

مقدمات آمار کاربردی

در مطالعات علوم پزشکی

کارگاه روش تحقیق

دوشنبه ۱۹ آبان ۹۵

دانشگاه علوم پزشکی استان سمنان

دکتر مجید میر محمدخانی

اپیدمیولوژیست

داده ، اطلاعات، آگاهی

• بین کلمات "داده" (data) ، "اطلاعات" (information) و "آگاهی" (intelligence) تفاوت وجود دارد.

• **داده** شامل مشاهدات جداگانه‌ای از صفات یا رویدادها است که وقتی به تنهایی مورد توجه قرار می‌گیرند دارای معنای کمی هستند؛ به طوری که داده‌های جمع‌آوری شده از نظام‌ها یا مراکز اداری سلامت برای برنامه‌ریزی کفایت لازم را ندارند.

• داده‌ها باید تبدیل به **اطلاعات** شوند و این کار با خلاصه‌سازی، جمع‌بندی و تطبیق آنها برای متغیرهایی نظیر ترکیب سنی و جنسی جامعه صورت می‌گیرد تا انجام مقایسه‌های زمانی و مکانی میسر گردد.

• اطلاعات با یکپارچه‌سازی و پردازش براساس درک و تجربیات برآمده از ارزش‌های اجتماعی و سیاسی تبدیل به **آگاهی** می‌گردد.

داده‌هایی که به اطلاعات و اطلاعاتی که به آگاهی تبدیل نمی‌شوند، ارزش زیادی برای راهنمایی تصمیم‌گیرندگان، سیاست‌گذاران، برنامه‌ریزان، مجریان و خود کارکنان مراقبت سلامت نخواهند داشت.

ارائه‌ی داده‌های آماری

□ در هر زمینه‌ای از تحقیق و جستجو ، ابتدا داده‌ها جمع‌آوری می‌شوند و پس از طبقه‌بندی و تحلیل ، درستی نتایج برآمده از آنها با استفاده از آزمون‌های آماری ارزیابی می‌گردد.

➤ بعد از این که داده‌های آماری جمع‌آوری شد، بایستی هدفمندانه مرتب شوند تا بتوان نکات مهم را به‌طور واضح و کامل از آنها استخراج نمود. روش‌های متفاوتی برای این کار وجود دارد:

- ❖ ترسیم جداول (tables) ، نمودارها (charts) ، طرح‌ها (diagrams) ، نگاره‌ها (graphs) ، تصاویر (pictures) ، منحنی‌ها (curves) ،
- ❖ استفاده از شاخص‌های آماری شامل آماره‌های مرکزی (معدل‌های آماری) و مقادیر پراکندگی

ترسیم جدول

- به وسیله ترسیم جداول می توان به سادگی انبوهی از داده های آماری را به نمایش گذارد.
- اولین اقدام قبل از تجزیه و تحلیل داده ها ساختن جداول می باشد.
- جدول ممکن است ساده یا پیچیده باشد که به تعداد یا اندازه ی مجموعه ی منفرد یا مجموعه های متعدد حاوی موضوعات مورد بررسی بستگی دارد.

Table 1. Means and Standard Deviations for the Duration of Retention in the Program per Week

Variable	Retention, wk ^a	P Value ^b
Prison		< 0.001
Isfahan	23.5 ± 0.9	
Hamadan	11.3 ± 1.0	
Tehran	26.4 ± 0.3	
Drug injection		0.008
Yes	25.9 ± 0.7	
No	21.9 ± 0.8	
Main used drug		0.001
Opium/opium extract	18.6 ± 3.1	
Heroin and crack	24.6 ± 0.6	
Ecstasy	25.9 ± 1.0	
Other drugs	9.8 ± 0.9	
HBV test		0.02
Negative	25.0 ± 0.6	
Positive/unknown	23.1 ± 0.8	
Occupation/income		0.05
Stable job / regular income	22.7 ± 1.1	
Instable job / irregular income	24.1 ± 0.6	
Main cause of imprisonment		< 0.001
Drug-related crime	25.8 ± 0.6	
Crimes not related to drugs	18.4 ± 1.5	
History of involvement in harm reduction programs		< 0.001
Yes	19.9 ± 1.2	
No	25.7 ± 0.6	

^a Data are presented as Mean ± SD.

^b Log-Rank Test.

اصول طراحی جدول

- الف- جداول باید **شماره گذاری** شوند مثل جدول ۱، جدول ۲ و نظیر آن.
- ب- برای هر جدول باید **عنوانی** در نظر گرفته شود. عنوان باید مختصر و به تنهایی گویا باشد.
- ج- **عناوین ستون‌ها و سطرها** باید روشن و مختصر باشند.
- د- داده‌ها بایستی **به ترتیب** براساس مقدار یا اهمیت، تقدم و تاخر زمانی، حروف الفبا و یا طبق ملاحظات جغرافیایی ارائه شوند.
- ه- اگر قرار است مقادیر به شکل درصد یا میانگین با هم **مقایسه** شوند تا آن جا که می شود بهتر است مقادیر آن‌ها نزدیک به هم ارائه شوند.
- و- **جدول نباید خیلی طولانی باشد.**
- ز- بیشتر مردم، جداول تنظیم شده به صورت **عمودی** را بهتر از افقی می‌دانند، چرا که مرور داده‌ها از بالا به پایین آسانتر است.
- ح- می توان برای جدول **زیرنویس** (Foot notes) تهیه کرد، تا در صورت لزوم ارائه‌دهنده‌ی اطلاعات اضافی یا نکات توضیحی باشد.

مثال هایی از جداول ساده

جدول ۲، جمعیت هند

Year	Population
1901	238,396,000
1921	251,321,000
1981	685,185,000
1991	843,930,000
2001	1027,015,247

* Source: Census of India, 2001

جمعیت ایران سال‌های ۱۳۵۵ تا ۱۳۹۰

سال	کشور
۱۳۵۵	۳۳.۷۰۸.۷۴۴
۱۳۶۵	۴۹.۸۵۷.۳۸۴
۱۳۷۵	۶۰.۰۵۵.۴۸۸
۱۳۸۵	۷۰.۴۹۵.۷۸۲
۱۳۹۰	۷۵.۱۴۹.۶۶۹

آ) جدول ۱، جمعیت برخی از ایالات هند*

States	Population 2001
Andhra Pradesh	75,727,541
Bihar	82,878,796
Madhya Pradesh	60,385,118
Uttar Pradesh	116,052,859

* Source: Census of India, 2001

جمعیت برخی از استان‌های ایران در سال ۱۳۹۰

استان	پیت سال
تهران	۱۲.۱۸۳.۳۹۱
خراسان رضوی	۵.۹۹۴.۴۰۲
اصفهان	۴.۸۷۹.۳۱۲
فارس	۴.۵۹۶.۶۵۸
آذربایجان شرقی	۳.۷۲۴.۶۲۰

جدول توزیع فراوانی

- در جدول توزیع فراوانی، ابتدا داده‌ها به گروه‌های مناسب (تحت عنوان فواصل طبقه‌ای یا Class intervals) تقسیم شده و تعداد افراد (فراوانی) هر گروه، در ستون کنار آن نمایش داده می‌شود.

نمونه: اعداد زیر سن بیماران را نشان می‌دهد که با تشخیص فلج کودکان در یک بیمارستان بستری شده‌اند.
جدول توزیع فراوانی آن را ترسیم نمایید.

۸، ۲۴، ۱۸، ۵، ۶، ۱۲، ۴، ۳، ۳، ۳، ۲، ۳، ۲۳، ۹، ۱۸، ۱۶، ۱، ۲، ۳، ۵، ۱۱، ۱۳، ۱۵، ۹، ۱۱، ۱۱، ۷، ۶، ۱۰، ۹، ۵، ۱۶، ۲۰، ۴، ۳، ۳، ۳، ۱۰، ۳، ۱، ۲، ۹، ۶، ۳، ۷، ۱۴، ۸، ۴، ۶، ۱۵، ۲۲، ۱، ۴، ۷، ۱۲، ۳، ۲۳، ۴، ۱۹، ۶، ۲، ۴، ۱۴، ۲، ۲۱، ۳، ۳، ۹، ۲، ۱، ۷، ۱۹

داده‌های مذکور را به راحتی می‌توان به صورت زیر مرتب نمود:

فراوانی	گروه سنی
	-
	-
	-
	-
	-

داده‌هایی که به شکل مذکور مرتب شده‌اند، به صورت جدول فراوانی نمایش داده می‌شوند:

جدول ۳: توزیع سنی بیماران مبتلا به فلج کودکان

شمار بیماران	سن
۳۵	۴-۰
۱۸	۹-۵
۱۱	۱۴-۱۰
۸	۱۹-۱۵
۶	۲۴-۲۰

جدول ۳، توزیع سنی بیماران مبتلا به فلج کودکان

سن	شمار بیماران
۴-۰	۳۵
۹-۵	۱۸
۱۴-۱۰	۱۱
۱۹-۱۵	۸
۲۴-۲۰	۶

- در مثال بالا، بیماران از نظر سن به ۵ گروه تقسیم شده‌اند و تعداد مشاهدات در هر گروه **فراوانی** نام دارد.

✓ سوالی که در رسم جداول توزیع فراوانی مطرح می‌شود این است که داده‌ها را در چند گروه طبقه‌بندی کنیم؟ و چه فواصلی باید برای طبقه‌بندی در نظر گرفته شود؟

✓ پاسخ: بر اساس یک قانون عملی وقتی داده‌های زیادی وجود دارد حداکثر ۲۰ گروه و وقتی داده‌ها چندان زیاد نیست، حداقل ۵ گروه قابل تعیین است.

- تا جایی که مقدور است فاصله‌ی طبقات باید برابر باشد تا بتوان مشاهدات را با هم مقایسه کرد.
- مزیت جدول توزیع فراوانی آن است که در یک نگاه نشان می‌دهد چه تعداد از مشاهدات در هر گروه قرار گرفته‌اند و تراکم اصلی در کدام گروه واقع شده‌است.

نمودارها و طرح ها

- نمودارها و طرح ها روش های مفیدی برای نمایش داده های ساده ی آماری محسوب می شوند.

➤ آنها تاثیر قوی بر ذهن مردم می گذارند. بنابراین برای تبیین داده های آماری ، به خصوص در روزنامه ها و مجله ها ، رسانه های متداولی محسوب می شوند.

- در مورد نمودار و طرح سه نکته مهم قابل ذکر است:

I. طرح ها بهتر از جداول آماری در ذهن می مانند.

II. داده هایی که قرار است به شکل طرح نمایش داده شوند ، باید ساده باشند. به این ترتیب نگران این نخواهیم بود که خواننده در فهم مطلب دچار مشکل شود.

III. سادگی طرح به قیمت کاهش صحت و دقت داده ها تمام خواهد شد. یعنی بسیاری از جزئیات داده های اولیه، در نمودارها و طرح ها از بین خواهد رفت. اگر ما در پی نکات واقعی یک مطالعه هستیم، بایستی به داده های اولیه رجوع کنیم.

نمودارهای ستونی

❖ نمودارهای ستونی (Bar charts) برای نشان دادن مجموعه ای از اعداد به کار می روند که در آن طول هر ستون، متناسب با بزرگی اندازه ای است که قرار است نمایش داده شود.

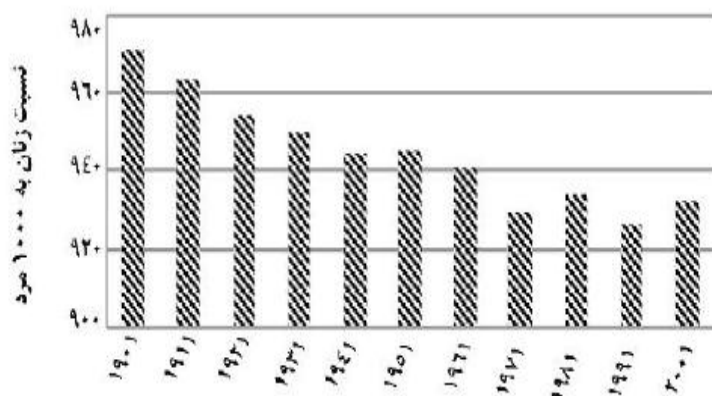
❖ از آنجایی که نمودارهای ستونی به راحتی قابل ترسیم هستند و با آنها می توان ارقام را به طور چشمی مقایسه نمود ، وسیله ی متداولی به شمار می آیند.

• انواع نمودارهای ستونی:

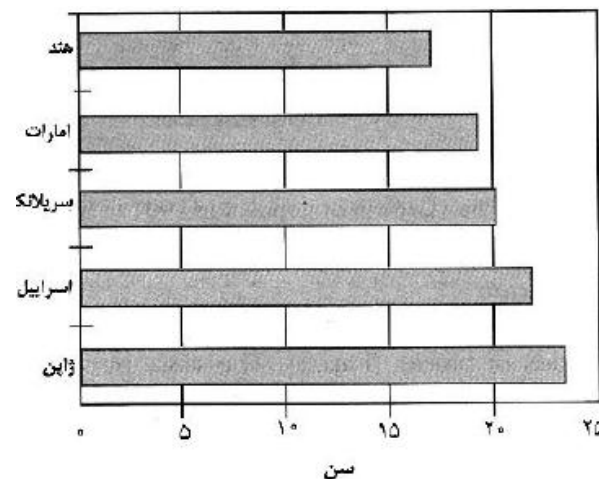
- الف- نمودار ستونی ساده (simple)
- ب- نمودار ستونی چندتایی (multiple)
- ج- نمودار ستونی مرکب (component)

نمودار ستونی ساده

- ستون‌ها یا عمودی هستند یا افقی.
- ستون‌ها معمولاً با فاصله‌هایی از هم جدا می‌شوند تا منظم‌تر و واضح‌تر به نظر برسند.
- از نظر اندازه لازم است مقیاس مناسبی برای نمایش ستون‌ها انتخاب شود.



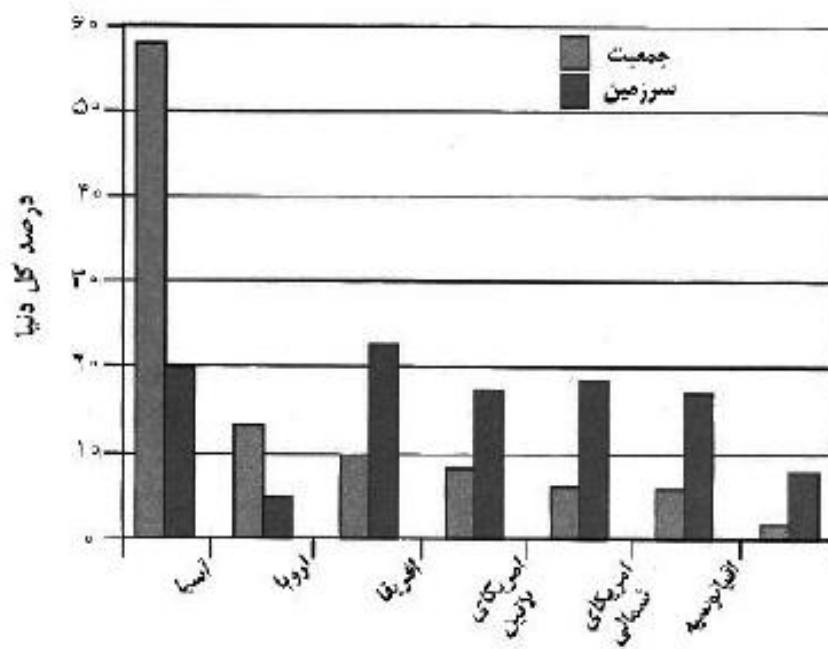
شکل ۱. کسر جنسی^۱ در هند، ۱۹۰۱ تا ۲۰۰۱



شکل ۲. میانگین سن ازدواج (زنان) در برخی کشورها

نمودار ستونی چندتایی

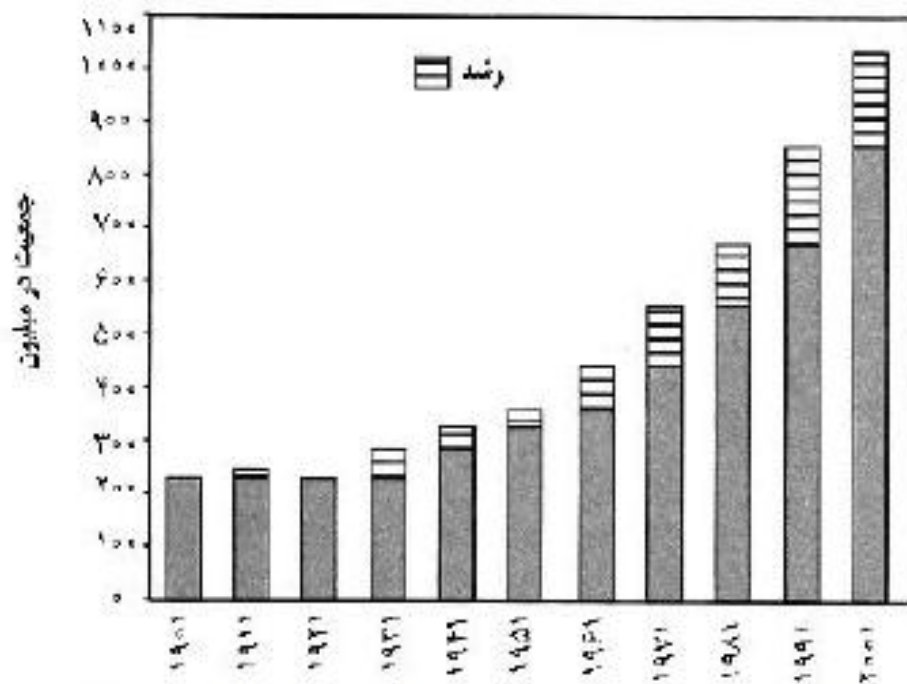
- شکل زیر ، نمونه‌ای از نمودار ستونی چندتایی یا نمودار ستونی ترکیبی را نشان می‌دهد. در این نوع از نمودار، دو ستون یا بیشتر ممکن است با هم یک گروه را تشکیل دهند. در این شکل ، جمعیت و وسعت سرزمین در مناطق مختلف دنیا با هم مقایسه شده‌اند.



شکل ۱۳: جمعیت و وسعت سرزمین به تفکیک منطقه

نمودار ستونی مرکب

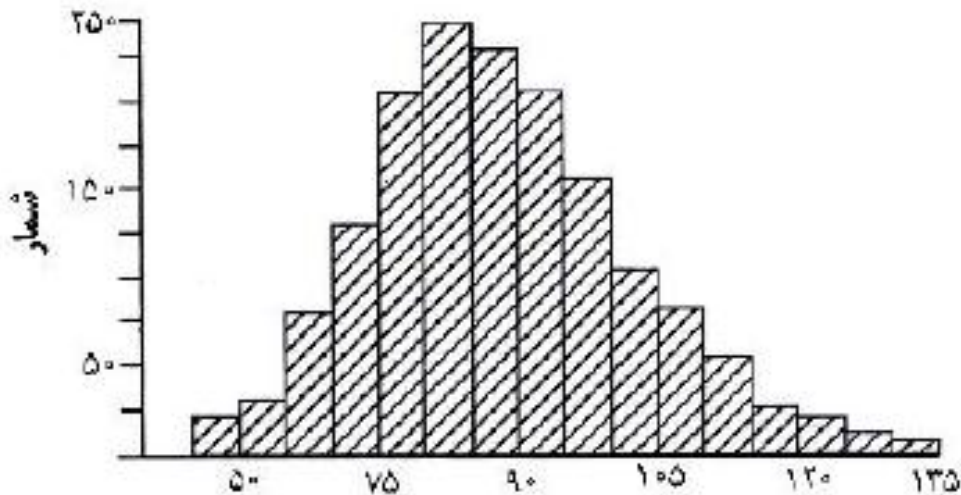
- ستون‌ها ممکن است به دو یا چند قسمت تقسیم شوند. هر قسمت مربوط به موضوع خاصی است و اندازه‌ی آن متناسب با بزرگی یا مقدار همان موضوع می‌باشد.



شکل ۴: رشد جمعیت در سال‌های ۱۹۰۱ تا ۲۰۰۱

هستوگرام

- هستوگرام (histogram) یک طرح تصویری از توزیع فراوانی است و دارای مجموعه‌ای از قطعات مستطیل شکل (blocks) می باشد.
 - ❖ فواصل طبقه‌ای در محور افقی و فراوانی‌ها در محور عمودی قرار می گیرند.
 - ❖ مساحت هر قطعه یا مستطیل متناسب با فراوانی موضوع متناظر است.

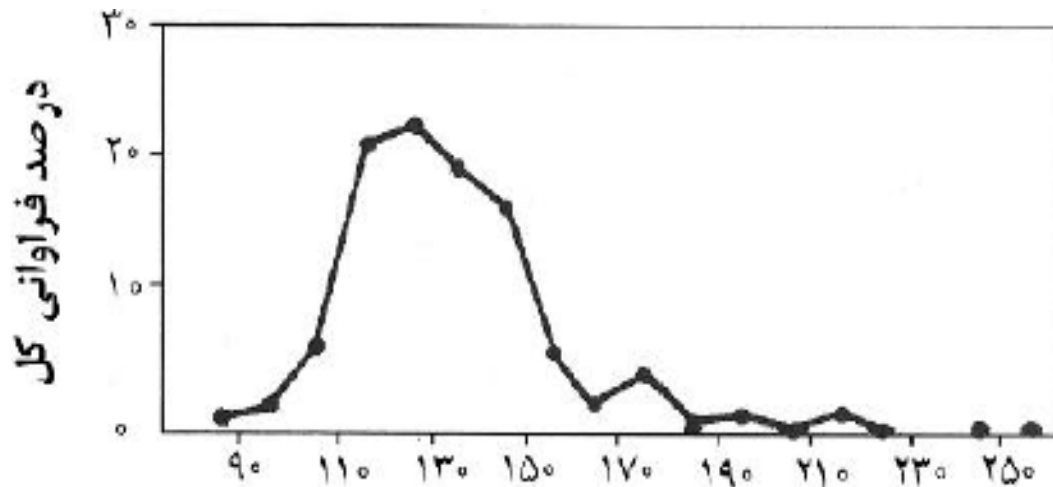


شکل ۵. توزیع فراوانی فشارخون دیاستولیک در زنان ۴۵ تا ۶۴ سال

چندگوشه‌ی فراوانی

- توزیع فراوانی را می‌توان در قالب طرحی موسوم به چندگوشه‌ی فراوانی (Frequency Polygon) نشان داد.

❖ این طرح از اتصال نقاط میانی قطعات هیستوگرام به دست می‌آید. شکل زیر، چندگوشه‌ی فراوانی فشار خون سیستولیک اندازه‌گیری شده را در یک جامعه نشان می‌دهد.

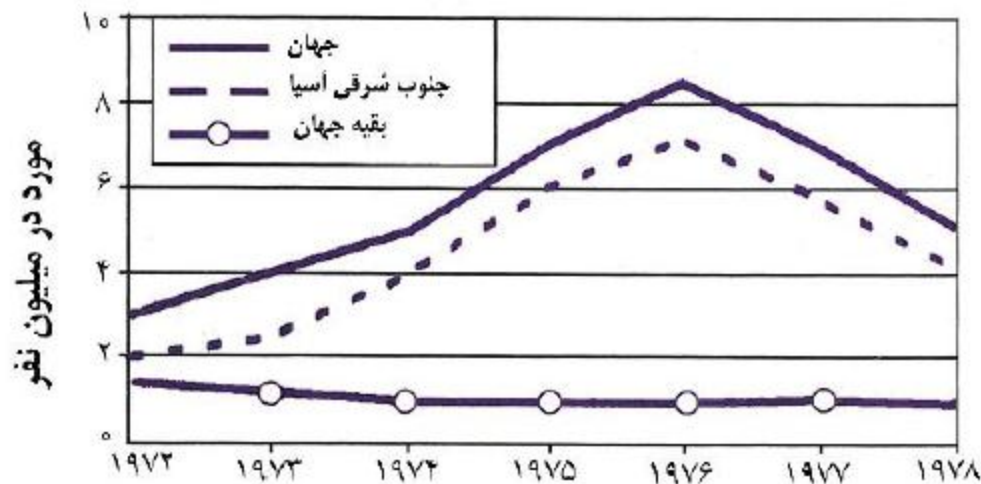


شکل ۶، توزیع فراوانی فشارخون سیستولیک اندازه‌گیری شده

طرح خطی

- طرح‌های خطی (Line Diagram) برای نشان دادن روند رویدادها در طول گذر زمان به کار می‌روند.

➤ شکل زیر نمونه‌ای از یک طرح خطی می‌باشد که روند گزارش موارد مالاریا در جهان (غیر از منطقه‌ی آفریقا) را طی سال‌های ۱۹۷۲ تا ۱۹۷۸ نشان می‌دهد.



شکل ۷، موارد مالاریای گزارش شده، ۱۹۷۱-۱۹۷۸ (به استثنای منطقه‌ی آفریقایی)

نمودار دایره ای

- در نمودار دایره ای یا کلوچه ای (Pie chart) برخلاف نمودار ستونی که در آن طول ستون‌ها با هم مقایسه می‌شد، مساحت بخش‌های یک دایره با هم مقایسه می‌شود.

- **مساحت هر بخش به زاویه‌ی آن بستگی دارد.**

❖ برخلاف مقبولیت زیاد این نوع نمودار میان عموم مردم، آمارشناسان نمودارهای ستونی را بهتر از نمودارهای دایره‌ای می‌دانند.

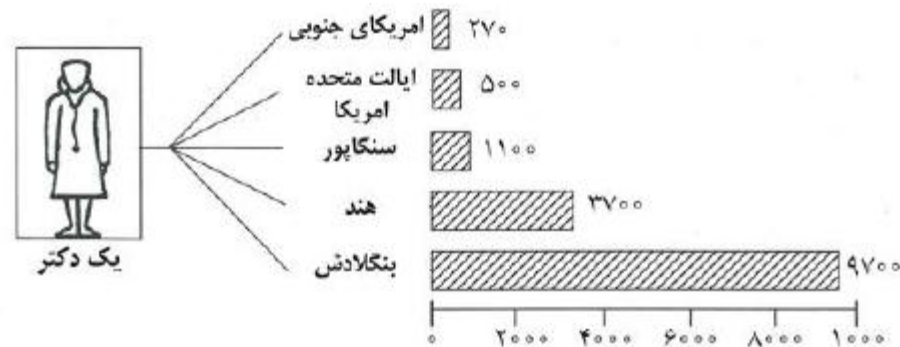
❖ اغلب لازم است برای هر بخش از دایره درصد فراوانی متناظر ذکر شود، چرا که گاهی مقایسه‌ی چشمی مساحت میان بخش‌ها به آسانی مقدور نمی‌باشد.



شکل ۸، جمعیت جهان

طرح تصویری

- طرح تصویری (Pictogram) میان مردم کوچه و بازار و برای کسانی که از نمودارهای دقیق و پیچیده کمتر سردر می آورند، یک روش مقبول محسوب می شود.
 - در این روش از نمادها یا عکس های کوچک استفاده می شود.
- ❖ مثلاً در طرح تصویری زیر، نسبت پزشک به جمعیت در کشورهای مختلف با استفاده از تصویر ساده ی یک پزشک نشان داده شده است.

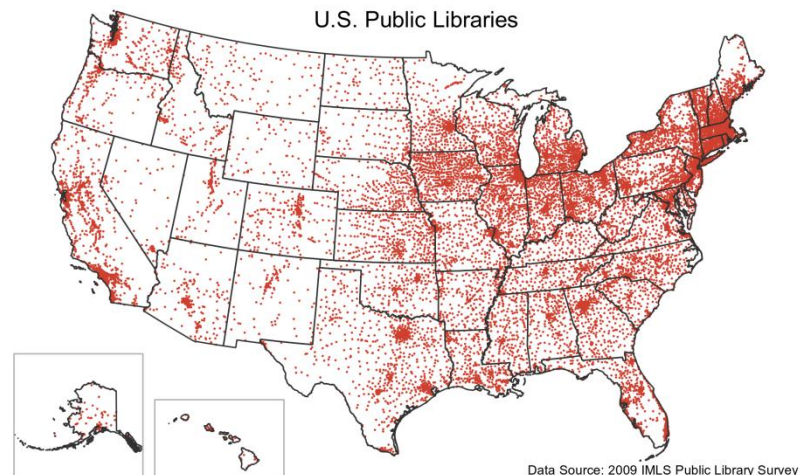
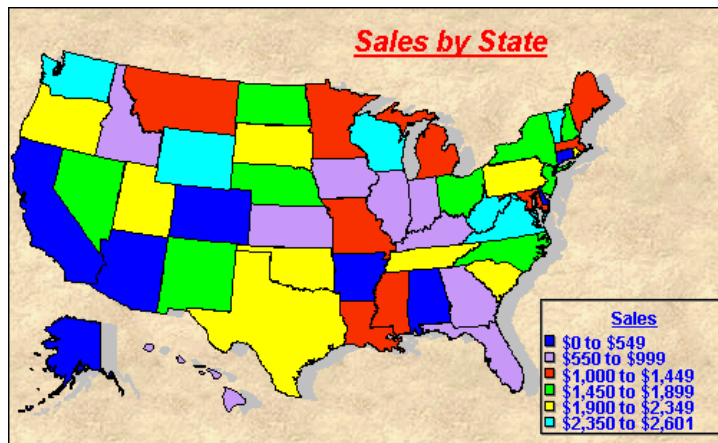


شکل ۹، نسبت جمعیت به پزشک، ۱۹۸۱

نقشه های آماری

- وقتی داده های آماری مربوط به مناطق جغرافیایی یا اجرایی باشند، می توان آن ها را با استفاده از "نقشه های هاشورزده" (Shaded maps) یا "نقشه های نقطه گذاری شده ی" (Dot maps) مناسب نیز نشان داد.

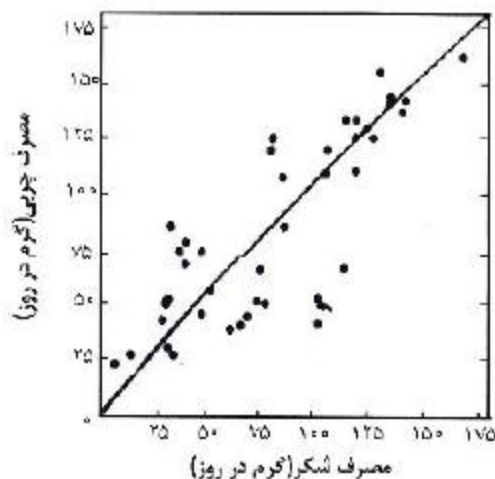
- نقشه های هاشورزده برای نشان دادن داده هایی استفاده می شوند که از نظر اندازه متفاوت باشند.
- قسمت های مختلف نقشه با استفاده از رنگ های مختلف یا با شدت های متفاوت از یک رنگ ، هاشورزده می شوند و در کلید نقشه برای راهنمایی بیننده تفسیر هر کدام آورده می شود.



طرح نقاط پراکنده

- طرح نقاط پراکنده (Scatter) برای نشان دادن ارتباط میان دو متغیر به کار می رود.
- به عنوان نمونه در شکل زیر، ارتباط مثبت بین متوسط مصرف غذایی چربی و مصرف قند در ۴۱ کشور نشان داده شده است.

❖ نکته: اگر تجمع نقاط بیشتر حول و حوش یک خط صاف باشد، شواهد به نفع وجود ارتباط با ماهیت خطی بین متغیرها خواهد بود. اما اگر پراکندگی نقاط به این صورت نباشد، احتمالاً ارتباطی بین متغیرها وجود ندارد.



شکل ۱۰. ارتباط بین متوسط مصرف چربی و قند در ۴۱ کشور

معدلهای آماری

- واژه "معدل آماری" (Statistical Averages) گویای مقداری از توزیع است که دیگر مقادیر حول و حوش آن پراکنده شده‌اند، به طوری که از نظر مقدار تصویری ذهنی از وسط (central) توزیع را به ما می‌دهد.
- انواع مختلفی از معدهای آماری وجود دارد که رایج‌ترین آنها عبارتند از:
 - (۱) میانگین حسابی
 - (۲) میانه
 - (۳) نما

میانگین

میانگین حسابی کاربرد زیادی در محاسبات آماری دارد. گاهی آن را برای سادگی فقط میانگین (Mean) می‌نامند. برای محاسبه‌ی میانگین ابتدا اندازه‌ی همه مشاهددها با هم جمع و سپس مجموع به‌دست آمده بر شمار کل مشاهددها بخش می‌شود. عمل جمع کردن مشاهددها را «جمع‌زنی» (summation) می‌نامند و با علامت Σ یا S نشان می‌دهند. شمار مشاهددها با علامت n و میانگین را با علامت $\bar{X}$ (که « X بار» خوانده می‌شود) نمایش می‌دهند.

میانگین ($\bar{X}$) به این ترتیب محاسبه می‌شود: فرض کنیم فشار خون دیاستولیک ۱۰ نفر ۸۳، ۷۵، ۸۱، ۷۹، ۷۱، ۹۵، ۷۵، ۷۷، ۸۴، ۹۰ بوده‌است. جمع کل فشار خون‌ها برابر است با ۸۱۰ و میانگین برابر است با ۸۱۰ بخش بر ۱۰ که می‌شود ۸۱. برتری میانگین این است که محاسبه و فهم آن آسان است. از عیب‌های آن یکی این است که گاهی ممکن است زیر تأثیر اندازه‌های غیر عادی قرار گیرد. دیگر این که گاهی ممکن است از نظر مفهوم مضحک به نظر برسد؛ مانند میانگین بچه‌های متولد شده از هر زن در یک منطقه خاص ۴/۷۶ محاسبه شده است، در حالی که این رقم هرگز در واقعیت معنی نمی‌دهد. با وجود این عیب‌ها، میانگین حسابی از مفیدترین معدل‌های آماری به‌شمار می‌آید.

میانه

میانه (median)، معدلی است که به نوعی از میانگین متفاوت است و به جمع مقادارها و شمار مشاهده‌ها وابسته نیست. برای به دست آوردن میانه، ابتدا داده‌ها بر اساس مقدارشان به صورت صعودی یا نزولی مرتب می‌شوند و سپس مشاهده‌ای که از نظر مقدار در وسط قرار گرفته است به عنوان میانه نام می‌گیرد. برای مثال فشار خون دیاستولیک ۹ نفر در شکل ۱۱ درج شده است. میانه عدد ۷۹ است که در وسط داده‌ها واقع شده است (شکل ۱۲).

فشارخون دیاستولیک	فشارخون دیاستولیک
۷۱	۸۳
۷۵	۷۵
۷۵	۸۱
۷۷	۷۹
۷۹ میانه	۷۱
۸۱	۹۵
۸۳	۷۵
۸۴	۷۷
۹۵	۸۴

شکل ۱۲، داده‌های مرتب شده بر اساس مقدار

شکل ۱۱، داده‌های مرتب نشده

اگر به جای ۹ مشاهده، ۱۰ مشاهده وجود می داشت، میانه با محاسبه ی میانگین دو رقم مربوط به مشاهددهای میانی به دست می آمد. بنابراین هرگاه شمار مشاهده ها زوج باشد، برای به دست آوردن میانه باید میانگین دو رقم میانی را حساب کرد. مانند فشار خون دیاستولیک ۱۰ نفر به شرح مندرج در شکل ۱۳ است:

فشارخون دیاستولیک	
۷۱	
۷۵	
۷۵	
۷۷	
۷۹	} $\frac{۷۹ + ۸۱}{۲}$
۸۱	
۸۳	
۸۴	=۸۰
۹۰	
۹۵	

فشارخون دیاستولیک
۸۳
۷۵
۸۱
۷۹
۷۱
۹۵
۷۵
۷۷
۸۴
۹۰

شکل ۱۳، داده های * در این نمونه میانه برابر است با ۷۹+۸۱ بخش بر ۲ که می شود ۸۰ تب شده بر اساس مقدار

نما

نما (mode) به مقداری از مشاهده‌ها می‌گویند که بیشترین شمار را در یک توزیع داده‌ای دارد. در واقع نما مقداری است که بیشتر از همه دیده می‌شود و در میان مشاهده‌ها از همه «رایج‌تر» است. برای نمونه فشار خون دیاستولیک ۲۰ نفر به شرح زیر است:

۸۵، ۷۵، ۸۱، ۷۹، ۷۱، ۹۵، ۷۵، ۷۷، ۷۵، ۹۰، ۷۱، ۷۵، ۷۹، ۷۵، ۷۷، ۸۴، ۷۵، ۸۱، ۷۵

نما یا پرتکرارترین مشاهده در میان این داده‌ها ۷۵ است. برتری نما این است که فهم آن آسان است و تحت تأثیر مقدرهای خیلی کم یا خیلی زیاد قرار نمی‌گیرد. از نقاط ضعف آن یکی این است که جایگاه واقعی آن مشخص نیست و دیگر این که اغلب دارای تعریف روشنی نیست. بنابراین نما معمولاً در آمارهای پزشکی و علوم زیستی استفاده نمی‌شود.

مقدارهای پراکندگی

مقدار کالری مورد نیاز برای فرد سالمی که در کارش تحرک چندانی وجود ندارد، حدود ۲۴۰۰ کالری در روز است. بدیهی است که این مقدار برای همه صادق نیست. بین انسان‌ها تفاوت‌های بسیاری وجود دارد. اگر مقدارهای فشارخون، قد یا وزن افراد را در یک گروه بزرگ اندازه‌گیری کنیم، خواهیم دید که این مقادیر از فردی به فرد دیگر متفاوت است. حتی در یک فرد هم، بسته به زمان اندازه‌گیری، مقدار آن متفاوت خواهد بود. پرسش‌هایی که مطرح می‌شود عبارت‌اند از این که: مقدار طبیعی این تغییرات چقدر است؟ و چگونه می‌توان این تغییرات را اندازه گرفت؟

چندین مقدار برای اندازه‌گیری تغییرات (یا در اصطلاح فنی «پراکندگی») وجود دارد که شماری از شناخته شده‌ترین آن‌ها عبارت‌اند از:

(آ) دامنه (Range) (ب) میانگین یا متوسط انحرافات (پ) انحراف معیار (Standard Deviation)

الف) دامنه

- دامنه (range) را شاید بتوان ساده ترین شاخص برای اندازه گیری پراکندگی دانست و عبارت است از اختلاف بین بیشترین و کمترین مقادیر در یک نمونه.
- برای مثال مقادیر فشار خون دیاستولیک ۱۰ نفر را به شرح ذیل در نظر بگیرید:

۸۳، ۷۵، ۸۱، ۷۹، ۷۱، ۹۰، ۷۵، ۹۵، ۷۷، ۹۴

- همان طور که مشاهده می شود بیشترین مقدار ۹۵ و کمترین مقدار ۷۱ بوده است. بنابراین دامنه برابر است با فاصله ی ۷۱ تا ۹۵ و یا تفاضل مطلق این دو مقدار یعنی ۲۴.
 - در داده های گروه بندی شده دامنه می تواند از تفاضل بین نقطه های میانی کوچک ترین و بزرگ ترین گروه ها نیز به دست آید.
- دامنه در عمل کاربرد زیادی ندارد ، چرا که فقط نشان دهنده دو مقدار انتهایی است و پراکندگی مشاهداتی که بین این دو مقدار قرار دارند در نظر گرفته نمی شود.

ب) میانگین انحرافات

میانگین انحرافات، معدل مقدارهای انحرافات هر مشاهده از میانگین حسابی داده‌هاست و از فرمول زیر محاسبه می‌شود:

$$M.D. = \frac{\sum - \bar{x}}{n}$$

نمونه: فشار خون دیاستولیک ۱۰ نفر به این شرح است: ۸۳، ۷۵، ۸۱، ۷۹، ۷۱، ۹۵، ۷۵، ۷۷، ۸۴، ۹۰. میانگین انحرافات را محاسبه کنید.

پاسخ (میانگین انحرافات)

فشارخون دیاستولیک	میانگین حسابی	انحراف از میانگین
x	$\bar{x}$	$(x - \bar{x})$
۸۳	۸۱	۲
۷۵	۸۱	-۶
۸۱	۸۱	۰
۷۹	۸۱	-۲
۷۱	۸۱	-۱۰
۹۵	۸۱	۱۴
۷۵	۸۱	-۶
۷۷	۸۱	-۴
۸۴	۸۱	۳
۹۰	۸۱	۹
جمع = ۸۱۰		(بدون توجه به علامت $\pm$) ۵۶ = جمع

$$\text{میانگین} = \frac{۸۱۰}{۱۰} = ۸۱$$

$$\text{میانگین انحرافات} = \frac{۵۶}{۱۰} = ۵/۶$$

پ) انحراف معیار

انحراف معیار رایج‌ترین مقدار پراکندگی به شمار می‌آید. انحراف معیار به زبان ساده «ریشه‌ی دوم میانگین مجذور انحرافات» است و با حرف یونانی σ [سیگما] و یا به شکل مخفف S.D.^۱ نوشته می‌شود. فرمول اصلی برای محاسبه‌ی انحراف معیار عبارت است از:

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

هنگامی حجم نمونه بزرگتر از ۳۰ باشد، فرمول یاد شده بدون هیچ تغییری به کار می‌رود؛ اما برای نمونه‌های کوچک‌تر فرمول مذکور انحراف معیار را کمتر از مقدار واقعی برآورد می‌کند و نیاز به اصلاح دارد که این کار با جانشینی $(n-1)$ به جای n در مخرج کسر انجام می‌شود. فرمول تصحیح شده عبارت است از:

$$S.D. = \sqrt{\frac{\sum (x - \bar{x})^2}{n-1}}$$

۱- به گونه‌ی کلی از حرف سیگمای یونانی (Σ) برای بیان انحراف معیار برای کل جامعه و از حروف (S یا SD) برای بیان انحراف معیار نمونه‌ها استفاده می‌گردد. (ویر)

مراحلی که باید در محاسبه انحراف معیار انجام شود بدین شرح است:

ا) پیش از هر کاری، انحراف مقدار هر مشاهده را از میانگین حسابی به دست آورید، $(x - \bar{x})$

ب) سپس مجذور هر انحراف را حساب کنید، $(x - \bar{x})^2$

پ) مجذور انحرافات را با هم جمع کنید، $\sum(x - \bar{x})^2$

ت) مقدار به دست آمده را بر شمار کل مشاهدات (n) تقسیم کنید [اگر حجم نمونه کمتر از ۳۰ باشد بر $(n-1)$]

تقسیم می کنیم]

ث) آنگاه از نتیجه جذر بگیرید تا انحراف معیار به دست آید.

نمونه: فشار خون دیاستولی ۱۰ نفر به این شرح است: ۸۳، ۷۵، ۸۱، ۷۹، ۷۱، ۹۵، ۷۵، ۷۷، ۸۴، ۹۰. انحراف معیار را

محاسبه کنید.

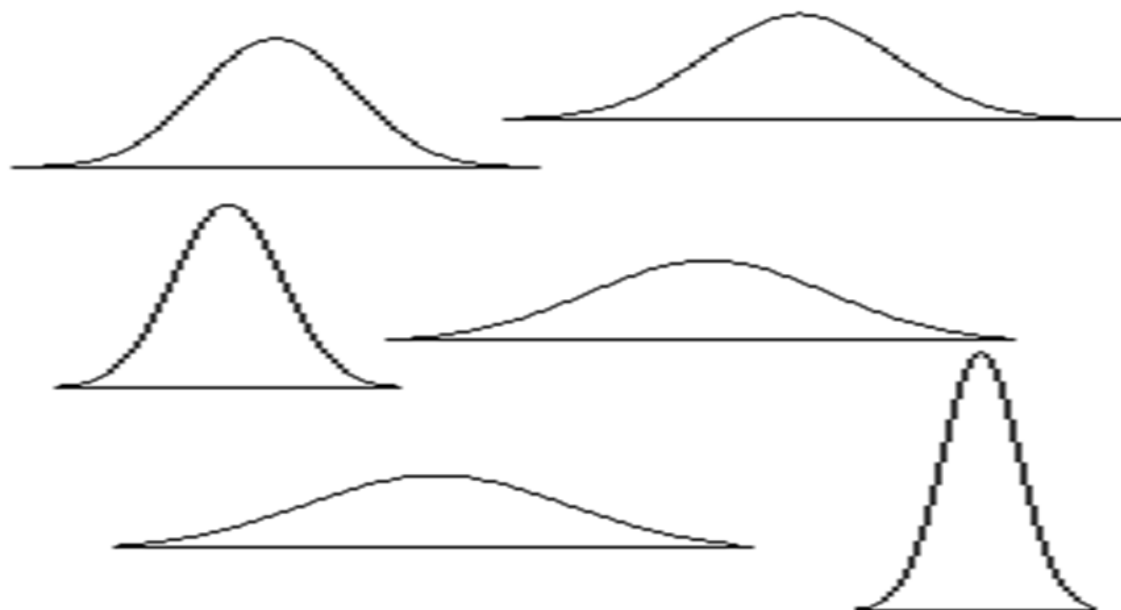
پاسخ

x	$(x - \bar{x})$	$\sum(x - \bar{x})^2$
۸۳	۲	۴
۷۵	-۶	۳۶
۸۱	۰	-
۷۹	-۲	۴
۷۱	-۱۰	۱۰۰
۹۵	۱۴	۱۹۶
۷۵	-۶	۳۶
۷۷	-۴	۱۶
۸۴	۳	۹
۹۰	۹	۸۱
$\bar{x} = 81 \quad n = 10$		جمع = ۴۸۲

$$S.D. = \sqrt{\frac{\sum(x - \bar{x})^2}{n-1}} = \sqrt{\frac{482}{10-1}} = \sqrt{\frac{482}{9}} = \sqrt{53.55} = 7.31$$

توزیع نرمال

توزیع نرمال یا «منحنی نرمال» یکی از مفاهیم مهم در تئوری آمار به شمار می‌آید. بیا بیا فرض کنیم مقدار هموگلوبین شمار زیادی از مردم را گردآوری کرده و شکل توزیع فراوانی آن را در حالی که میان طبقات فاصله‌ی کمی وجود دارد، ترسیم نموده‌ایم. در این صورت به احتمال زیاد خواهیم دید یک منحنی متقارن و هموار به دست آمده است. چنین منحنی را توزیع یا منحنی نرمال می‌نامند. شکل منحنی بستگی به میانگین و انحراف معیار داده‌ها دارد که خود بستگی به شمار و طبیعت مشاهده‌ها دارد؛ بنابراین می‌توان نتیجه گرفت که شمار بی‌شماری منحنی نرمال وجود دارد.



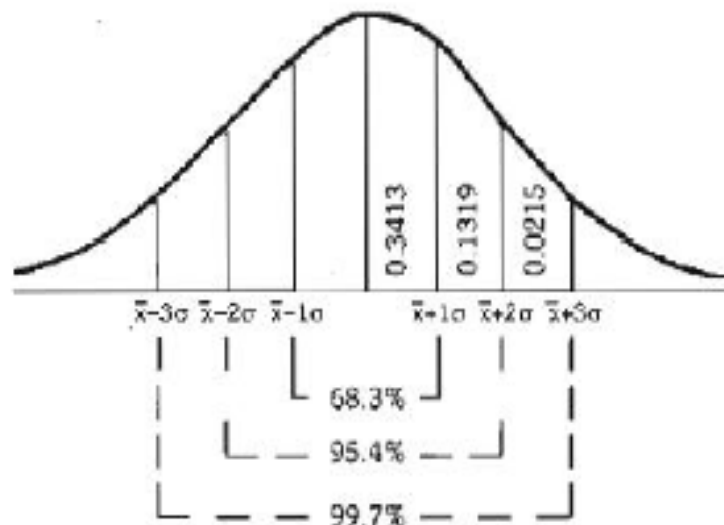
چند نکته در خصوص توزیع نرمال

دانستن برخی نکته‌ها درباره‌ی منحنی نرمال مفید است (شکل ۱۵):

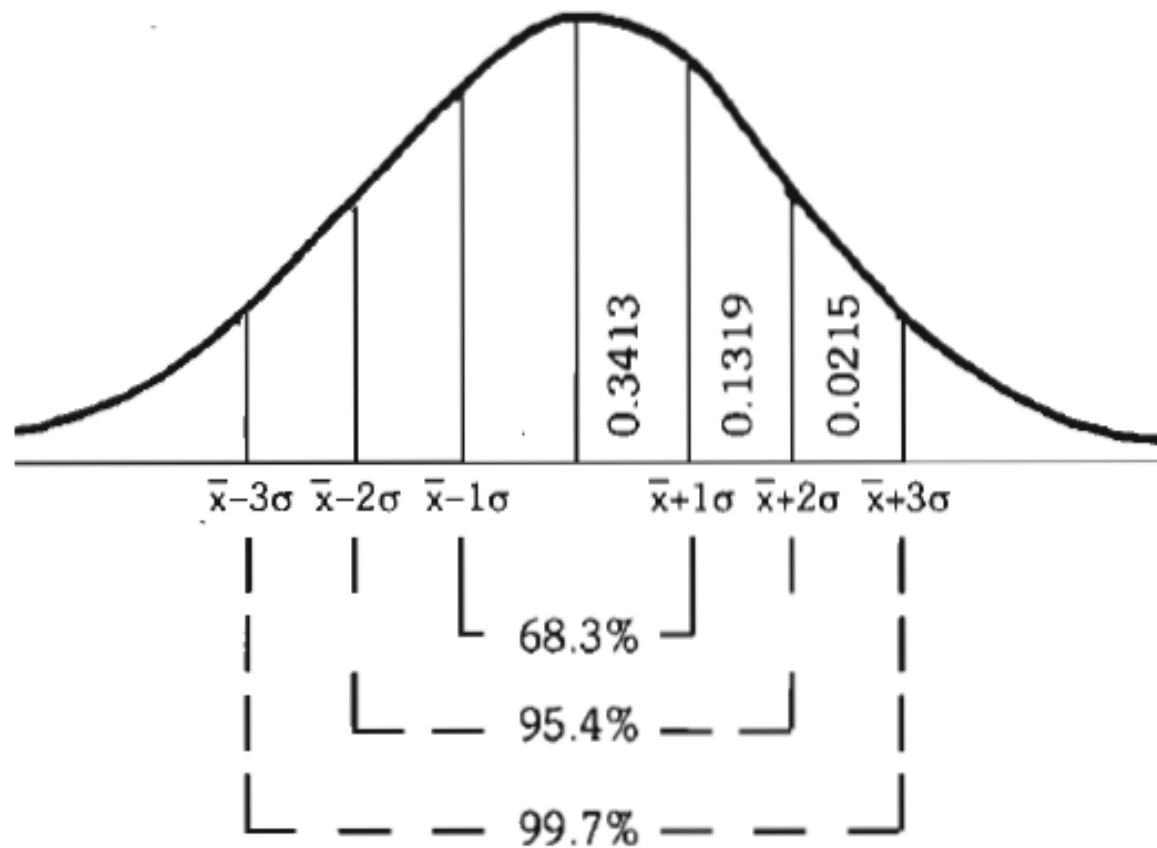
ا) مساحت محدود بین یک انحراف معیار $(\bar{x} \pm 1\sigma)$ از طرفین میانگین تقریباً ۶۸٪ از مقادیر را در توزیع شامل می‌شود،

ب) مساحت محدود بین دو انحراف معیار $(\bar{x} \pm 2\sigma)$ از طرفین میانگین اغلب مقادیر، یعنی تقریباً ۹۵٪ آن‌ها را در توزیع شامل می‌شود.

پ) مساحت محدود بین سه انحراف معیار $(\bar{x} \pm 3\sigma)$ از طرفین میانگین ۹۹/۷٪ از مقادیر را شامل می‌شود. حدود مذکور در طرفین میانگین را «حدود اطمینان» می‌نامند که در شکل ۱۵ نشان داده شده است.



شکل ۱۵. منحنی نرمال



حدود اطمینان ۹۵٪ ($\bar{x} \pm 2\sigma$) را در نظر بگیرید. این عبارت به ۹۵٪ از کل مساحت زیر منحنی نرمال اشاره می‌کند و به این معناست که ۹۵٪ از کل مشاهدات در فاصله ($\bar{x} \pm 2\sigma$) قرار می‌گیرند. پس احتمال این که یکی از مشاهدات خارج از حدود اطمینان ۹۵٪ قرار گیرد، یک بیستم است ($P=0.05$).

منحنی نرمال استاندارد

اگر چه با توجه به تفاوت در میانگین و انحراف معیار، بی شمار منحنی نرمال وجود دارد، اما فقط یک منحنی نرمال استاندارد (standard) می توان رسم کرد که توسط آمارشناسان تعریف شده است، به گونه ای که می توان به راحتی مساحت بین هر دو نقطه از سطح زیر منحنی نرمال را محاسبه کرد. منحنی نرمال استاندارد یک منحنی هموار، زنگوله ای شکل و کاملاً قرینه است که بر اساس شمار بی شماری از مشاهددهای به دست می آید. مساحت کل زیر منحنی برابر با یک، میانگین آن صفر و انحراف معیار آن یک است. در منحنی نرمال میانگین، میانه و نما روی هم می افتند. در این منحنی، فاصله ی میانگین ($\bar{X}$) از یک مقدار (X) تقسیم بر انحراف معیار «انحراف نسبی» یا «متغیر نرمال استاندارد» نامیده می شود که آن را با z نشان می دهند. متغیر نرمال استاندارد یا z از فرمول زیر محاسبه می شود:

$$z = \frac{(X - \bar{X})}{\sigma}$$

یک متغیر تصادفی (x) وقتی استاندارد خوانده می‌شود که میانگین آن صفر و انحراف معیار آن یک باشد. برای استاندارد کردن ابتدا میانگین x را از خود x کم کرده و حاصل تفاضل را بر σ یعنی انحراف معیار x تقسیم می‌کنیم. بنابراین $\frac{(x - \bar{x})}{\sigma}$ یک متغیر استاندارد شده است. از مفاهیم مهم آمار ریاضی این است که متغیر جدید z هم مثل متغیر x از توزیع نرمال پیروی می‌کند. میانگین توزیع متغیر جدید صفر و انحراف معیار (σ) آن یک است. آگاهی از مقدارهای سطح زیر منحنی نرمال استاندارد معمولاً نیاز است، که بر حسب مقدارهای مختلف، انحراف نسبی محاسبه می‌شود. نتایج محاسبه این مقادیر در جدول ۴ آمده است.

جدول ۴: مقادیر سطح زیر منحنی نرمال استاندارد با میانگین صفر و انحراف معیار یک، منبع (۲۰)

سطح زیر منحنی نرمال از مقدار میانگین تا نقطه مورد بررسی	انحراف نسبی (Z) $\frac{(x - \bar{x})}{\sigma}$
۰	۰
۰.۱۹۱۵	۰.۵
۰.۳۴۱۳	۱
۰.۴۳۳۲	۱.۵
۰.۴۷۷۲	۲
۰.۴۹۸۷	۳
۰.۴۹۹۹۷	۴
۰.۴۹۹۹۹۹۸	۵

برآورد احتمال (نمونه)

فرض کنید ضربان قلب گروهی از مردان سالم ۷۲ با انحراف معیار ۲ باشد. احتمال این که ضربان قلب مردی که شانس از این جمعیت انتخاب شده است، ۸۰ یا بیشتر از ۸۰ باشد، چقدر است؟

$$(Z) = \frac{(x - \bar{x})}{\sigma}$$

$$= \frac{80 - 72}{2} = 4$$

سطح زیر منحنی نرمال متناظر با انحراف نسبی ۴، برابر است با ۰/۴۹۹۹۷ (جدول ۴). از آن جایی که فقط با نصف سطح کل سروکار داریم (یعنی ۰/۵) بقیه سطح یعنی ۰/۴۹۹۹۷-۰/۵ یا ۰/۰۰۰۰۳ میزان احتمال مذکور را بیان می‌کند. به عبارت دیگر فقط ۳ نفر از ۱۰۰.۰۰۰ نفر احتمال دارد ضربان قلبی برابر با ۸۰ یا بیشتر داشته باشند.

جدول ۴، مقادیر سطح زیر منحنی نرمال استاندارد با میانگین صفر و انحراف معیار یک، منبع (۲۰)

سطح زیر منحنی نرمال از مقدار میانگین تا نقطه مورد بررسی	انحراف نسبی (Z) $\frac{(x - \bar{x})}{\sigma}$
۰	۰
۰.۱۹۱۵	۰.۵
۰.۳۴۱۳	۱
۰.۴۳۳۲	۱.۵
۰.۴۷۷۲	۲
۰.۴۹۱۷	۳
۰.۴۹۹۹۷	۴
۰.۴۹۹۹۹۸	۵

خطای معیار (Standard error)

اگر از یک جمعیت نمونه‌ای به حجم (n) انتخاب و نمونه‌گیری به گونه‌ای مشابهی بارها تکرار شود، هر نمونه میانگین $(\bar{x})$ خاص خود را خواهد داشت. چنانچه منحنی توزیع فراوانی میانگین‌های نمونه‌ای را رسم کنیم، می‌بینیم که منحنی دارای توزیع تقریباً نرمال خواهد بود و عملاً میانگین‌های نمونه‌ای برابر با میانگین خود جمعیت (μ) است. این مشاهده که میانگین‌های نمونه‌ای دارای توزیع نرمال در اطراف میانگین جمعیت هستند، موضوع بسیار مهمی است. انحراف معیار میانگین‌های نمونه‌ای که معیاری برای سنجش خطای نمونه‌گیری است و با فرمول $\sigma / \sqrt{n}$ محاسبه می‌شود، خطای معیار یا خطای معیار میانگین نامیده می‌شود. از آنجایی که میانگین نمونه‌ها از توزیع نرمال پیروی می‌کند به سادگی مشاهده‌پذیر است که ۹۵٪ از میانگین‌های نمونه‌ای در محدوده دو انحراف معیار $[\mu \pm 2(\sigma / \sqrt{n})]$ از طرفین میانگین جمعیت قرار می‌گیرند. بنابراین خطای معیار (S.E) شاخصی است که با استفاده از آن می‌توانیم قضاوت نماییم که آیا میانگین یک نمونه‌ی خاص درون حدود اطمینان واقع شده است یا نه.

آزمون‌های معنی‌داری

- در واقع **خطای معیار** نشان می‌دهد تا چه حد **بر آورد** یک میانگین ممکن است درست باشد. لذا با کمک آن می‌توان **آزمون‌های معنی‌داری** یا همان (Tests of significance) را انجام داد.
- مفهوم خطای معیار با استفاده از فرمول‌های مناسب برای همه‌ی آماره‌ها (statistics) کاربرد دارد که برخی از پرکاربردترین آنها عبارتند از:

(الف) خطای معیار میانگین

(ب) خطای معیار نسبت

(ج) خطای معیار تفاوت دو میانگین

(د) خطای معیار تفاوت دو نسبت

خطای معیار میانگین

- پیش از این درباره‌ی معنای "خطای معیار میانگین" که به آن به طور ساده **خطای معیار** هم گفته می‌شود و نیز در خصوص توزیع میانگین‌های نمونه‌ای در اطراف میانگین واقعی جمعیت مطالبی بیان شد.
 - در عمل معمولاً نمونه‌گیری‌های تکراری از جمعیت انجام نمی‌دهیم، بلکه با گرفتن فقط یک نمونه از جمعیت، میانگین و خطای معیار را محاسبه می‌کنیم.
- ❖ در این جا سوالاتی به این شرح مطرح می‌شود: میانگین نمونه‌ی گرفته‌شده تا چه اندازه صحیح است؟ چه اظهار نظری درباره‌ی میانگین واقعی جمعیت می‌توان کرد؟
- ❖ برای پاسخ به این سوالات باید **خطای معیار میانگین** محاسبه شود و بر اساس آن "حدود اطمینانی" که میانگین واقعی جمعیت (که فقط نمونه‌ای از آن مورد مطالعه قرار گرفته‌است) محتمل است در آن قرار گیرد، تعیین شود.

مثال از محاسبه خطای معیار میانگین

- مثال: فرض کنید در مطالعه‌ی نمونه‌ای تصادفی شامل ۲۵ مرد ۲۰ تا ۲۴ ساله ، میانگین درجه حرارت بدن آن‌ها ۹۸/۱۴ درجه‌ی فارنهایت با انحراف معیار ۰/۶ به دست آمده باشد. در مورد میانگین درجه‌ی حرارت بدن افراد متعلق به جامعه‌ای که این نمونه از آن انتخاب شده‌است چه اظهار نظری می‌توان کرد؟

برای پاسخ به این پرسش از خطای معیار کمک می‌گیریم: $S.E.\bar{x} = S / \sqrt{n} = 0.6 / \sqrt{25} = 0.12$
حال می‌توانیم حدود اطمینان را بر اساس منحنی توزیع نرمال به دست آوریم. اگر حدود اطمینان در فاصله‌ی دو خطای معیار از میانگین محاسبه شود (یعنی حدود اطمینان ۹۵٪) دامنه‌ی میانگین جمعیت بین $98.14 - (2 \times 0.12) = 97.90$ تا $98.14 + (2 \times 0.12) = 98.38$ خواهد بود. احتمال این که میانگین جامعه خارج از این محدوده باشد، فقط ۱ در ۲۰ ($P=0.05$) است.

خیلی وقت‌ها در متون علمی به واژه‌ی معنی‌دار (significant) برخورد می‌کنیم. زمانی می‌گوییم «این تفاوت معنی‌دار است» منظورمان این است که بعید است این تفاوت شانس رخ داده باشد. به عنوان یک قرارداد، هنگامی از کلمه‌ی معنی‌دار استفاده می‌کنیم که $P < 0.05$ (۱ در ۲۰) باشد و هنگامی که $P < 0.01$ (۱ در ۱۰۰) باشد، یعنی معنی‌داری بیشتری وجود دارد.

حدود معنی داری

حدود معنی داری	$(N.D.) = \frac{x - \mu}{\sigma / \sqrt{n}}$ انحراف نرمال	حدود اطمینان
$P < 0.05$ معنی دار در سطح کمتر از ۵٪	$N.D. > 2$	μ بیرون از حدود اطمینان ۹۵٪ قرار دارد
$P = 0.05$ معنی دار درست در سطح ۵٪	$N.D. = 2$	μ درست روی حدود اطمینان ۹۵٪ قرار دارد
$P > 0.05$ در سطح ۵٪ معنی دار نیست	$N.D. < 2$	μ درون حدود اطمینان ۹۵٪ قرار دارد

خطای معیار نسبت

- فرض کنید نسبت مردان در جمعیت ساکن در یک روستا ۵۲ درصد باشد. از این روستا نمونه‌ای تصادفی شامل ۱۰۰ نفر گرفته‌ایم و نسبت مردان در این نمونه ۴۰ درصد است. از این نمونه چه نتیجه‌ای می‌گیریم؟ از یک نمونه ۱۰۰ تایی، دامنه‌ی احتمالی قابل انتظار برای نسبت مردان با حدود اطمینان ۹۵ درصد چقدر است؟

- در این نوع مثال‌ها دیگر نه با میانگین‌ها، بلکه با نسبت‌ها در نمونه و در جامعه اصلی سروکار داریم. این نسبت‌ها را می‌توانیم با علامت p و q نشان دهیم و انحراف معیار را حول و حوش نسبت موردانتظار، یعنی ۵۲ درصد محاسبه نماییم که به آن خطای معیار نسبت می‌گویند و از فرمول زیر محاسبه می‌شود:

$$\text{خطای معیار نسبت} = \sqrt{\frac{pq}{n}}$$

که در آن: حجم نمونه = n ؛ نسبت زنان = q ؛ نسبت مردان = p

$$\text{معیار خطای نسبت} = \sqrt{\frac{52 \times 48}{100}} = 5$$

- دو خطای معیار را در طرفین نسبت ۵۲ درصد به عنوان ملاک در نظر می گیریم. به این معنی که اگر واقعا نمونه‌ی انتخاب شده نماینده‌ی خوبی برای جامعه اصلی باشد انتظار داریم نسبت برآورد شده از آن یک عدد تصادفی در محدوده‌ی $۵۲ \pm ۲(۵) = ۴۲$ تا $۵۲ \pm ۲(۵) = ۶۲$ درصد باشد.
- از آن جایی که نسبت مردان در نمونه مورد مطالعه فقط ۴۰ درصد و کاملاً خارج از حدود اطمینان بوده‌است، تفاوت بین مقادیر نسبت مشاهده شده و نسبت مورد انتظار "معنی دار" تلقی می شود و این تفاوت نمی تواند شانس باشد. در این مورد خاص، *انحراف نسبی* بر حسب خطای معیار، یعنی ۵ برابر خواهد بود با:

$$\text{انحراف نسبی} = \frac{۵۲ - ۴۰}{۵} = ۲/۴$$

حدود معنی داری	$(N.D.) = \frac{x - \mu}{\sigma / \sqrt{n}}$ انحراف نرمال	حدود اطمینان
$P < ۰/۰۵$ معنی دار در سطح کمتر از ۵٪	$N.D. > ۲$	μ بیرون از حدود اطمینان ۹۵٪ قرار دارد
$P = ۰/۰۵$ معنی دار درست در سطح ۵٪	$N.D. = ۲$	μ درست روی حدود اطمینان ۹۵٪ قرار دارد
$P > ۰/۰۵$ در سطح ۵٪ معنی دار نیست	$N.D. < ۲$	μ درون حدود اطمینان ۹۵٪ قرار دارد

- چون مقدار انحراف نسبی از عدد ۲ بیشتر شده است ، لذا انحراف معنی دار می باشد.
- چنین آزمونی در مواردی کاربرد پیدا می کند که جمعیت فقط شامل دو گروه یا دو نسبت باشد ، برای مثال زنان و مردان، بیماران و افراد سالم ، موفقیت ها و شکست ها و مشابه آن.

خطای معیار تفاوت دو میانگین

- در پژوهش‌های زیست‌شناختی زیاد پیش می‌آید که پژوهشگر با مشکل مقایسه میان دو گروه شاهد و تجربه مواجه شود. سوالی که در این موارد مطرح است این است که آیا تفاوت بین دو گروه به اندازه‌ای معنی‌دار هست که نشان دهد این دو گروه متعلق به دو جمعیت متفاوت می‌باشند. پاسخ این سوال با محاسبه‌ی **خطای معیار تفاوت دو میانگین** داده خواهد شد.

فرض کنید می‌خواهیم تأثیر یک داروی خاص را بر موش‌ها آزمایش کنیم. ۲۴ موش که از هر لحاظ با هم مقایسه‌پذیر هستند شانسی به دو گروه تقسیم می‌شوند. گروه (ا) را به عنوان گروه شاهد در نظر می‌گیریم به این معنی که موش‌های این گروه هیچ درمانی دریافت نمی‌کنند. ولی موش‌های گروه (ب) که همان گروه آزمایش‌اند، برای مدت مشخصی زیر درمان با داروی مورد بررسی قرار می‌گیرند. در پایان، همه‌ی موش‌ها کشته شده، و وزن کلیه‌ی هر یک از آن‌ها بر حسب میلی‌گرم اندازه‌گیری می‌شود. تأثیر درمان بر وزن کلیه‌ها به شرح زیر بوده است:

گروه	شمار	میانگین	انحراف معیار
گروه شاهد	۱۲	۳۱۸	۱۰/۲
گروه تجربه	۱۲	۳۷۰	۲۴/۱

حال بیاید ببینیم تفاوت دیده شده بین وزن کلیه‌های دو گروه، معنی‌دار بوده است یا نه. برای محاسبه‌ی خطای معیار تفاوت دو میانگین از این فرمول به شرح زیر استفاده می‌شود:

$$S.E.(d) = \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$$
$$= \sqrt{\frac{(10/2)^2}{12} + \frac{(24/1)^2}{12}} = \sqrt{8/67 + 48/4} = \sqrt{57/70} = 7/5$$

به این ترتیب خطای معیار تفاوت دو میانگین ۷/۵ به دست می‌آید. تفاوت واقعی میان میانگین‌ها $52 = (370 - 318)$ می‌باشد که بیشتر از دو برابر خطای معیار تفاوت دو میانگین است و بنابراین "معنی‌دار" می‌باشد. پس نتیجه می‌گیریم درمان بر وزن کلیه‌ها موثر بوده است.

خطای معیار تفاوت دو نسبت

- گاهی اوقات به جای تفاوت دو میانگین، معنی داری تفاوت بین دو نسبت یا کسر (ratio) آزمون می شود که آیا تفاوت قابل مشاهده بین دو نسبت یا کسر شانس است یا نه. در این موارد **خطای معیار تفاوت دو نسبت** را محاسبه می کنیم.

فرض کنید قصد داریم در یک کارآزمایی میدانی دو نوع واکسن سیاه سرفه را با هم مقایسه کنیم. نتایج مطالعه بدین شرح است:

واکسن	شمار افراد واکسینه	شمار افراد مواجهه یافته	شمار موارد بیماری	میزان حمله ^۱
آ	۲۴۰۰	۹۰	۲۲	۲۴/۴٪
ب	۲۳۰۰	۸۶	۱۴	۱۶/۲٪

از نتایج مزبور این گونه به نظر می رسد که واکسن (ب) بهتر از واکسن (آ) است. آیا این تفاوت، معنی دار است؟ برای پاسخ به این پرسش ابتدا خطای معیار تفاوت دو نسبت را به ترتیب زیر محاسبه می کنیم:

$$S.E. = \sqrt{\frac{P_1 Q_1}{n_1} + \frac{P_2 Q_2}{n_2}}$$
$$= \sqrt{\frac{24/4 \times 75/6}{90} + \frac{16/2 \times 83/8}{86}} = \sqrt{20/49 + 15/78} = \sqrt{36/27} = 6/0.2$$

خطای معیار تفاوت دو نسبت ۶ است در حالی که تفاوت مشاهده شده (۲۴/۴-۱۶/۲) یعنی ۸/۲ است. به عبارت دیگر تفاوت دیده شده بین دو گروه کمتر از دو برابر خطای معیار تفاوت بین دو نسبت یعنی ۲×۶ است. پس استنتاج می شود که شواهد قوی دال بر تفاوت میان تأثیر دو واکسن وجود ندارد. بنابراین تفاوت دیده شده می تواند شانسی باشد. در این موردها از آزمون مجذورکای (χ^2) نیز می توان بهره برداری نمود.

آزمون مجذور کای (CHI-SQUARE)

آزمون مجذور کای (χ^2) روشی جایگزین برای آزمون تفاوت بین دو نسبت است. مزیتش این است که برای انجام مقایسه نسبت‌ها در بیش از دو گروه نیز قابل استفاده است. همان نمونه پیش را در نظر بگیرید: دو واکسن سیاه سرفه را در یک کارآزمایی میدانی با هم مقایسه کرده‌ایم. نتایج بدین قرار هستند:

واکسن	شمار افرادی که دارای حمله‌ی بیماری بوده‌اند	شمار افرادی که بدون حمله‌ی بیماری بوده‌اند	جمع	میزان حمله
آ	۲۲	۶۸	۹۰	۲۴/۴٪
ب	۱۴	۷۲	۸۶	۱۶/۲٪
جمع	۳۶	۱۴۰	۱۷۶	–

در ظاهر ملاحظه می‌شود که واکسن (ب) بهتر از واکسن (آ) بوده‌است. اما این پرسش وجود دارد که آیا واکسن ب از واکسن آ بهتر است، یا تفاوت مشاهده‌پذیر فقط از روی شانس بوده‌است.

واکسن	شمار افراد واکسینه	شمار افراد مواجهه یافته	شمار موارد بیماری	میزان حمله ^۱
آ	۲۴۰۰	۹۰	۲۲	۲۴/۴٪
ب	۲۳۰۰	۸۶	۱۴	۱۶/۳٪

۱. آزمون فرضیه‌ی صفر^۱

در ابتدا فرضیه‌ای را تنظیم می‌کنیم که به نام فرضیه‌ی صفر (Null Hypothesis) خوانده می‌شود. این فرضیه می‌گوید که هیچ تفاوتی بین دو واکسن وجود ندارد. فرضیه مذکور باید به گونه کمی آزمون شود. آزمون به شکل زیر اجرا می‌شود:

ابتدا نتایج مربوط به دو واکسن را گردآوری می‌کنیم. نسبت کسانی که مورد حمله بیماری قرار گرفته‌اند $۰/۲۰۴ = ۳۶ \div ۱۷۶$ و نسبت کسانی که مورد حمله بیماری قرار نگرفته‌اند $۰/۷۹۵ = ۱۴۰ \div ۱۷۶$ است. با بهره‌گیری از این نسبت‌ها، شمار مورد انتظار بیماران یا همان موارد حمله و نیز شمار مورد انتظار کسانی که دچار حمله‌ی بیماری نشده‌اند را می‌توان محاسبه کرد. به این ترتیب که برای واکسن آ، شمار مورد انتظار برای کسانی که دارای حمله‌ی بیماری بوده‌اند $۱۸/۳۶ = ۹۰ \times ۰/۲۰۴$ و برای کسانی که حمله‌ی بیماری نداشتند $۷۱/۵۵ = ۹۰ \times ۰/۷۹۵$ محاسبه می‌شود. به همین صورت برای واکسن ب، شمار بیماران مورد انتظار $۱۷/۵۴۴ = ۸۶ \times ۰/۲۰۴$ و شمار کسانی که حمله‌ی بیماری نداشتند $۶۸/۳۷ = ۸۶ \times ۰/۷۹۵$ برآورد می‌شود. حال می‌توان جدول پیشین را دوباره نوشت. در این جدول شمار مشاهده‌ها (O)، موارد انتظار (E) و تفاضل این دو مقدار (O-E)^۲ در هر خانه از جدول آمده است.

واکسن	افرادی که دارای حمله‌ی بیماری بوده‌اند	افرادی که حمله‌ی بیماری نداشتند
آ	O = ۲۲	O = ۶۸
	E = ۱۸/۳۶	E = ۷۱/۵۵
	O - E = +۳/۶۴	O - E = -۳/۵۵
ب	O = ۱۴	O = ۷۲
	E = ۱۷/۵۴	E = ۶۸/۳۷
	O - E = -۳/۵۴	O - E = +۳/۶۳

۲. به کارگیری آزمون مجذور کای

$$\chi^2 = \frac{\sum(O - E)^2}{E}$$
$$= \frac{(3/64)^2}{18/36} + \frac{(3/55)^2}{71/55} + \frac{(3/54)^2}{17/54} + \frac{(3/63)^2}{68/37}$$
$$= 0.72 + 0.17 + 0.71 + 0.19 = 1.79$$

۳. یافتن درجه‌ی آزادی^۱

گام بعدی، محاسبه درجه‌ی آزادی (d.f.) است. درجه‌ی آزادی (Degree freedom) به شمار سطرها و ستون‌های جدول اصلی بستگی دارد و از فرمول زیر به دست می‌آید:

$$d.f. = (c - 1)(r - 1)$$

شمار ستون‌ها = C

شمار سطرها = r

جدول مزبور دو ستون (افرادى که داراي حمله‌ی بیماری بوده‌اند و افرادى که بدون حمله‌ی بیماری بوده‌اند) و دو سطر (واکسن آ و واکسن ب) دارد. این جدول یک جدول توافقی ۲×۲ است و درجه‌ی آزادی آن برابر است با:

$$d.f. = (c - 1)(r - 1) = (2 - 1)(2 - 1) = 1$$

استفاده از جدول مجذور کای

۳. جداول احتمال

مرحله‌ی بعدی مراجعه به جدول احتمال است که برای این منظور از قبل آماده شده و نمونه‌ای از آن در آخر همین فصل آورده شده است. با مراجعه به این جدول می‌بینیم در صورتی که درجه‌ی آزادی ۱ باشد، مقدار χ^2 برای احتمال ۰/۰۵، برابر با ۳/۸۴ است. از آن جایی که مقدار مشاهده شده (۱/۷۹) بسیار کوچک‌تر از ۳/۸۴ است، چنین نتیجه‌گیری می‌شود که فرضیه صفر درست است و واکسن ب از واکسن آ بهتر نیست. این آزمون فقط هنگامی معتبر است که هیچ‌یک از اعداد مورد انتظار در خانه‌های جدول از عدد ۲ کمتر نباشد.

Probability (P)

D.F.	.50	.10	.05	.02	.01	.005	.001
1.	0.45	2.71	3.84	5.41	6.64	7.88	10.83
2.	1.39	4.61	5.99	7.82	9.21	10.60	13.82
3.	2.37	6.25	7.82	9.84	11.34	12.64	16.27
4.	3.36	7.78	9.49	11.67	13.28	14.86	18.47
5.	4.35	9.24	11.07	13.39	15.09	16.75	20.51
6.	5.35	10.65	12.59	15.03	16.81	18.55	22.46
7.	6.35	12.02	14.07	16.62	18.48	20.28	24.32
8.	7.34	13.36	15.51	18.17	20.09	21.96	26.13
9.	8.34	14.68	16.92	19.68	21.67	23.59	27.88
10.	9.34	15.99	18.31	21.16	23.21	25.19	29.59

مفهوم همبستگی

- گاهی اوقات می‌خواهیم بدانیم آیا بین دو متغیر مثلاً وزن و قد، دمای بدن و ضربان نبض، سن و ظرفیت حیاتی ریه و مواردی مانند این‌ها **ارتباط خطی** وجود دارد یا خیر.

□ برای پاسخ به این سوال که آیا بین دو متغیر (مثلاً x و y) ارتباط معنی‌داری وجود دارد یا نه، باید **ضریب همبستگی** (*Coefficient Correlation*) را محاسبه کنیم که با حرف "r" نمایش داده می‌شود.

- فرض کنید دو متغیر به نام‌های x و y داریم که هر دوی آن‌ها برای افراد یک نمونه به حجم n بررسی شده‌اند. ضریب همبستگی با استفاده از فرمول ذیل محاسبه می‌شود:

$$r = \frac{\sum (x - \bar{x})(y - \bar{y})}{\sqrt{\sum (x - \bar{x})^2 \sum (y - \bar{y})^2}}$$

تفسیر ضریب همبستگی

- مقدار ضریب همبستگی بین $+1$ و -1 تغییر می کند .
- اگر ضریب همبستگی به $+1$ نزدیک تر باشد، نشان می دهد رابطه ی بین x و y مثبت و قوی است، به این معنی که اگر یکی از متغیرها افزایش یابد دیگری هم افزایش می یابد.
- اگر نزدیک به -1 باشد ، نشان دهنده ی رابطه ی منفی و قوی بین دو متغیر است ؛ یعنی با افزایش یکی دیگری کاهش می یابد.
- اگر $r=0$ باشد ، نشان دهنده ی این است که بین x و y رابطه ی خطی وجود ندارد.
- آزمون هایی وجود دارند که با استفاده از آن ها می توان نشان داد آیا همبستگی مشاهده شده شانسی است یا نه.

لازم است اشاره شود وجود همبستگی الزاما علیت (Causation) را ثابت نمی کند.

مفهوم رگرسیون

در مورد یک فرد اگر بخواهیم به مقدار نامعلوم یک متغیر [متغیر وابسته]، از روی مقدار معلوم یک متغیر دیگر [متغیر مستقل] پی ببریم، باید ضریب رگرسیون^۱ یکی را نسبت به دیگری حساب کنیم. به گونه قراردادی متغیر مستقل را با x و متغیر وابسته را با y نشان می دهند. فرمول محاسبه ضریب رگرسیون به شرح زیر است:

$$y = \bar{y} + b(x - \bar{x})$$

$$\bar{y} = \text{میانگین } Y_1 + Y_2 + Y_3 \dots Y_n \quad \text{که در آن:}$$

$$\bar{x} = \text{میانگین } X_1 + X_2 + X_3 \dots X_n$$

$$b = \frac{\sum (x - \bar{x})(y - \bar{y})}{\sum (x - \bar{x})^2}$$

مقدار b به نام ضریب رگرسیون y نسبت به x خوانده می شود رگرسیون x نسبت به y را نیز با همین روش می توان حساب کرد:

$$x = \bar{x} + b'(y - \bar{y})$$

$$b' = \frac{\sum (x - \bar{x})(y - \bar{y})}{\sum (y - \bar{y})^2}$$

که در آن b' ضریب رگرسیون x نسبت به y است. کاربرد رگرسیون، محاسبه ی میانگین برآوردهای مقدار یک متغیر از روی مقدار متغیر دیگر است.

خلاصه مطالب

QUICK SUMMARY



■ در هر پژوهش، ابتدا داده‌ها جمع‌آوری می‌شوند. داده‌هایی که مستقیم از افراد مورد بررسی به‌دست می‌آیند، داده‌های ابتدایی و آن‌هایی که از منابع خارجی به‌دست می‌آیند، داده‌های بعدی (دومین داده‌ها) نامیده می‌شوند. نخستین داده‌ها معمولاً از داده‌های بعدی دقیق‌تر هستند. پس از جمع‌آوری داده‌ها، بایستی آن‌ها را هدفمندانه مرتب و خلاصه نمود تا استخراج نکته‌های مهم به‌گونه‌ی شفاف و کامل از آن‌ها ممکن شود.

■ با جدول می‌توان به‌سادگی انبوهی از داده‌های آماری را به نمایش گذارد. در یک جدول توزیع فراوانی، ابتدا داده‌ها به گروه‌های مناسب تقسیم شده و شمار افراد هر گروه، در ستون کنار آن نمایش داده می‌شود. نمودارها و طرح‌ها در مقایسه با جدول‌ها ساده‌ترند و تأثیر نیرومندتری بر ذهن می‌گذارند. از میان آن‌ها می‌توان به انواع نمودارهای ستونی، دایره‌ای، هیستوگرام، چندگوشه‌ی فراوانی، طرح نقاط پراکنده و طرح‌های خطی و تصویری اشاره کرد. وقتی داده‌های آماری مربوط به مناطق جغرافیایی باشند، می‌توان آن‌ها را با بهره‌برداری از نقشه‌های هاشور زده یا نقطه‌گذاری شده نیز نشان داد.

■ انواع گوناگونی از معدل‌های آماری وجود دارد که رایج‌ترین آن‌ها میانگین، میانه و نما هستند. برای محاسبه‌ی میانگین ابتدا اندازه‌ی همه مشاهده‌ها با هم جمع و سپس مجموع به‌دست آمده بر کل مشاهدات بخش می‌شود. برای به‌دست آوردن میانه، ابتدا داده‌ها بر اساس مقدارشان مرتب می‌شوند و سپس مشاهده‌ای که از نظر مقدار در وسط قرار گرفته‌است، میانه نام می‌گیرد. نما به پرتکرارترین مشاهده در یک توزیع داده‌ای نسبت داده می‌شود.

■ چندین مقدار برای اندازه‌گیری پراکندگی وجود دارد که شناخته شده‌ترین آن‌ها دامنه، میانگین انحرافات و انحراف معیار هستند. دامنه عبارت است از اختلاف بین بیشترین و کمترین مقدارها در یک نمونه. میانگین انحرافات، معدل مقادیر انحرافات هر مشاهده از میانگین داده‌هاست. انحراف معیار که رایج‌ترین شاخص پراکندگی به شمار می‌آید، ریشه‌ی دوم میانگین مجذور انحرافات است.

■ توزیع یا منحنی نرمال با شکل متقارن و زنگوله مانند خود یکی از مفاهیم مهم و پرکاربرد آماری به شمار می‌رود. در این منحنی میانگین، میانه و نما روی هم می‌افتند. مساحت محدود بین یک، دو و سه انحراف معیار از طرفین میانگین به ترتیب تقریباً ۶۸، ۹۵ و ۹۹/۷٪ از مقادیر را در توزیع شامل می‌شود. حدود مزبور در طرفین میانگین را حدود اطمینان می‌نامند.

■ منحنی نرمال استاندارد نوع خاصی از توزیع نرمال است که در آن مساحت کل زیر منحنی برابر با یک، میانگین آن صفر و انحراف معیار آن یک است. در این منحنی، فاصله‌ی میانگین از یک مقدار بخش بر انحراف معیار، متغیر نرمال استاندارد نامیده می‌شود.

■ وقتی قرار باشد شمار زیادی از مردم یا اشیاء یا هر واحد دیگر مورد بررسی قرار گیرند، نمونه‌گیری می‌کنیم. چارچوب نمونه‌گیری فهرستی از افراد جمعیت مورد نظر است که قرار است نمونه از بین آن‌ها انتخاب شود. روش نمونه‌گیری باید به گونه‌ای باشد که در آن نمونه‌ی انتخاب شده نماینده خوبی از کل جمعیت مورد مطالعه باشد. رایج‌ترین روش‌های نمونه‌گیری عبارت‌اند از تصادفی ساده، تصادفی منظم و تصادفی طبقه‌بندی شده.

■ اگر از یک جمعیت همواره نمونه‌گیری شود، نتیجه‌ی بررسی بر پایه هر نمونه متفاوت خواهد بود. به این تفاوت خطای نمونه‌گیری می‌گویند که تحت تأثیر حجم نمونه و تفاوت‌های طبیعی موجود میان اندازه‌گیری‌ها است. اگر از یک جمعیت با یک حجم معین بارها نمونه‌گیری انجام شود، هر نمونه میانگین خاص خود را خواهد داشت که توزیع آن‌ها نرمال و میانگین آن‌ها برابر با میانگین خود جمعیت خواهد بود.

■ انحراف معیار میانگین‌های نمونه‌ای که معیاری برای سنجش خطای نمونه‌گیری است، خطای معیار نامیده می‌شود. با برآورد خطای معیار برای میانگین، نسبت، تفاوت دو میانگین و یا تفاوت دو نسبت، می‌توان برای این مقادارها آزمون آماری انجام داد. آزمون مجذورکای روشی جایگزین برای آزمون تفاوت بین دو نسبت است که برای انجام مقایسه نسبت‌ها در بیش از دو گروه نیز قابل بهره‌برداری است.

■ برای بررسی ارتباط خطی میان دو متغیر باید ضریب همبستگی آن‌ها را محاسبه کنیم که مقدار آن بین $+1$ و -1 تغییر می‌کند. اگر ضریب همبستگی به $+1$ نزدیک‌تر باشد، نشان می‌دهد رابطه‌ی بین دو متغیر مثبت و توانمندتر است، به این معنی که اگر یکی از متغیرها افزایش یابد دیگری هم افزایش می‌یابد. اگر نزدیک به -1 باشد، نشان دهنده‌ی رابطه‌ی منفی و توانمند بین دو متغیر است. ضریب همبستگی صفر به معنای نبود همبستگی خطی میان آن دو است.

■ در مورد یک فرد اگر بخواهیم به مقدار نامعلوم یک متغیر، از روی مقدار معلوم یک متغیر دیگر پی ببریم، باید ضریب رگرسیون یکی را نسبت به دیگری حساب کنیم. کاربرد رگرسیون، محاسبه‌ی میانگین برآوردهای مقدار یک متغیر از روی مقدار متغیر دیگر است.

منبع:

درس نامه پزشکی پیشگیری و پزشکی اجتماعی

✓تالیف ک. پارک

- ویراست ۲۱
- فصل سوم جلد اول با عنوان:

“اطلاعات سلامت و آمار پایه پزشکی”

