



به نام خداوند بخشنده مهربان

آمار و روش تحقیق

(پیشرفته کمی و کیفی)

مجموعه:

علوم اجتماعی

مؤلف:

گروه مولفین ماهان



آمادگی آزمون دکتری

گروه مولفین ماهان

آمار و روش تحقیق پیشرفته کمی و کیفی مجموعه علوم اجتماعی / دکتر نظام بهرامی کمیل

مشاوران صعود ماهان: ۱۴۰۴

۴۵۷ص: جدول، نمودار (آمادگی آزمون دکتری مجموعه علوم اجتماعی)

ISBN/N: 978-600-458-650-4

فهرست نویسی بر اساس اطلاعات فیپا.

فارسی - چاپ اول

۱- آمار و روش تحقیق پیشرفته کمی و کیفی

۲- آزمونها و تمرینها

۲- آزمون دوره های تحصیلات تکمیلی

۴- دانشگاه ها و مدارس عالی - ایران - آزمونها

ج - عنوان

کتابشناسی ملی ۴۴۱۸۲۴۰

انتشارات مشاوران صعود ماهان



موسسه آموزش عالی آزاد
www.mahan.ac.ir

☐ نام کتاب: آمار و روش تحقیق پیشرفته کمی و کیفی

☐ مدیر مسئولان: مجید و هادی سیاری

☐ مولف: گروه مولفین ماهان

☐ مدیر تولید و محتوا: سمیه بیگی

☐ ناشر: مشاوران صعود ماهان

☐ نوبت و تاریخ چاپ: اول / ۱۴۰۴

☐ تیراژ: ۱۰۰۰ نسخه

☐ قیمت: ۶/۷۲۰/۰۰۰ ریال

☐ شابک: ISBN ۹۷۸-۶۰۰-۴۵۸-۶۵۰-۴

انتشارات مشاوران صعود ماهان: تهران - خیابان ولیعصر، بالاتر از تقاطع ولیعصر مطهری، پلاک ۲۰۵۰

تلفن: ۸۸۱۰۰۱۱۳ و ۸۸۴۰۱۳۱۳

کلیه حقوق مادی و معنوی این اثر متعلق به موسسه آموزش عالی آزاد ماهان می باشد. و هرگونه اقتباس و کپی برداری از این اثر بدون اخذ مجوز پیگرد قانونی دارد.

بنام خدا

ایمان داریم که هر تغییر و تحول بزرگی در مسیر زندگی بدون تحول معرفت و نگرش میسر نخواهد بود. پس بیایید با اندیشه توکل، تفکر، تلاش و تحمل در توسعه دنیای فکریمان برای نیل به آرامش و آسایش توأمان اولین گام را برداریم. چون همگی یقین داریم دانایی، توانایی می‌آورد.

شاد باشید و دلی را شاد کنید

برادران سیاری

فصل اول: آمار، احتمالات و نمونه‌گیری.....	۱۱
گفتار اول: آمار.....	۱۱
گفتار دوم: احتمالات.....	۲۱
گفتار سوم: نمونه‌گیری.....	۳۷
فصل دوم: مفاهیم و گونه‌شناسی تحقیق.....	۴۷
مقدمه.....	۴۸
گفتار اول: مفاهیم.....	۵۱
روش.....	۵۱
علم.....	۵۱
اهداف علم.....	۵۱
استراتژی‌های پژوهش.....	۵۲
شرایط و ویژگی‌های پژوهش.....	۵۳
ویژگی‌های پژوهشگر.....	۵۳
ارزش‌ها.....	۵۴
نظریه.....	۵۴
نظریه برد متوسط.....	۵۷
مدل.....	۵۸
گفتار.....	۵۹
قانون.....	۵۹
آگزیوم.....	۵۹
مفهوم.....	۶۰
شاخص.....	۶۱
پیشنهاد.....	۶۱
فرضیه.....	۶۱
انواع فرضیه.....	۶۱
مراحل فرضیه‌سازی.....	۶۲
متغیر.....	۶۳
انواع متغیر.....	۶۳
سطوح سنجش متغیرها.....	۶۵
انواع تعریف.....	۶۶
تبیین.....	۶۷
علیت.....	۶۹
راه‌های کسب شناخت.....	۷۰
مغالطه محیط‌گرایی و تقلیل‌گرایی.....	۷۱
گفتار دوم: گونه‌شناسی تحقیقات.....	۷۳
مقدمه.....	۷۳
۱. تحقیقات به لحاظ مبانی هستی‌شناختی (کمی و کیفی).....	۷۴
۲. تحقیقات بر اساس کاربرد.....	۷۵

۳.	تحقیقات از نظر زمان	۷۷
۴.	تحقیقات از نظر عمق تحقیق	۷۸
۵.	تحقیقات بر اساس نوع رابطه	۷۸
۶.	تحقیقات بر اساس روش گردآوری داده‌ها	۷۹
۷.	تحقیقات بر اساس نقش متغیرها	۷۹
۸.	تحقیقات بر اساس هدف	۷۹
۸۱	گفتار سوم: روش‌ها و تکنیک‌های تحقیق	
۱.	روش پیمایش	۸۱
	تکنیک پرسش	۸۱
۲.	روش آزمایش	۸۴
	انواع آزمایش	۸۵
۳.	تحقیق میدانی	۸۶
	انواع مشاهده	۸۸
۴.	تحلیل محتوا	۸۸
	انواع تحلیل محتوا	۸۹
۵.	روش دلفی	۹۲
۶.	ایده‌آل تایپ	۹۲
۷.	سنجش‌های بدون مزاحمت	۹۳
۸.	تحلیل ثانوی	۹۴
۹.	تحلیل آمارهای موجود	۹۴
۱۰.	روش کیو	۹۴
	خلاصه مطالب	۹۵
۹۹	فصل سوم: تحقیقات کمی	
۱۰۱	گفتار اول: روایی و پایایی	
	روایی ۱۰۱	
	انواع روایی	۱۰۱
	پایایی	۱۰۳
	انواع پایایی	۱۰۳
	خلاصه مطالب	۱۰۵
۱۰۷	گفتار دوم: سنجش و طیف‌ها	
	سطوح سنجش	۱۰۷
	شاخص	۱۰۸
	انواع تعریف شاخص	۱۰۸
	اصول شاخص‌سازی	۱۰۹
۱۱۷	گفتار سوم: آزمون‌های آماری	
	مقدمه	۱۱۷
	تقسیم‌بندی آزمون‌های ناپارامتریک	۱۱۹
	ضرایب همبستگی	۱۲۰
	ضریب همبستگی رتبه‌ای اسپیرمن	۱۲۱
	گاما	۱۲۳
	تقسیم‌بندی آزمون‌های آماری براساس تعداد و نوع متغیرها	۱۲۵
	الف. آزمون‌های دو متغیری اسمی	۱۲۵

ب. آزمون‌های دو متغیره ترتیبی	۱۲۶
پ. آزمون‌های متغیرهای فاصله‌ای و نسبی	۱۲۷
رگرسیون	۱۳۵
رگرسیون چند متغیری	۱۳۶
تحلیل رگرسیون لجستیک	۱۳۷
ضریب همبستگی درون طبقه‌ای Eta تحلیل واریانس	۱۳۹
تحلیل عامل	۱۴۲
مراحل تحلیل عامل	۱۴۲
روش‌های کاهش عوامل	۱۴۵
مفهوم کنترل آماری متغیرها	۱۴۸
همبستگی پاره‌ای	۱۴۸
همبستگی جزئی	۱۴۹
تحلیل مسیر	۱۵۰
مفروضات تحلیل مسیر	۱۵۱
مراحل تحلیل مسیر	۱۵۱
خلاصه روش تحلیل مسیر	۱۵۸
خلاصه آزمون‌های پارامتری و ناپارامتری	۱۵۸
آزمون یک دامنه و دو دامنه	۱۶۱
روش تحلیل مناسب برای یک متغیر مستقل اسمی و یک متغیر وابسته فاصله‌ای	۱۶۴
مقایسه میانگین‌ها	۱۶۴
آزمون تی استیودنت	۱۶۴
فصل چهارم: فلسفه تحقیق و تحقیقات کیفی.....	۱۶۷
مقدمه	۱۶۸
گفتار اول: کلیات.....	۱۶۹
۱. هستی‌شناسی	۱۶۹
۲. شناخت‌شناسی	۱۶۹
۳. انسان‌شناسی	۱۷۰
۴. روش‌شناسی	۱۷۱
۵. دوگرایی	۱۷۱
۶. جبر و اختیار	۱۷۲
۷. داده‌های کمی و کیفی	۱۷۲
۸. سطح خرد و کلان	۱۷۳
۹. استقراء و قیاس	۱۷۴
۱۰. عینی و عین‌گرایی در مقابل ذهنی و ذهن‌گرایی	۱۷۵
۱۱. مطلق‌گرایی و نسبی‌گرایی	۱۷۶
۱۲. ذات‌گرایی و غیر ذات‌گرایی	۱۷۶
۱۳. اثبات‌گرایی و غیر اثبات‌گرایی	۱۷۷
۱۴. تجربه‌گرایی	۱۷۸
۱۵. رفتارگرایی	۱۸۰
۱۶. ماتریالیسم دیالکتیکی	۱۸۱
۱۷. عقل‌گرایی انتقادی	۱۸۲
۱۸. هرمنوتیک	۱۸۲

۱۸۴	۱۹. ساخت‌گرایی
۱۸۵	۲۰. درک عمومی
۱۸۶	۲۱. مقایسه رویکردهای اثباتی، انتقادی و تفسیری
۱۸۸	نتیجه‌گیری
۱۹۱	گفتار دوم: فیلسوفان روش
۱۹۱	۱. فرانسیس بیکن
۱۹۱	۲. رنه دکارت
۱۹۲	۳. کارل پوپر
۱۹۴	۴. ایمره لاکاتوش
۱۹۵	۵. توماس کوهن
۱۹۷	۶. پاول کارل فیرابند
۲۰۰	۷. لودویک ویتگنشتاین
۲۰۲	۸. پست‌مدرنیسم
۲۰۵	گفتار سوم: روش‌های کیفی
۲۰۵	۱. تفاوت‌های روش‌های کمی و کیفی
۲۰۶	۲. تفاوت‌های روش کیفی و کمی در جریان پژوهش
۲۰۶	۳. نظریه زمینه‌ای (مینایی یا خاکی)
۲۰۷	۴. تحلیل تاریخی
۲۰۸	۵. قوم‌نگاری و پدیدارشناسی
۲۱۰	۶. گروه‌های کانونی
۲۱۱	۷. فراتحلیل
۲۱۵	۸. تحلیل گفتمان
۲۱۸	۹. تحقیق کنشی (اقدام پژوهی)
۲۱۸	۱۰. نمایش‌پردازی
۲۱۹	۱۱. مثلث‌سازی
۲۲۳	پرسش‌های تستی
۲۴۳	پاسخنامه
۲۴۵	سوالات کنکور سراسری سال ۹۸ الی ۱۴۰۴
۴۵۷	منابع

پیشگفتار:

سابقه تدریس نگارنده در درس روش تحقیق به سال ۱۳۸۰ و آموزشگاه بسیج دانشجویی دانشگاه تهران (کوثر) باز می‌گردد. در آن سال تدریس روش تحقیق برای متقاضیان کارشناسی ارشد رشته پژوهش علوم اجتماعی به‌عهده گرفتم. در سال‌های بعد جزوات آموزشی به نام «گام به گام محقق شویم» را برای سازمان عقیدتی سیاسی نیروی انتظامی تألیف کردم و بالاخره تدریس روش تحقیق را از سال ۱۳۸۵ تا ۱۳۸۹ در کنار تدریس دروس نظریه‌ها و حوزه‌های جامعه‌شناسی در آموزشگاه آزاد علمی ماهان ادامه دادم. با تستی شدن آزمون دکترا، تدریس در مقطع دکترا دیگر چندان منطقی نبود زیرا پرسش‌های باز و تحلیلی جای خود را به پرسش‌های چهار گزینه‌ای داده بود. از آنجا که امکان تدریس در دانشگاه هم فراهم نشد و خاک خوردن مطالبی که در طول چند سال تهیه کرده بودم درست نبود، در نهایت تصمیم به انتشار آن در قالب یک کتاب کمک آموزشی گرفتم. هر چند در تهیه این کتاب حداقل پنجاه کتاب روش تحقیق بطور کامل و ده‌ها کتاب و جزوه هم به‌طور گذار مطالعه شده است اما شاید بتوان گفت حدود نیمی از مطالب کتابی که در دست دارید از چند منبع بسیار خوبی است که در انتهای کتاب فهرست شده است. امیدوارم این اثر بتواند ضمن کمک به داوطلبان کنکور مقطع دکترا و همچنین کارشناسی ارشد، ترس و وحشت درس روش تحقیق را برای خوانندگان به علاقه و دوستی تبدیل کند.

فصل اول

آمار، احتمالات و نمونه‌گیری

آنچه که در این فصل می‌خوانیم

- ✓ گفتار اول: آمار
- ✓ گفتار دوم: احتمالات
- ✓ گفتار سوم: نمونه‌گیری

گفتار اول

آمار

آمار: (Statistics)

آمار به ما کمک می‌کند تا داده‌های فراوان را در چند عدد خلاصه کنیم. برخی نیز گفته‌اند آمار، علم گردآوری، تنظیم، نشان دادن و تحلیل داده‌های عددی مربوط به جامعه است. آمار توصیف پدیده‌ها به شکل کمی و عددی است که ابتدا محقق مسأله‌ای را مطرح می‌کند و سپس یک پاسخ حدسی به نام فرضیه را به این مسأله و سوال می‌دهد. بعد فرضیه‌آزمایی می‌کند که در این مرحله به گردآوری اطلاعات پرداخته می‌شود. تکنیک‌های مختلفی برای گردآوری اطلاعات هست و سپس به کمک آمار اطلاعات به دست آمده را طبقه‌بندی می‌کنیم.

۱- **آمار توصیفی: (Descriptive Statistics)** به توصیف یافته‌ها یا اطلاعات تهیه شده از جوامع آماری می‌پردازد و به صورت جدول توزیع فراوانی‌ها، نمودارها، شاخص‌های گرایش به مرکز یعنی نما، نیمساز و میانگین و شاخص‌های پراکندگی یعنی واریانس، انحراف معیار و دامنه قابل مشاهده است.

۲- **آمار استنباطی: (Inferential Statistics)** به تعمیم یافته‌ها از نمونه به جامعه می‌پردازد.

- آمار استنباطی، برآورد خصوصیتی از جمعیت بر پایه خصوصیتی در نمونه است.

- آمار استنباطی، آزمون فرضیه درباره رابطه بین دو متغیر در جمعیت بر پایه نمونه است.

فرضیه‌هایی که ما داریم یا در کل جمعیت مورد آزمایش و استفاده قرار می‌گیرد و یا در یک نمونه‌ای از جمعیت؛ بنابراین تحقیق دو نوع است:

(۱) سرشماری یا همه‌پرسی یا تمام‌شماری: محقق به تمام افراد رجوع می‌کند و از یک یک آنها اطلاعات می‌گیرد و به کمک آمار توصیفی به تجزیه و تحلیل می‌پردازد.

(۲) نمونه‌گیری: در این شیوه از کل افراد نمونه‌ای انتخاب می‌کنیم. به کمک آمار توصیفی نمونه را توصیف می‌کنیم و بعد به کمک آمار استنباطی نتیجه‌ای که از نمونه به دست آوردیم را به کل افراد تعمیم می‌دهیم؛ بنابراین هنگامی که آمار استنباطی نیاز داریم که تحقیق از نوع نمونه‌گیری است.

تفاوت آمار توصیفی با استنباطی:

در آمار توصیفی محقق از طریق استفاده از موازین مشخص آماری، به توصیف و توجیه ویژگی‌های موردنظر می‌پردازد، بدون آنکه قصد تعمیم نتایج به جوامع بزرگتر را داشته باشد ولی در آمار استنباطی با استفاده از آزمون‌های آماری و برآوردها، نتایج نمونه را به کل تعمیم می‌دهیم.

جمعیت آماری: کلیه اجزا و عناصر که حداقل در یک صفت مشترک باشند یا به عبارت دیگر مجموعه‌ای از افراد جامعه که در یک یا چند صفت مشترک باشند.

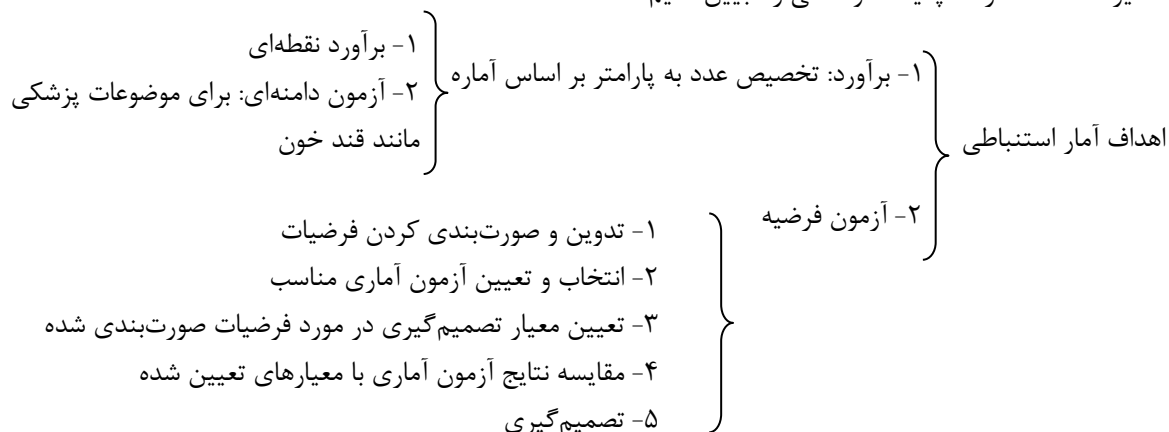
جمعیت نامحدود: جمعیتی است که قابل شمارش نیستند یا یک یک آنها را برداریم به انتها نمی‌رسیم. مثال: برگهای درختان کره زمین، جمعیت ستارگان آسمان، جمعیت ایران که در واقع جمعیتی است که هر قدر می‌شماریم و جلوتر می‌رویم از آن طرف اضافه می‌شود.

جمعیت محدود: جمعیتی است که قابل شمارش است. مثال: دانشجویان دانشکده علوم اجتماعی علامه یا درختان خیابان ولی‌عصر.

نمونه: مجموعه نسبتاً کوچک افراد یا اعضای جمعیت که به شکل خاص انتخاب می‌شوند.

عنصر یا واحد آماری: به هر یک از اعضای جمعیت، واحد آماری گفته می‌شود.

فرضیه در تحقیق و فرضیه آماری: اصل صرفه‌جویی در آمار مطرح است، یعنی با کمترین منابع حداکثر نتایج را به‌دست آوریم. به‌همین خاطر به نمونه‌گیری می‌پردازیم یا سرشماری را هر چند سال یکبار انجام می‌دهیم. در همین راستا برای تبیین پدیده خودکشی نمی‌توان ۳۰۰ متغیر را مطرح کرد و آنها را در امر خودکشی موثر دانست. بلکه باید بتوانیم حداکثر با ۴ یا ۵ متغیر ۸۰ تا ۹۰ درصد پدیده خودکشی را تبیین کنیم.



فرضیه آماری:

فرضیه پاسخ حدسی و موقتی، اما معقول و اندیشیده به مساله تحقیق است. هر تحقیق فرضیه ندارد و ممکن است پرسش داشته باشد. برای مثال: ازدواج در فلان روستا چگونه است؟ مردم در مقابل خرافات چگونه عمل می کنند؟ تحقیقات توصیفی فرضیه ندارند. اما بعضی تحقیقات، فرضیه آزمایی دارند که در آنها به دنبال رابطه بین متغیرها می باشیم که تحقیقات استنباطی و تحلیلی از این جمله اند. مثال: آیا بین خانم ها و آقایان دانشجو از لحاظ میزان مطالعه فرقی است؟

- با کلمات نوشته نمی شود، بلکه با نمادها و سمبل های آماری و ریاضی نوشته می شود.

- فرضیه آماری بر حسب پارامترها نوشته می شود، نه آمارها.

- فرضیه آماری از سه جنبه باید مورد توجه باشد:

الف) شکل منطقی ب) شکل جهت ج) شکل مقایسه ای

۱. شکل منطقی:

هر فرضیه آماری وقتی که نوشته می شود بلافاصله فرضیه ای خنثی نیز با آن خود به خود نوشته می شود. یعنی هر فرضیه ای همزاد خود را یا نقیض خود را به همراه دارد. به فرضیه خنثی یا نفی کننده یا پوچ اصطلاحاً فرضیه صفر می گویند و آن را با H_0 نشان می دهند. نظر محقق در قالب فرضیه ای عنوان می شود که به آن فرضیه مقابل صفر یا مخالف صفر یا جانشین صفر گفته می شود و آن را با H_1 نمایش می دهند.

H_1 : بین جنس دانشجو و میزان مطالعه رابطه وجود دارد.

H_0 : بین جنس دانشجو و میزان مطالعه رابطه وجود ندارد. در مسایل آماری متفاوت درباره جمعیت بر حسب H_0 صورت می گیرد. مثلاً می گوئیم نمی توان H_0 را رد کرد یا در این تحقیق می توان H_0 را رد کرد. اگر متغیر از جنس اسمی باشد آن را با کلمات می نویسیم.

بین جنس دانشجو و میزان مطالعه رابطه وجود ندارد. ($\rho_{xy} = 0$).

بین جنس دانشجو و میزان مطالعه رابطه وجود دارد ($\rho_{xy} \neq 0$)

همبستگی پیرسن است که همیشه بر حسب پارامترها نوشته می شود:

مثال: بین میانگین هوش دانشجویان علوم اجتماعی و میانگین هوش دانشجویان علوم طبیعی تفاوت وجود دارد.

$$H_1 = \mu_s \neq \mu_N \quad H_0 = \mu_s = \mu_N$$

۲. شکل جهتی: فرضیه آماری یا بدون جهت است یا جهت دار فرضیه جهت دار فرضیه ای است که سمت و سوی حرکت متغیرها را نیز تعیین می کند و فرضیه بدون جهت به سمت و سوی حرکت متغیرها اشاره ای ندارد.

مثال برای جهت دار: افراد مسن حافظه ضعیف تری از افراد جوان دارند. یعنی بین سن و حافظه رابطه معکوس یا منفی برقرار است. سن $\uparrow$ حافظه $\downarrow$

بین میزان مطالعه دانشجو و میزان حضور او در کلاس رابطه مثبت یا مستقیم وجود دارد. میزان مطالعه $\uparrow$ میزان حضور در کلاس $\downarrow$

مثال برای فرضیه بدون جهت دار: بین سن و حافظه رابطه وجود دارد یا بین میزان مطالعه و میزان حضور در کلاس رابطه است. علامت فرضیه جهت دار آن است که در آن لفظ رابطه مثبت یا منفی و یا لفظ بیشتر و یا کمتر آمده است. مثال، دختر بچه ها علاقه بیشتری به عروسک بازی دارند تا پسر بچه ها میانگین هوش کارگران بیشتر از میانگین هوش کشاورزان است.

۳. شکل مقایسه ای: از جنبه مقایسه ای فرضیه آماری به دو دسته تقسیم می شود:

الف) مقایسه نمونه با جمعیت ب) مقایسه دو یا چند نمونه با یکدیگر

مثال برای الف: افراد مجرد هوش بیشتری از متوسط هوش جمعیت دارند.

مثال برای ب: مهارت شغلی کارمندان مرد بیشتر از مهارت شغلی کارمندان زن است که باید به صورت پارامترهای آماری نوشته شود:

$$H_1 = \mu_{\text{مردان}} > \mu_{\text{زنان}} \quad H_0 = \mu_{\text{مردان}} = \mu_{\text{زنان}}$$

هر فرضیه‌ای که شکل منطقی، جهت‌ی و مقایسه‌ای دارد:

$$\rho_{xy} = 0 = \text{فعل مثبت}$$

$$\rho_{xy} \neq 0 = \text{فعل منفی}$$

H_1 : معمولاً فعل مثبت دارد ولی H_0 معمولاً فعل منفی دارد.

(۱) متوسط هوش کودکان دبستان ایثار بیشتر از متوسط کودکان مدارس تهران است. H_1 - جهت‌دار - مقایسه نمونه با جمعیت

(۲) بین میزان تحصیلات افراد نمونه با میزان تحصیلات جمعیت که این نمونه از آن گرفته شده است، تفاوت وجود دارد. H_1 - بدون جهت - مقایسه نمونه با جمعیت

(۳) بین شدت علاقه دانشجویان زن و مرد تفاوت وجود دارد. H_1 - بدون جهت - دو نمونه با هم

(۴) بین متوسط IQ مردان دانشجوی و متوسط IQ جمعیت مردان اختلاف وجود دارد. H_1 - بدون جهت - نمونه با جمعیت

(۵) بین جنس دانشجوی و میزان مشارکت سیاسی او رابطه وجود دارد. H_1 - بدون جهت - دو نمونه با هم

همان‌طور که گفتیم آمار استنباطی برداشت‌های حدسی از جمعیت براساس نمونه است.

اشتباه در تصمیم‌گیری‌های استنباطی: تصمیم استنباطی براساس اطلاعاتی که از نمونه به دست می‌آید صورت می‌گیرد، اما اینکه تصمیم‌گیری استنباطی یا حدسی ما در مورد جمعیت درست یا غلط است تنها وقتی معلوم می‌شود که از وضعیت واقعی جمعیت با اطلاع باشیم. مثلاً فرض کنید که اطلاعات حاصل از نمونه ما را به این تصمیم استنباطی برساند که بگوییم در جمعیت دانشجویان، دانشجویان زن از دانشجویان مرد بیشتر مطالعه می‌کنند. اما اگر درون جمعیت آماری عکس این مطلب درست باشد یعنی مردان دانشجوی از زنان دانشجوی بیشتر مطالعه کنند، آنگاه تصمیم استنباطی ما در مورد پذیرش فرضیه خود در مورد جمعیت اشتباه خواهد بود. البته معمولاً پژوهشگر از وضعیت واقعی جمعیت آماری اطلاعی ندارد و اگر اطلاع داشت دیگر نیازی به تحقیق کردن نبود.

شاخص‌های تمرکز و پراکندگی:

برای خلاصه کردن داده‌های یک سری مشاهده دو گونه اندازه‌گیری وجود دارد: اندازه‌های گرایش به مرکز و تغییرپذیری (پراکندگی).

۱. اندازه‌های گرایش به مرکز:

هر مجموعه داده‌ای چند شاخص گرایش به مرکز دارد. مهمترین اندازه‌های گرایش به مرکز که نقطه تمرکز توزیع صفت را مشخص می‌کنند، عبارتند از: نما، میانگین و میانه. چنین عددی به یک معنا با یافتن مقداری در نقطه میانی داده‌ها، آنها را خلاصه می‌کند.

نما: مقداری از متغیر است که تعداد مشاهدات آن حداکثر است. برای مثال در یک کلاس ۲۰ نفره که ۷ نفر نمره ۱۵، ۳ نفر نمره ۱۷، ۵ نفر نمره ۱۲، ۲ نفر نمره ۸ و ۳ نفر نمره ۶ گرفته باشند، نما نمره ۱۵ است، زیرا بیشترین فراوانی را دارد.

میانگین: مهمترین شاخص گرایش به مرکز، مقدار متوسط (Average) یا میانگین (Mean) یعنی $\bar{x}$ است. برای محاسبه آنها تمام مقادیر داده‌ها را با هم جمع می‌کنیم و سپس حاصل را بر تعداد کل مقادیر یعنی n تقسیم می‌کنیم.
مثال: شمار سالهایی که شش بیمار پس از تشخیص ابتلا به سرطان زنده مانده‌اند به شرح زیر است:

۲/۸ ۱/۴ ۴/۶ ۲/۱ ۳/۲ ۱/۷

بنابراین میانگین مدت زنده ماندن (به سال) این بیماران عبارت است از:

$$\bar{x} = \frac{1}{6} (1/7 + 3/2 + 2/1 + 4/6 + 1/4 + 2/8) = \frac{15/8}{6} = 2/63$$

یعنی اینک:

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

که $x_6 = 2/8, \dots, x_7 = 3/2, x_1 = 1/7, n = 6$.

حرف یونانی سیگما (Σ) به این معنی است: «آنچه را که در پی می‌آید جمع کن»؛
i اندیسی است که در این مثال مقادیری بین $i = 1$ تا $i = n$ دارد.

$$\sum_{i=1}^6 x_i^2 = (1/7)^2 + (3/2)^2 + (2/1)^2 + (4/6)^2 + (1/4)^2 + (2/8)^2$$

$$= 2/49 + 9/4 + 4/1 + 16/9 + 1/16 + 1/16 = 48/5$$

$$\sum_{i=1}^i c = c + c + \dots + c \quad (n) \quad \Rightarrow nc$$

$$\sum_{i=1}^n 1 = n$$

میانگین وزنی: (Weighted Average) میانگین وزنی را هنگامی به کار می‌بریم که بخواهیم میانگین تقریبی را پیدا کنیم.

در مثال زیر میانگین وزنی چهار مقدار x_1 تا x_4 را محاسبه کرده‌ایم:

۴	۳	۲	۱	j
۱	۶	۱	۲	w_j
۴	۲	۴	۳	x_j

میانگین وزنی مقادیر x_j ، با ضرایب وزنی w_j به شکل زیر محاسبه می‌شود:

$$= \frac{\sum_{j=1}^f w_j x_j}{\sum_{j=1}^f w_j} = \frac{(2 \times 3) + (1 \times 4) + (6 \times 2) + (1 \times 4)}{2 + 1 + 6 + 1} = \frac{26}{10} = 2.6$$

میانگین وزنی ۲/۶

یعنی میانگین وزنی عبارت است از: مجموع وزنی تقسیم بر مجموع وزنها. (اگر هر $w_j = 1$ باشد، میانگین وزنی همان میانگین معمولی می‌شود.)

داده‌های جدول زیر را در نظر بگیرید:

طبقه (x_j)	فراوانی (f_j)	$f_j x_j$
۳	۵	۱۵
۴	۱۲	۴۸
۵	۲	۱۰
۶	۱	۶
جمع	$\sum_{j=1}^f f_j = 20$	$\sum_{j=1}^f f_j x_j = 79$

میانگین وزنی داده‌های فوق، با گذاشتن فراوانی هر عدد به جای وزن آن، عبارت است از:

$$\bar{x} = \frac{\sum_{j=1}^f f_j x_j}{\sum_{j=1}^f f_j} = \frac{79}{20} = 3.95$$

راه دیگر آن است که همه ۲۰ داده را با هم جمع و تقسیم بر ۲۰ کنیم.

توجه کنید که چگونه در میانگین وزنی اهمیت نسبی هر یک از مقادیر به حساب می‌آید و نیز اینکه مقدار ۳/۹۵ مقداری «نوعی» برای مجموعه داده‌ها محسوب می‌شود، زیرا شاخص خوبی برای گرایش به مرکز است.

میانه: شاخص دیگر گرایش به مرکز میانه (Median) است، یعنی مقدار میانی مجموعه‌ای از داده‌ها که از کوچک به بزرگ یا به عکس مرتب شده‌اند. مثال زیر مفهوم میانه را روشن می‌سازد.

تعدادی داده داریم که عبارتند از: ۹۵، ۸۶، ۷۸، ۹۰، ۶۲، ۷۳ و ۸۹. این داده‌ها را از کوچک به بزرگ مرتب می‌کنیم:

۶۲ ۷۳ ۷۸ ۷۶ ۸۹ ۶۰ ۹۵

میانه، عدد ۷۶ یعنی مقدار وسطی است.

اگر مثلاً به جای عدد ۶۲ عدد ۵ را بگذاریم میانه هیچ تغییری نمی‌کند اما میانگین کاهش می‌یابد.

اگر تعداد داده‌ها زوج باشد، میانگین دو مقدار وسط را به جای میانه می‌گیریم. مثلاً میانه داده‌های ۸۹، ۸۶، ۷۸، ۷۳، ۷۳، ۶۲، ۷۳، ۷۸، ۸۶، ۸۹ و ۹۰ عبارت است از:

$$\frac{78 + 86}{2} = 82$$

میانۀ داده‌های طبقه‌بندی شده

ط. ب. قه	f_j
$۹۴/۵ - (۹۹/۵)$	۳
$۸۹/۵ - (۹۴/۵)$	۴
$۸۴/۵ - (۸۹/۵)$	۱۲
$۷۹/۵ - (۸۴/۵)$	۱۰
$۷۴/۵ - (۷۹/۵)$	۹
$۶۹/۵ - (۷۴/۵)$	۶
$۶۴/۵ - (۶۹/۵)$	۲
$۵۹/۵ - (۶۴/۵)$	۴
جمع	$n = ۵۰$

میانۀ داده‌های طبقه‌بندی شده را قبلاً نشان داده‌ایم. مقدار میانی $\frac{n}{2} = ۲۵$ است. (اگر n فرد باشد این مقدار عدد صحیح نخواهد بود.) در اینجا ۲۱ نمره کمتر از $۷۹/۵$ و ۳۱ نمره کمتر از $۸۴/۵$ است؛ بنابراین میانۀ در طبقه $(۸۴/۵ - ۷۹/۵)$ قرار دارد. در واقع $\frac{۴}{۱۰}$ این فاصله ۵ نمره‌ای مورد نیاز است؛ بنابراین:

$$\text{میانۀ داده‌های طبقه‌بندی شده} = ۷۹/۵ + \frac{۴}{۱۰}(۵) = ۸۱/۵$$

این مقدار را با میانگین $۸۰/۶$ که در صفحه قبل به دست آوردیم مقایسه کنید. در کل:

$$\text{میانۀ داده‌های طبقه‌بندی شده} = L + \frac{\left(\frac{۴}{۲}n - f_c\right)}{f_j}W$$

که:

L : حد پایینی طبقه‌ای است که میانۀ در آن قرار دارد.

f_j : فراوانی آن طبقه

W : فاصله آن طبقه

f_c : فراوانی تراکمی پایین‌تر از L است.

$\left(\frac{۱}{۲}n - f_c\right)$ تعداد مقادیر کمتر از میانۀ در طبقه حاوی میانۀ است.

۲. اندازه‌های تغییرپذیری:

یکی از اندازه‌های ساده و ابتدایی تغییرپذیری، دامنه Range داده‌هاست که از تفاضل کمترین و بیشترین مقدار مجموعه داده‌ها به دست می‌آید. این شاخص به مقادیر کرانی Extreme بسیار حساس است.

اندازه دیگر، میانگین مجذور انحراف‌ها است و از فرمول $\frac{1}{n} = \sum_{i=1}^n (x_i - \bar{x})^2$ به دست می‌آید که n تعداد داده‌ها و $\bar{x}$ میانگین آنهاست. اگر همه انحراف‌های از میانگین را با هم جمع کنیم، حاصل آن همیشه مساوی صفر است و نمی‌توان از آن به عنوان شاخص تغییرپذیری استفاده کرد. به همین دلیل هر یک از انحراف‌ها را قبل از جمع کردن به توان دوم می‌رسانیم و سپس میانگین آنها را محاسبه می‌کنیم. اگر برخی از مقادیر، بسیار دور از میانگین باشند مجذور انحراف‌ها بزرگ خواهد شد.

واریانس: مهمترین اندازه تغییرپذیری، یعنی واریانس (Variance) از تقسیم مجموع مجذور انحرافها بر $(n-1)$ به جای n به دست می آید. فرمول واریانس که آن را با s^2 نشان می دهیم عبارت است از:

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

با تقسیم کردن مجموع مجذور انحرافها بر $(n-1)$ ، برآوردی بدون تورش از واریانس به دست می آید.

در داده های طبقه بندی شده، فرض بر این است که همه مقادیر هر فاصله در مرکز آن فاصله قرار دارند؛ بنابراین:

$$s^2 = \frac{1}{n-1} \sum_{j=1}^k (x_j - \bar{x})^2 f_j$$

که فاصله های طبقه ای را با $j=1$ تا $j=k$ شماره گذاری می کنیم.

خواص واریانس: ۱. اگر یک مقدار ثابت (مثلا عدد ۳) را از تمام مقادیر مورد بررسی کم یا به همه آنها اضافه کنیم، واریانس هیچ تغییری نمی کند. ۲. اگر همه مقادیر متغیر را بر یک عدد ثابت تقسیم کنیم، واریانس توزیع جدید ۳ به توان دو برابر از واریانس قدیم کوچکتر می شود. مثلا اگر همه اعداد یک مجموعه را بر عدد ۳ تقسیم کنیم، واریانس آنها بر ۹ تقسیم می شود. ۳. اگر همه مقادیر را در یک عدد ثابت ضرب کنیم، واریانس جدید ۳ به توان دو برابر واریانس قدیم می شود. یعنی واریانس جدید ۹ برابر واریانس قدیم می شود.

انحراف معیار: یکی دیگر از اندازه های تغییرپذیری انحراف معیار (Standard Deviation) است که با فرمول زیر محاسبه می شود:

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

که $(n-1)$ را درجه آزادی می نامند.

مزیت انحراف معیار بر واریانس آن است که واحد اندازه گیری s با واحد داده ها یکی است چون جذر گرفتن از واریانس، اثر به توان دو رساندن داده ها را خنثی می کند.

مثال: انحراف معیار اعداد ۱، ۴، ۵ و ۲ را حساب کنید.

$(x_i - \bar{x})^2$	$x_i - \bar{x}$	x_i	i
۹	$1-4=-3$	۱	۱
۱۶	$8-4=4$	۸	۲
۰	$4-4=0$	۴	۳
۱	$5-4=1$	۵	۴
۴	$2-4=-2$	۲	۵
۳۰	۰	۲۰	جمع

$$\bar{x} = \frac{20}{5} = 4 \text{ و } s = \sqrt{\frac{30}{4}} = \sqrt{7.5} = 2.74$$

انحراف معیار، انحرافی «نوعی» یا متوسط از میانگین است. حتما توجه کرده اید که مجموع انحراف های از میانگین همواره مساوی صفر است.

معایب انحراف معیار: ۱. انحراف معیار هیچ وقت کمتر از صفر نمی شود. ۲. حد بالای انحراف معیار باز است. ۳. در توزیع های دارای چولگی زیاد کارایی آن کم است.

انحراف چارکی: به تفاضل چارک سوم از چارک اول اشاره دارد و یکی از اندازه‌های تغییرپذیری است که با این فرمول

$$\text{مشخص می‌شود: } \text{انحراف چارکی} = \frac{Q_3 - Q_1}{2}$$

چارک اول نشان دهنده ۲۵٪ داده‌ها، چارک دوم که همان میانه است نشان دهنده ۵۰٪ داده‌ها و چارک سوم نشان دهنده ۷۵٪ داده‌ها است. انحراف چارکی همان کار انحراف استاندارد را انجام می‌دهد با این تفاوت که برای زمانی بکار می‌رود که ما از میانه Md برای برآورد استفاده می‌کنیم. (انحراف استاندارد برای زمانی است که ما از میانگین استفاده می‌کنیم).
تذکر: وقتی منحنی نرمال چولگی داشته باشد بهترین آماره میانه Md است و وقتی منحنی نرمال باشد، میانگین بهترین آماره است.

نکته: ۱. نما برای متغیرهای اسمی، میانه برای متغیرهای ترتیبی و میانگین برای متغیرهای فاصله‌ای و نسبی به کار برده می‌شود. ۲. مجموع انحراف نمره‌ها از میانگین، همیشه برابر صفر است، ولی مجموع انحراف نمره‌ها از میانه یا نما همیشه برابر با صفر نیست. ۳. ثبات میانگین بیشتر از نما و میانه است. ۴. با میانگین عملیات ریاضی بیشتری می‌توان انجام داد. ۵. در توزیع‌های چوله بهترین معیار مرکزگرایی میانه است ولی در توزیع‌های نرمال میانگین است.
نکته: میانگین Z (نمرات استاندارد) همیشه صفر و انحراف معیار آن همیشه یک است.
ضریب تغییرات:

$$\text{نسبت انحراف معیار بر میانگین را ضریب تغییرات می‌نامند و آن را چنین نمایش می‌دهند: } CV = \frac{S}{\bar{X}}$$

ضریب تغییرات برای از بین بردن واحد اندازه‌گیری‌های متفاوت بکار می‌رود. برای مثال فرض کنید در یک کارخانه میانگین و انحراف معیار مزد (X) و ساعت کار (y) هفتگی کارگران به شرح زیر است:

$$\begin{aligned} \bar{y} &= 50 \text{ ساعت} & S_y &= 12/5 \text{ ساعت} \\ \bar{X} &= 180 \text{ دلار} & S_x &= 36 \text{ دلار} \end{aligned}$$

در نتیجه ضریب تغییرات مزد و ساعت کار چنین محاسبه می‌شود:

$$CV_y = \frac{12/5}{50} = 0/25 \text{ ساعت کار}$$

$$CV_x = \frac{36}{180} = 0/2 \text{ مزد}$$

از آنجا که هرچه مقدار ضریب تغییرات به صفر نزدیکتر باشد، داده‌ها همگون‌تر است بنابراین مقایسه مان‌شان می‌دهد که ضریب تغییرات ساعت کار بیشتر از ضریب تغییرات مزد است در حالی که اگر انحراف معیارها را با هم مقایسه می‌کردیم باید عکس آن را نتیجه می‌گرفتیم که کار درستی نبود.

نمرات استاندارد شده: در آمار به نمره‌ای که یک آزمودنی کسب می‌کند نمره خام گفته می‌شود. برای پی بردن به ارزش و رتبه هر نمره خام باید آن را استاندارد کنیم. یکی از روش‌های استاندارد کردن نمره‌های خام تبدیل نمره خام (X) به نمره استاندارد (Z) براساس فرمول زیر است:

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

مثال: در یک آزمایش میانگین و انحراف معیار نمره‌های یک گروه به ترتیب ۳۲ و ۶ است. نمره خام دو آزمودنی A و B نیز به ترتیب ۲۶ و ۴۴ است. نمره استاندارد هر دو را محاسبه کنید.

$$Z_A = \frac{26 - 32}{6} = -1$$

$$Z_B = \frac{44 - 32}{6} = 2$$

گفتار دوم

احتمالات

عدم قطعیت جزء لاینفکی از زندگی روزمره است. رویدادهای بسیار معدودی را می‌توان به قطع و یقین پیش‌بینی کرد؛ باقی رویدادهای قابل پیش‌بینی قطعی نیست، مثلاً واکنش بیمار به نوعی درمان خاص، طول دوره بیماری، معنای نتایج آزمونی تشخیصی و از این قبیل. احتمال، شاخص یا اندازه عدم قطعیت است. احتمال، عدم قطعیت را به زبان کمیت و بر روی مقیاسی بین ۰ و ۱ بیان می‌کند که ۰ به معنی مطلقاً ناممکن و ۱ به معنی صد در صد قطعی است. احتمال را می‌توان به صورت درصد یا تعداد شانس‌ها در ۱۰۰ بیان کرد که در هر صورت مقدار آن بین ۰٪ به معنی مطلقاً ناممکن و ۱۰۰٪ به معنی کاملاً قطعی متغیر است.

لغات و عبارات بسیاری هست که بیان‌کننده نوعی ارزیابی ذهنی از احتمال است. لغات و عباراتی چون «گاهی»، «قطعاً»، «بعید»، «هرگز»، «احتمالاً»، «نامعقول»، «همیشه» و از این قبیل که همگی حاوی نظری راجع به احتمال وقوع رویدادی خاص است. در تحقیقی از برخی متخصصان علوم بهداشتی خواسته شد تا کلمات فوق را به صورت عددی تعبیر کنند. ارزش‌های عددی که آنان برای کلمات قائل می‌شدند کاملاً متغیر بود و این نشان می‌داد که بهتر است هنگامی که می‌خواهند مفهوم احتمال را به دیگران انتقال دهند، از مقادیر عددی استفاده کنند.

همان‌طور که بازی‌های مبتنی بر شانس لااقل سه هزار سال قدمت دارند مفهوم احتمال نیز جدید نیست. با این همه، تا پیش از قرن ۱۷ میلادی، مفهوم احتمال به صورت مستقل یا به زبان ریاضی بیان نشده بود. در اوایل قرن ۱۸، دو موآور که از طریق محاسبه شانس برد برای قماربازانی که در کافه‌ها قمار می‌کردند گذران زندگی می‌کرد، فرصت یافت تا برخی از راه‌حل‌های نظری خود را با استفاده از نتایج بازی‌های تاس و ورق که در دور و برش انجام می‌شد بازبینی نماید.

در بازی‌های شانس، با ملاحظه مساله و برقرار دانستن مفروضات خاصی می‌توان برآوردی نظری یا پیشینی از احتمال به دست آورد. به عنوان مثال، در شیر یا خط کردن با سکه، فرض بر این است که سکه تقلبی نیست و احتمال پیشامد شیر و خط هر دو مساوی است چون در این مورد خاص تنها دو پیشامد متصور است احتمال شیر آمدن مساوی ۰/۵ است. برای اینکه درستی این حکم را در مورد یک سکه خاص بررسی کنیم، باید بارها و بارها با آن سکه شیر یا خط کنیم و تعداد شیرها را بشماریم. نسبت دفعات آمدن شیر، برآوردی از احتمال شیر به دست می‌دهد، به این ترتیب که:

$$\Pr(\text{شیر}) = \frac{\text{تعداد دفعات آمدن شیر}}{\text{تعداد کل پرتاب سکه}}$$

اگر تعداد دفعات پرتاب سکه کم باشد، این برآورد بسیار متغیر و بی‌ثبات است، اما هر چه تعداد دفعات شیر یا خط کردن افزایش یابد این برآورد بیشتر به سمت 0.5 میل می‌کند، البته اگر سکه بی‌اشکال باشد. در مورد سکه دو پیشامد ممکن وجود دارد که این دو پی‌آمد جامعند، چون مجموع احتمال این دو مساوی $1/0$ است و نیز ناسازگارند، چون در یک بار شیر یا خط کردن وقوع هر دو پیشامد با هم ممکن نیست.

به عنوان دومین مثال، پیشامدهای ممکن آزمایش با 52 ورق بازی را در نظر بگیرید. اگر یک ورق از بین یک دسته ورق بیرون بکشیم، 52 پیشامد ممکن متصور است. احتمال بیرون آمدن ده خشت $\frac{1}{52}$ است. احتمال آمدن یک ورق با خال خشت $\frac{13}{52}$ است، چون 13 ورق با خال خشت وجود دارد و احتمال همه آنها نیز یکی است. آمدن خال خشت رویدادی است که از 13 پیشامد تشکیل شده است. سه رویداد دیگر عبارتند از: آمدن خال دل، خال پیک و خال گشنیز. این چهار رویداد با هم ناسازگارند چون هر ورق تنها می‌تواند از یک خال باشد و نیز جامعند چون مجموع احتمال این چهار رویداد مساوی $1/0$ است. بعضی از رویدادها ناسازگار نیستند، مثلاً رویداد «آمدن ورق ده» و «آمدن خال خشت»، چون یک ورق می‌تواند هم این باشد و هم آن. احتمال این قبیل حالات را در بخش بعدی یعنی بخش ترکیب احتمال‌ها بررسی خواهیم کرد. در کل، رویدادها نتیجه تجارب تصادفی است که در موقعیت‌های مختلف با آن مواجه هستیم، از قبیل:

- آزمایش ریختن تاس، که یکی از رویدادهای ممکن «آمدن عدد فرد» است.
- آزمایش با گزینش 500 رأی‌دهنده که در یک نظرسنجی نمونه‌گیری شده‌اند؛ یکی از رویدادهای ممکن «دست کم 300 نفر موافق حزب محافظه‌کار» است.
- آزمایش با 50 دانش‌آموز که سر جلسه امتحان نشسته‌اند؛ رویدادهای ممکن این است که «متوسط نمره افراد بالاتر از 75% باشد».
- آزمایش با 20 بیمار که با داروی جدیدی درمان می‌شوند؛ یکی از رویدادهای ممکن این است که «دست کم ده بیمار کاملاً معالجه شوند».
- دو گروه نمونه مساوی از بیماران سرطانی، یک گروه از طریق عمل جراحی و گروه دیگر از طریق شیمی‌درمانی، تحت درمان قرار گرفته‌اند؛ یکی از رویدادهای ممکن این است که «تعداد بیماران بهبود یافته در گروه شیمی‌درمانی بیشتر باشد».

چون احتمال‌های نظری یا پیشینی در بسیاری از موقعیت‌های این چنینی معلوم نیست، این احتمال‌ها را باید از داده‌ها و با مشاهده پیشامدهای آزمایش‌ها و محاسبه نسبت دفعات وقوع رویدادهای مورد نظر برآورد کرد.

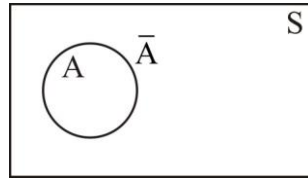
ترکیب احتمال‌ها

کسانی که به حرفه‌های بالینی می‌پردازند معمولاً به بیش از یک رویداد و به ترکیب احتمال‌های آن رویدادها توجه دارند؛ مثلاً، احتمال اینکه زنی که غده‌ای در سینه دارد مبتلا به سرطان باشد با این شرط که سنش از 40 بیشتر باشد و خواهری مبتلا به سرطان داشته باشد.

یکی از راه‌های درک مفهوم احتمال، استفاده از نمودار ون (Venn diagram) است. طبق تعریف فضای نمونه‌ای S مجموعه تمامی پیشامدهای ممکن یک آزمایش است و A مجموعه‌ای از این پیشامدهاست که رویداد موردنظر را تشکیل می‌دهد. در نمودار ون، فضای نمونه (سطح داخل مستطیل) و مجموعه‌ای از پیشامدها که رویداد A را تشکیل می‌دهند (سطح داخل دایره) و رویداد «نه A » یا $\bar{A}$ (سطح بیرون از دایره) که متمم A نامیده می‌شود نشان داده شده است. چون فضای نمونه S مجموعه همه پیشامدهای ممکن است و یکی از این پیشامدها حتماً رخ می‌دهد، $\Pr(S) = 1$ است. بر همین منوال چون A و $\bar{A}$

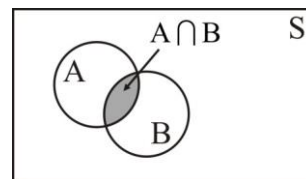
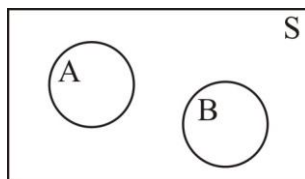
مجموعه فضای نمونه را تشکیل می‌دهد $\Pr(A) + \Pr(\bar{A}) = 1$ است و در نتیجه $\Pr(\bar{A}) = 1 - \Pr(A)$. به این ترتیب رویدادهای A و $\bar{A}$ هم ناسازگار و هم جامعند.

نمودار ون رویداد A و متمم آن



چندین رویداد مرکب را به راحتی می‌توان با نمودار ون نشان داد. اولین آنها رویداد « A یا B » است که به آن اجتماع (Union) رویدادهای A و B می‌گویند و در نظریه مجموعه‌ها آن را با نماد $A \cup B$ نشان می‌دهند. در این رویداد دست کم یکی از رویدادهای A یا B رخ می‌دهد. رویداد مرکب دوم، رویداد « A و B » است که به آن اشتراک (Intersection) رویدادهای A و B می‌گویند و آن را با نماد $A \cap B$ نشان می‌دهند و به معنای آن است که رویدادهای A و B هر دو با هم رخ دهند. نمودارهای زیر، $A \cap B$ را در دو حالت نشان داده است؛ یکی حالتی که اشتراک $A \cap B$ تهی است و دیگری حالتی که اشتراک $A \cap B$ تهی نیست.

رویداد مرکب $A \cap B$



(۲) اشتراک $A \cap B$ تهی است؛ A و B ناسازگارند

(۱) اشتراک $A \cap B$ تهی نیست

در حالت دوم هیچ همپوشانی‌ای بین دو رویداد A و B وجود ندارد. یعنی اینکه این دو رویداد نمی‌توانند همزمان رخ دهند. بر طبق قاعده جمع احتمال‌های مرکب:

$$\Pr(B \mid A) \neq \Pr(A \cup B) = \Pr(A) + \Pr(B) - \Pr(A \cap B)$$

چون در جمع احتمال رویداد A و B ، فصل مشترک آنها دو بار محاسبه می‌آید، احتمال اشتراک A و B را باید از جمع احتمال آن دو کم کرد تا $\Pr(A \cup B)$ به دست آید. اگر رویدادهای A و B ناسازگار باشند جمع احتمال دو رویداد به شکل ساده‌تر زیر درمی‌آید.

$$\Pr(A \cup B) = \Pr(A) + \Pr(B)$$

مفهوم دیگری را که باید توضیح دهیم احتمال شرطی *Conditiona Probability* است. گفتیم که همپوشی رویدادهای A و B ، اشتراک آن دو رویداد است و آن را با $A \cap B$ نشان می‌دهند. احتمال وقوع B به عنوان رویداد دوم بسته به این است که رویداد اول یعنی A رخ داده باشد یا نه. این احتمال یعنی $\Pr(B \mid A)$ یا $\Pr(A \mid B)$ که به شکل $\Pr(A \mid B)$ نوشته می‌شود، نوعی احتمال شرطی است و با استفاده از نمودار (۱) در شکل قبل خواهیم داشت:

$$\Pr(B \mid A) = \frac{\Pr(A \cap B)}{\Pr(A)}$$

$$\Pr(A \mid B) = \frac{\Pr(A \cap B)}{\Pr(B)}$$

در نتیجه به قاعده ضرب (Multiplication Rule) احتمال‌ها می‌رسیم که عبارت است از:

$$\Pr(A \cap B) = \Pr(A|B) \times \Pr(B) = \Pr(B|A) \times \Pr(A)$$

برای فهم بهتر قاعده جمع و ضرب احتمال‌ها به مثال ۵۲ ورق بازی برمی‌گردیم. یک ورق از میان دسته ورق بیرون می‌کشیم و می‌خواهیم بدانیم احتمال اینکه ورق مزبور هم ده و هم خشت باشد چقدر است. اگر رویداد «بیرون آمدن ده» را با A و رویداد «بیرون آمدن خال خشت» را با B نشان دهیم، آنگاه $\Pr(A) = \frac{4}{52}$ (چون ۴ ورق از ۵۲ ورق ده است) و $\Pr(B) = \frac{13}{52}$ (چون ۱۳ ورق از ۵۲ خشت است)، پس $\Pr(A \cap B) = \frac{1}{52}$ و در نتیجه براساس قاعده جمع احتمال‌ها $\frac{16}{52} = 0.31$.

حال فرض کنید دو ورق از میان دسته ورق بیرون کشیده‌ایم و هر دو را بدون اینکه نگاه کنیم در دست نگاه داشته‌ایم. مساله محاسبه احتمال آن است که هر دو ورق ده باشد. اگر A رویداد «ده بودن ورق اول» و B رویداد «ده بودن ورق دوم» باشد، آنگاه $\Pr(A) = \frac{4}{52}$ و $\Pr(B) = \frac{3}{51}$ خواهد بود چون پس از بیرون کشیدن ورق ده اول، تنها سه ورق ده در میان ۵۱ ورق باقی می‌ماند؛ بنابراین:

$$\Pr(A \cap B) = \Pr(B|A) \Pr(A) = \left(\frac{3}{51}\right) \times \left(\frac{4}{52}\right) = 0.005$$

مثال: احتمال اینکه دختر جوانی برنامه تلویزیونی خاصی را نگاه کند 0.4 ، احتمال اینکه دخترعمویش آن برنامه را نگاه کند 0.5 و احتمال اینکه آن دختر جوان برنامه مزبور را نگاه کند با این فرض که دخترعمویش نیز نگاه می‌کند مساوی 0.7 است. احتمال این را بیابید که:

(الف) هم دختر جوان و هم دخترعمویش، هر دو برنامه را نگاه کنند؛

(ب) دختر عمو برنامه را نگاه کند با این شرط که دختر جوان نیز برنامه را نگاه کند؛

(ج) هیچ‌کدام برنامه را نگاه نکنند.

اگر A رویداد «نگاه کردن دختر جوان» و B رویداد «نگاه کردن دختر عمو» باشد، آنگاه $\Pr(A) = 0.4$ و $\Pr(B) = 0.5$ و $\Pr(A|B) = 0.7$.

$$\Pr(A \cap B) = \Pr(A|B) \Pr(B) = 0.7 \times 0.5 = 0.35 \quad (\text{الف})$$

$$\Pr(B|A) = \frac{\Pr(B \cap A)}{\Pr(A)} = \frac{0.35}{0.4} = 0.875 \quad (\text{ب})$$

(ج) احتمال اینکه هم دختر جوان برنامه را نگاه کند و هم دخترعمویش عبارت است از:

$$\Pr(A \cup B) = \Pr(A) + \Pr(B) - \Pr(A \cap B) = 0.4 + 0.5 - 0.35 = 0.55$$

پس احتمال اینکه هیچ‌کدام برنامه را تماشا نکنند $1 - \Pr(A \cup B) = 0.45$ است.

یادآوری: احتمال‌های اولیه در این مثال باید پس از مشاهده رفتار تماشای برنامه در دختر جوان و دخترعمویش تعیین گردد. چنین ارقامی را نمی‌توان بدون دلیل موجهی از خود درآورد.

مثال: در جدول زیر تعداد کودکان اوریون گرفته و اوریون نگرفته و نیز تعداد کودکانی که واکسن اوریون زده‌اند و آنانی که واکسن نزده‌اند آمده است:

مایه‌کوبی شده	اوریون گرفته	اوریون نگرفته	جمع
۱۴	۱۰۵	۱۱۹	
۶۰	۷۱	۱۳۱	
۷۴	۱۷۶	۲۵۰	جمع

رویداد M آن است که «کودک به تصادف انتخاب شده اوریون گرفته باشد» و رویداد V «کودک به تصادف انتخاب شده

مایه‌کوبی شده باشد». پس براساس داده‌های فوق، $\Pr(V) = \frac{119}{250} = 0/476$ و $\Pr(V|M) = \frac{14}{74} = 0/189$.

همچنین $\Pr(V \cap M) = \frac{14}{250} = 0/056$.

راه دیگرش این است که چون $\Pr(M) = \frac{74}{250}$ ، در نتیجه:

$$\Pr(V \cap M) = \Pr(M) \Pr(V|M) = \left(\frac{74}{250}\right) \times \left(\frac{14}{74}\right) = 0/056$$

همچنین $\Pr(V|M) = \frac{14}{119} = 0/118$ ، احتمال مبتلا شدن کودک به اوریون است با این شرط که مایه‌کوبی شده باشد.

رویدادهای مستقل و وابسته

فرض کنید می‌خواهیم بدانیم چقدر احتمال دارد وزن دو کودک که پشت سر هم در یک بیمارستان متولد شده‌اند و دوقلو هم نیستند کمتر از ۲/۵ کیلوگرم باشد. این دو رویداد مستقلند چون رویداد دوم به هیچ وجه ارتباطی به رویداد اول ندارد.

رویداد B هنگامی مستقل از رویداد A دانسته می‌شود که احتمال وقوع B از وقوع یا عدم وقوع A تاثیری نپذیرد یا به عبارت دیگر، احتمال B مساوی با احتمال شرطی B به فرض A باشد؛ یعنی اینکه $\Pr(B) = \Pr(B|A)$. در نتیجه قاعده ضرب $\Pr(A \cap B) = \Pr(A) \Pr(B|A)$ به شکل ساده‌تر زیر درمی‌آید:

$$\Pr(A \cap B) = \Pr(A) \Pr(B)$$

این فرمول در حالتی صادق است که دو رویداد A و B مستقل از هم باشند. همچنین $\Pr(A) = \Pr(A|B)$ خواهد بود چون براساس قاعده دوم ضرب احتمال‌ها $\Pr(A \cap B) = \Pr(B) \Pr(A|B)$.

مثال: A و B دو رویدادند و $\Pr(A) = 0/60$ ، $\Pr(B) = 0/40$ و $\Pr(A \cap B) = 0/10$ است. مقدار $\Pr(A \cup B)$ و $\Pr(A|B)$ را پیدا کنید و معلوم نمایید که دو رویداد A و B مستقلند یا نه.

در اینجا $\Pr(A \cup B) = \Pr(A) + \Pr(B) - \Pr(A \cap B) = 0/60 + 0/40 - 0/10 = 0/90$

و $\Pr(A|B) = \frac{\Pr(A \cap B)}{\Pr(B)} = \frac{0/10}{0/40} = 0/25$ چون $\Pr(A|B) = 0/25 \neq \Pr(A)$ پس رویدادهای A و B مستقل از هم

نیستند.

نمودار درختی برای محاسبه احتمال

نمودار درختی محاسبه احتمال رویدادهای مرکب را راحت تر می کند، به خصوص اگر آزمایش چند مرحله ای باشد. به عنوان مثال، می خواهیم رویدادهای تعداد فرزندان پسر را در سه نسل یک خانواده نشان دهیم و احتمال این رویدادها را نیز محاسبه کنیم. برای ساده تر کردن مساله فرض می کنیم که هیچ کدام از فرزندان دوقلو یا چند قلو نیستند و احتمال تولد پسر ۰/۵ است. (در واقع این احتمال اندکی کمتر از ۰/۵ است.)

در شکل زیر حالات ممکن گسترش خانواده را در سه نسل پی در پی نشان داده ایم. احتمال پسر بودن (B) یا دختر بودن (G) در هر تولد مساوی ۰/۵ است. چون تولد هر کودک مستقل از تولد قبلی است، بر طبق قاعده ضرب احتمال رویدادهای مستقل:

$$\Pr(B \text{ و در پی آن } B) = \Pr(BB) = 0.5 \times 0.5 = 0.25$$

نمودار درختی احتمال

تعداد پسران	$\Pr(\dots)$	پایه شامد	پایه سوم	نسل دوم	نسل اول
۳	$0.5 \times 0.5 \times 0.5 = 0.125$	(B,B,B)	B	۰/۵	
۲	۰/۱۲۵	(B,B,G)	G	۰/۵	
۲	۰/۱۲۵	(B,G,B)	B	۰/۵	
۱	۰/۱۲۵	(B,G,G)	G	۰/۵	
۲	۰/۱۲۵	(G,B,B)	B	۰/۵	
۱	۰/۱۲۵	(G,B,G)	G	۰/۵	
۱	۰/۱۲۵	(G,G,B)	B	۰/۵	
۰	۰/۱۲۵	(G,G,G)	G	۰/۵	

۹

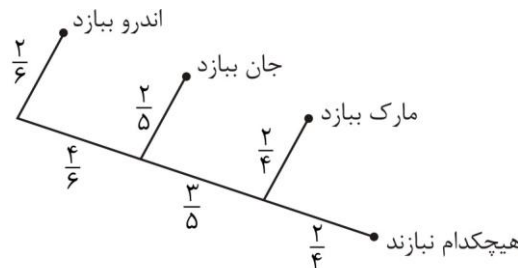
$$\Pr(G \text{ و در پی آن } B) = \Pr(BBG) = 0.5 \times 0.5 \times 0.5 = 0.125$$

در نمودار هشت رویداد ممکن نشان داده شده است که این هشت رویداد هم جامعند و هم ناسازگارند. در نتیجه:

$$\Pr(\text{دو پسر در خانواده}) = \Pr(BBG) + \Pr(BGB) + \Pr(GBB) = 0.125 + 0.125 + 0.125 = 0.375$$

به همین ترتیب $\Pr(\text{یک پسر در خانواده}) = 0.375$ و $\Pr(\text{هیچ پسر در خانواده}) = 0.125$ خواهد بود. در مثال بعدی احتمالها مشروط به وقوع رویدادهای وابسته است یعنی اینکه برای به دست آوردن احتمال رویداد نهایی باز هم باید احتمال شاخه ها را در هم ضرب کرد.

مثال: اندرو، جان و مارک دست به یک ماجراجویی زده‌اند. شش اتومبیل شبیه هم وجود دارد که دو تا از آنها ترمز ندارد. هر کدام از این سه نفر باید یک اتومبیل را به تصادف انتخاب کند و با سرعت زیاد به سمت پرتگاه براند و در زمان مناسب ترمز کند. احتمال‌های (یک نفر ببازد) Pr و (هیچ کدام نبازند) Pr را محاسبه کنید، با این فرض که با سقوط یکی از این سه نفر از دره بازی متوقف شود.



$$Pr(\text{اندرو اتومبیل بدون ترمز انتخاب کند}) = Pr(\text{باختن اندرو}) = \frac{2}{6}$$

$$Pr(\text{اندرو اتومبیل سالمی بردارد و جان اتومبیل بدون ترمز}) = Pr(\text{باختن جان}) = \left(\frac{4}{6}\right)\left(\frac{2}{5}\right) = \frac{4}{15}$$

$$Pr(\text{اندرو و جان اتومبیل سالمی انتخاب کنند و مارک اتومبیل بدون ترمز}) = Pr(\text{باختن مارک}) = \left(\frac{4}{6}\right)\left(\frac{3}{5}\right)\left(\frac{2}{4}\right) = \frac{3}{15}$$

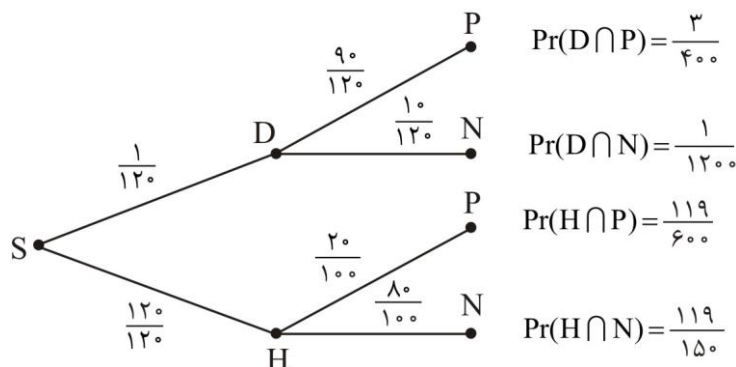
$$Pr(\text{هیچکدام نبازند}) = \left(\frac{4}{6}\right)\left(\frac{3}{5}\right)\left(\frac{2}{4}\right) = \frac{3}{15}$$

همان‌طور که می‌بینید جمع این چهار احتمال مساوی یک است. همچنین احتمال هر شاخه مشروط به وقوع پیش‌آمد قبل است.

مثال: بررسی بیماری‌رانی که در ایالات متحده آمریکا دچار سرطان پروستات شده‌اند نشان می‌دهد که تا پایان سال ۱۹۹۶ از هر ۱۲۰ مرد ۵۰ ساله و بالاتر، ۱ نفر به این بیماری مبتلا شده است. با یک آزمایش ساده خون میزان آنتی‌ژن‌های خاص پروستات اندازه‌گیری می‌شود. در صورت وجود سرطان، مقدار PSA شدیداً افزایش می‌یابد، اما مقدار PSA به طور طبیعی با افزایش سن نیز افزایش می‌یابد.

در یک آزمایش تشخیصی، سطح خاصی برای این آنتی‌ژن تعیین شده که در ۹۰٪ بیماران مقدار PSA از این سطح فراتر می‌رود (مثبت بودن آزمایش) و در ۸۰٪ افراد سالم مقدار PSA از این سطح پایین‌تر است (منفی بودن آزمایش). تعدادی از مردان بالای ۵۰ سال به تصادف برای انجام این آزمایش تشخیصی انتخاب شده‌اند.

(الف) افراد آزمایش شده یا مبتلا به سرطان پروستات هستند (D) یا سالمند (H) و در هر کدام از این دو حالت یا جواب آزمایش مثبت (P) است و یا منفی (N). نمودار درختی پیشامدهای مختلف را رسم کنید. احتمال هر یک از شاخص‌های این نمودار درختی را بنویسید و احتمال هر یک از پیشامدهای ممکن را حساب کنید و در کنار خط هر پیشامد یادداشت نمایید.



(ب) احتمال اینکه جواب آزمون تشخیصی مثبت (P) باشد چقدر است؟

$$\Pr(\text{مثبت شدن آزمایش}) = \Pr(D \cap P) + \Pr(H \cap P) = \frac{3}{400} + \frac{119}{600} = 0/2058$$

(ج) احتمال اینکه فردی که جواب آزمایش مثبت بوده بیمار (H) باشد چقدر است؟

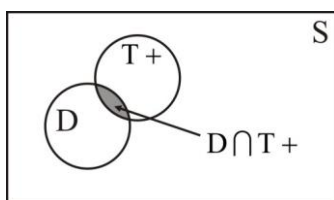
$$\Pr(H|P) = \frac{\Pr(H \cap P)}{\Pr(P)} = \frac{\frac{119}{600}}{0/2058} = 0/9636$$

آزمایش‌های تشخیصی

پزشکان از آزمون‌های تشخیصی برای تعیین بیماری اشخاص استفاده می‌کنند. پزشک غالباً اطلاعاتی در خصوص آزمون تشخیص و برآوردی از میزان شیوع بیماری در جامعه بیماران دارد. مساله این است که:

- احتمال بیمار بودن به شرط مثبت بودن آزمایش چقدر است، یا
- احتمال بیمار نبودن به شرط منفی بودن آزمایش چقدر است.

رویداد وضعیت بیماری و رویداد جواب آزمایش



این وضعیت را در شکل بالا با نمودار ون نشان داده‌ایم. D رویداد «وجود بیماری» و $\bar{D}$ متمم آن یعنی «فقدان بیماری»، T^+ رویداد «مثبت بودن آزمایش» و $\bar{T}$ متمم آن یعنی «منفی بودن یا بی‌نتیجه بودن آزمایش» است و بالاخره S فضای نمونه جمعیت موردنظر است. بنابراین شیوع بیماری عبارت است از:

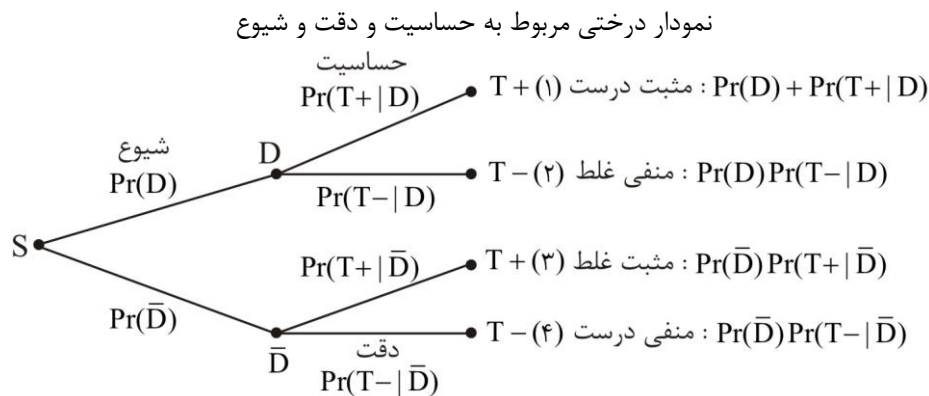
$$\Pr(D) = \frac{\text{تعداد افراد مبتلا به بیماری در گروه تعریف شده}}{\text{تعداد افراد گروه تعریف شده}}$$

که در واقع احتمال آن است که فرد به بیماری مبتلا شود در هنگامی که هیچ اطلاعات دیگری از آن فرد در دسترس نباشد. همچنین:

نسبت افراد بیماری که جواب آزمایششان مثبت بوده است = حساسیت (Sensitivity) آزمایش $\Pr(T^+ | D)$

نسبت افراد سالمی که جواب آزمایششان منفی بوده است = دقت (Specificity) $\Pr(T^- | \bar{D})$ = آزمایش

این مقادیر هرگز از پیش معلوم نیستند و باید آنها را برآورد کرد. شیوع را از طریق مطالعات همه‌گیرشناسی (اپیدمیولوژیکی) و یا گاهی با بررسی فهرستی طولانی از بیماران برآورد می‌کنند. حساسیت و دقت را می‌توان با استفاده از گروه نمونه‌ای از افرادی که وضعیت بیماریشان معلوم و مسجل است به بهترین نحو برآورد کرد. اما برای اینکه وضعیت بیماری آنها بر ما معلوم شود به آزمایش امتحان شده و معتبری نیازمندیم و به ندرت نیز چنین آزمایشی وجود دارد. نبود چنین آزمایشی به علاوه خطای نمونه‌برداری، موجب ایجاد خطا در برآورد حساسیت و دقت می‌شود.



مراحل چهارگانه اجرای آزمون آماری به ترتیب زیر است:

- (۱) فرضیه صفر (H_0) در مورد یکی از پارامترهای جامعه طرح می‌شود.
- (۲) فرضیه مقابل یا فرضیه پژوهش (H_A) پیشنهاد می‌شود. این فرضیه وقتی پذیرفته می‌شود که H_0 رد شود.
- (۳) آماره آزمون با استفاده از داده‌های نمونه محاسبه می‌شود. این آماره مقدار استاندارد شده میانگین یا نسبت یا تفاوت نمونه است و این همان نمره Z یا نمره t است که به صورت زیر به دست می‌آید:

$$\text{مقدار فرضیه صفر} - \text{آماره نمونه‌ای مشاهده شده} = \frac{\text{آماره آزمون}}{(\text{برآورد}) \text{ خطای معیار آماره}}$$

(۴) ناحیه رد (Rejection Region) تعیین می‌شود. این ناحیه شامل مجموعه‌ای از مقادیر آماره آزمون است که با H_0 مغایرت دارد و متضمن رد آن است.

ناحیه رد معمولاً ناحیه دوم یا دنباله توزیع Z یا t است و احتمال اینکه یک مقدار در ناحیه رد قرار گیرد «سطح معنی‌داری» نامیده و با علامت α نشان داده می‌شود. مقدار α معمولاً مساوی ۰/۰۵ یا ۰/۰۱ گرفته می‌شود.

مثال: فرض کنید ضربان قلب دختران جوان سالم در حالت استراحت توزیعی نرمال با میانگین ۶۶ ضربان در دقیقه داشته باشد و انحراف معیار آن نیز $\sigma = 9/2$ باشد. دارویی به ۱۰۰ زن جوان و سالم داده شده و ضربان قلب آنها پس از مصرف دارو اندازه‌گیری شده است. میانگین $\bar{x} = 68$ ضربان در دقیقه برای این گروه نمونه به دست آمده است. گمان بر این است که این دارو میزان ضربان قلب دختران را افزایش داده ولی انحراف معیار آن تغییری نکرده است. این فرضیه را آزمون کنید.

دو ادعا مطرح است:

(الف) این دارو اثری در ضربان قلب ندارد ($\mu = 66$).

(ب) این دارو میزان ضربان قلب را افزایش می دهد ($\mu > 66$) .
مرحله (۱). ادعای (الف) به عنوان فرضیه صفر گرفته می شود:

$$H_0: \mu = 66$$

مرحله (۲): ادعای (ب) به عنوان فرضیه پژوهش یا فرضیه مقابل گرفته می شود:

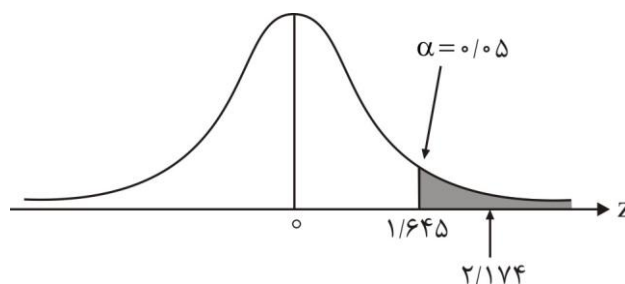
$$H_A: \mu > 66$$

در روش آزمون فرضیه، H_0 فرضیه صحیح انگاشته می شود و براساس همین فرض، میانگین نمونه $\bar{X} = 68$ استاندارد می شود تا معلوم گردد که آیا این مقدار در ناحیه رد با سطح معنی داری α قرار می گیرد یا نه.

مرحله (۳): $\bar{X} = 68$ است و با استاندارد کردن آن، آماره آزمون به دست می آید:

$$Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} = \frac{68 - 66}{\frac{9}{\sqrt{100}}} = \frac{2}{0.9} = 2.22$$

مرحله (۴): مقدار α را مساوی ۰/۰۵ می گیریم. چون فرضیه مقابل $\mu > 66$ است آزمون را یک سویه می گویند. ناحیه رد در شکل زیر نشان داده شده و مقدار $Z = 1.645$ به این دلیل انتخاب شده که سطح زیرمنحنی در سمت راست $Z = 1.645$ برابر ۰/۰۵ است.



این منحنی توزیع مقادیر Z است که $Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}}$ و در شرایط H_0 ، $\mu = 66$ است. مقدار $Z = 2.22$ در ناحیه رد H_0 قرار

می گیرد. بنابراین آماره آزمون در سطح ۰/۰۵ معنی دار است. پس H_0 را به نفع H_A رد می کنیم و نتیجه می گیریم که داروی موردنظر متوسط ضربان قلب را افزایش داده است.

اگر در این مثال، σ یعنی انحراف معیار جامعه معلوم نبود، می بایست مقدار آن را از روی مقدار s نمونه با $n - 1$ درجه آزادی برآورد کنیم. در اینجا دیگر مقدار معادل $Z = 1.645$ در جدول t برابر ۱/۶۵۸ است.

نکاتی در مورد آزمون آماری

- سطح معناداری یا α ، ناحیه رد را تعیین می کند. مثلاً اگر در مثال پیشین $\alpha = 0.01$ باشد، ناحیه رد، مقادیر $Z \geq 2.33$ خواهد بود. این ناحیه در شکل زیر نشان داده شده است. در این مورد، آماره به دست آمده معنی دار نیست و بنابراین نمی توان فرضیه صفر را به طور کامل مستدل در این سطح معنی داری رد کرد.