 1, 2, \dots, k$$

$$j = 1, 2, \dots, n$$

ای امکان
بر آنکه تقسیم حاصل اثر به مدل های پیچیده تر کلیل واریانس وجود داشته باشد، معمولاً آن را به صورت زیر نمایش می دهند:

$$X_{ij} = \mu + \tau_i + \varepsilon_{ij}$$

$$i = 1, 2, \dots, k$$

$$j = 1, 2, \dots, n$$

که در اینجا μ میانگین کل و τ_i را اثر جامعه نام ما اثر تیار نام می برند.

$$\sum_{i=1}^k \tau_i = 0$$

برای انجام آزمون، از تغییر پذیری موجود در داده ها استفاده می شود، به این منظور ابتدا تغییر پذیری کل موجود در نمونه های همه جوامع را به صورت زیر نشان می دهیم.

$$\sum_{i=1}^k \sum_{j=1}^n \left(X_{ij} - \bar{X} \right)^2$$

میانگین نمونه ای
کل داده ها

$$\bar{X} = \frac{\sum_{i=1}^k \sum_{j=1}^n X_{ij}}{kn}$$

ضریب تغییرات

اگر $X_1, X_2, \dots, X_n$ مقادیر بدست آمده از یک نمونه باشند، آنگاه ضریب تغییرات نمونه را می توان از رابطه زیر بدست آورد.

$$C.V = \frac{S}{\bar{X}}$$

نکته: چنانچه از رابطه ضریب تغییرات بدست آید، این معیار بدون واحد است در تقیید برای مقایسه پراکندگی دو مشخصه آماری که واحد هایشان یکسان نباشد می توان از این معیار استفاده کرد، همچنین برای مقایسه پراکندگی نسبی دو مشخصه که میانگین های متفاوتی دارند نیز از این معیار استفاده می شود.

نمونه: میانگین و انحراف معیار وزن و قد یک گروه از افراد به صورت جدول زیر من باشد، کدام یک از این دو مشخصه (قد و وزن) پراکندگی بیشتری دارد؟

	وزن (پوند)	قد (اینچ)
$\bar{X}$	۱۴۵	۴۹
S	۲۵	۱۲

$$C.V_1 = \frac{12}{49} = 0.2449$$

$$C.V_2 = \frac{25}{145} = 0.1724$$

نکته: چون در صفحه قد و وزن واحد متفاوتی دارند، بهترین شاخص برای اندازه گیری پیکاندی، ضریب تغییرات است. بنابراین پیکاندی نسبت قد و وزن در این مورد تقریباً یکسان است، اما پیکاندی نسبت قد کمی بیشتر از وزن است.

معیارهای شکل تابع

۵- هی اوقات دو توزیع مختلف معیارهای مکانی و پیکاندی یکسانی دارند اما شکل آنها کاملاً باهم متفاوت است، به عبارت دیگر شمار اول و دوم این توزیع ها یکسان بوده اما شمار هاس بالاتر آنها باهم فرق می کند. وجود معیار هاسی که بتواند گستره هاس بالاتر را اندازه گیری کند می تواند در بسیاری از تحلیل ها مفید باشد.

چولگی (Skewness)

اگر $x_1, x_2, \dots, x_n$ یک نمونه n تایی از یک جامعه باشد آنگاه چولگی نمونه به صورت زیر تعریف می شود.

$$b \frac{m_3}{(m_2)^{3/2}}$$

Date:

Subject:

۲۹، ۹، ۵۴

آزمون: هر کس که یک نمره شش وجهه استاندارد را بگیرد، پخش کند.
احتمال اینکه عدد ۴ یا بیشتر بیوفته چقدر است؟

$$P(4) = \frac{3}{6} = \frac{1}{2}$$

$\frac{1}{2}$ احتمال دارد عدد ۴ یا بیشتر بیاید.
تعداد حالات مطلوب ۳ و ۵ و ۶

نمونه: یک سکه را ۳ بار پرتاب می‌کنیم، احتمال اینکه دقیقاً دو بار
پایست سکه بیوفته را محاسبه کنید.

$$2^3 = 8 \quad P(2) = \frac{3}{8}$$

نمونه: یک دانشجو نمرات زیر را در کلاس بدست آورده است، میانگین
میانگین و واریانس نمرات این دانشجو را حساب کنید.

$$\begin{array}{cccccc} 12 & 14 & 14 & 11 & 20 & 2 \end{array} \quad \sum (x_i - \mu)^2$$

$$\frac{12 + 14 + 14 + 11 + 20}{5} = \frac{71}{5} = 14.2$$

$$(12 - 14)^2 = 4$$

$$(14 - 14)^2 = 0$$

$$(11 - 14)^2 = 9$$

$$(20 - 14)^2 = 36$$

$$(2 - 14)^2 = 144$$

$$\frac{169}{5} = 33.8$$

$$\sqrt{33.8} = 5.81$$

کمریای آن قد افراد در یک جامعه به طور نرمال توزیع شده باشد با میانگین ۱۷۰ سانتی متر و انحراف معیار ۱۰ سانتی متر احتمال اینکه فردی که تصادفی از این جامعه انتخاب می شود قدش بین ۱۴۰ و ۱۸۰ سانتی متر داشته باشد چقدر است؟ با توجه به آنکه می دانیم $P(-1 < Z < 1) = 0.2420$

$$\frac{170 - 180}{10} = \frac{-10}{10} = (-1)$$

$$\frac{170 - 140}{10} = \frac{30}{10} = (3)$$

کمریای در یک مسابقه شطرنج هر بازیکن دو حرکت می تواند انجام دهد اگر دو بازیکن در هر حرکت، در انتخاب مختلف باشند مقدار کل حالت ها برای یک دور بازی چقدر خواهد بود.

$$2 \times 2 = 4$$