

آمارزستی مقدماتی

تهیه و تنظیم : طیه معظمی

۱۳۹۲

بیش تر مردم با واژه آمار به مفهومی که جهت ثبت و نمایش اطلاعات به کار می رود و در یک مفهوم وسیع تر ارائه پاره ای از مشخصات عددی چون میانگین ، درصدها و ... است ، آشنا هستند. مواردی همانند تعداد بیماران مراجعه کننده به بیمارستان ، تعداد افراد مبتلا به یک بیماری خاص و ... ولی این مفهوم منطبق با موضوع اصلی مورد بحث آمار نیست. تعریف جامع تر آمار را می توان بصورت زیر بیان کرد.

آمار علمی است که مشخصات جامعه را از نظر کمی با در نظر گرفتن مشخص کننده های آن جامعه مورد بررسی قرار می دهد. در واقع آمار داده های عددی را جمع آوری ، نمایش و تحلیل می کند. در مرحله تحلیل آماری با مسئله قضاوت درباره فرضیه های مختلف مواجهیم. قضاوتهای آماری با قضاوتهایی که در علوم ریاضی بکار می رود متفاوت است. قضاوت آماری صرفاً بر اساس مشاهدات استوار است. هدف از آمار عمدتاً سه چیز است:

- ۱- توصیف و خلاصه کردن اطلاعات از طریق کم کردن آنها به یک سری داده های کوچک تر و معنادارتر. هنگامی که هدف ما خلاصه کردن داده هاست تعداد زیادی از داده های سازمان دهی نشده را می گیریم و آنها را به یک مجموعه کوچکتر می شکنیم تا بدون فدا شدن عناصر مهم به توصیف داده های اصلی پردازیم.
- ۲- پیشگویی و تعمیم دادن رخدادها بر اساس مشاهدات. هنگامی که هدف ما تعمیم در مورد رخ دادن داده ها باشد از آمار به عنوان ابزار استنباطی استفاده میشود. بدین وسیله می توانیم درجه اطمینان از صحت در یک زمینه تحقیقاتی خاص را بیان کنیم
- ۳- تعیین ارتباطات ، همبستگیها یا اختلافات بین یک سری مشاهده. وقتی می خواهیم ارتباط ، همبستگی یا اختلاف بین متغیرهای مورد نظر را تعیین میکنیم از دانش موجود در مورد یک سری داده برای پیشگویی یا استنباط ویژگیهای یک سری داده دیگر استفاده می کنیم.

تعریف آمار:

آمار علمی است که از روشهای جمع آوری ، مرتب کردن ، خلاصه نمودن ، نمایش اطلاعات و درنهایت تجزیه و تحلیل و استخراج نتایج آن ، در شرایط عدم حتمیت (قطعیت) بحث می کند. به عبارت دیگر آمار داده های عددی را جمع آوری ، نمایش و تحلیل می کند.

هرگاه از داده هایی که مربوط به علوم زیستی و پزشکی است در تجزیه و تحلیل آماری استفاده گردد، بدان "آمار زیستی" اطلاق می گردد. آمار زیستی شاخه ای از آمار (کاربردی) است که تمرکز و تأکید آن بر توسعه و استفاده از روشهای آماری است که در راستای حل مسائل و پاسخ به سؤالاتی که در بهداشت، پزشکی و ژنتیک و بیولوژی انسانی مطرح می شوند، که

می توان از آن در اطلاع از جدیدترین دستاوردهای پزشکی، کاربرد نتایج پژوهشها در درمان بیماران، شناخت شاخصهای زیستی، تفسیر وقایع اپیدمیولوژیک، هدایت و بررسی طرحهای پژوهشی و تحقیقاتی استفاده کرد.

جمع آوری داده ها:

برای بررسی داده ها از دو روش آمار توصیفی و آمار تحلیلی یا استنباطی استفاده می شود.

آمار استنباطی و آمار توصیفی:

۱- آمار توصیفی (**descriptive**) با خلاصه کردن داده ها به یک سری اصطلاحات قابل فهم تر و بدون از دست دادن یا مخدوش کردن اطلاعات به توصیف یا مشخص کردن داده ها می پردازد. در یک پژوهش جهت بررسی و توصیف ویژگیهای عمومی پاسخ دهندگان از روش های موجود در آمار توصیفی مانند جداول توزیع فراوانی، درصد فراوانی، تجمعی و میانگین استفاده میگردد. به طور کلی از سه روش در آمار توصیفی برای خلاصه سازی داده ها استفاده می شود:

- استفاده از جداول
- استفاده از نمودار
- محاسبه مقادیری خاص که نشان دهنده خصوصیات مهمی از داده ها باشند.

۲- آمار استنباطی (**inferential**): شامل یک سری روشهای آماری است که بر اساس اطلاعات نمونه درباره ویژگیهای جمعیت، پیشگویی هایی را فراهم می کند. چنانچه به جای مطالعه کل اعضای جامعه، بخشی از آن با استفاده از فنون **نمونه گیری** انتخاب شده، و مورد مطالعه قرار گیرد و بخواهیم نتایج حاصل از آن را به کل جامعه **تعمیم** دهیم از روش هایی استفاده می شود که موضوع آمار استنباطی (**statistics Inferential**) است. آن چه که مهم است این است که در گذر از آمار توصیفی به آمار استنباطی یا به عبارت دیگر از نمونه به جامعه بحث و نقش احتمال شروع می شود. در واقع احتمال، پل رابط بین آمار توصیفی و استنباطی به حساب می آید. نشان می دهد که شانس و تصادف تا چه حد می تواند روابط بین متغیرهای مشاهده شده را توجیه نماید.

تعاریف:

جامعه: در هر بررسی آماری، مجموعه عناصر مورد نظر را جامعه می نامند. به عبارت دیگر، جامعه مجموعه تمام مشاهدات ممکن است که می توانند با تکرار یک آزمایش حاصل شوند به طور کلی جامعه عبارت است از مجموعه کاملی از افراد یا اشیاء یا اجزاء که حداقل در یک صفت مورد علاقه مشترک باشند، گفته می شود.

سرشماری: سرشماری از جامعه متناهی، بررسی است که تمام واحدهای جامعه را دربرمی گیرد. در بسیاری از موارد، اجرای سرشماری در یک جامعه متناهی، کاری است شدنی.

نمونه: نمونه بخشی از جامعه تحت بررسی است که با روشی که از پیش تعیین شده است انتخاب می‌شود. به قسمی که می‌توان از این بخش، استنباطهایی درباره کل جامعه بدست آورد.

پارامتر و آماره: به مقادیری که خصوصیات جامعه را توصیف می‌کنند، پارامتر گفته می‌شود. به مقادیری هم که خصوصیات نمونه را توصیف می‌کنند، آماره یا شاخص آماری می‌گویند. برای تمیز قائل شدن بین دو مفهوم پارامتر و شاخص آماری معمولاً پارامترها را با حروف یونانی و شاخص‌های آماری را با حروف لاتین نمایش می‌دهند. به عنوان مثال برای نمایش دادن میانگین جامعه از حرف یونانی (مو = μ) و برای نشان دادن میانگین گروه نمونه از حرف لاتین $\bar{X}$ (بخوانید ایکس بار) و برای نشان دادن واریانس جامعه از حرف یونانی σ^2 (مجذور زیگما) و برای نشان دادن واریانس نمونه از S^2 استفاده می‌شود.

برآورد و آزمون فرض: برآوردیابی و آزمون فرض دو روشی هستند که برای استنباط درمورد پارامترهای مجهول جامعه به کار می‌روند.

متغیرها و روش اندازه گیری آن :

افراد جامعه حداقل دارای یک صفت مشترک هستند. هر صفت که از یک فرد به فرد دیگر جامعه تغییر کند را متغیر می‌نامیم. بنابراین متغیرها عبارتند از ویژگیهای افراد، اشیاء، واحدها و غیره که می‌توانند مقادیر کیفی یا کمی اختیار کنند. از جمله این مقادیر می‌توان قد، وزن، بهره هوشی و امثال آن را نام برد.

به طور کلی داده‌های متغیرها به دو دسته تقسیم می‌شوند:

داده‌های کیفی (Qualitative data)، مانند رنگ چشم یا جنسیت، و داده‌های کمی (**Quantitative data**) مانند طول، قد، وزن و بهره هوشی.

داده‌های کیفی با **کلمات** و داده‌های کمی با **اعداد** نشان داده می‌شوند.

داده‌های کیفی خود به دو دسته دیگر تقسیم می‌شوند:

الف: داده‌های اسمی (**Nominal data**)

ب: داده‌های رتبه‌ای (**Ordinal data**)

داده‌های کمی نیز خود به دو دسته دیگر تقسیم می‌شوند:

الف: داده‌های فاصله‌ای (**Interval data**)

ب: داده‌های نسبتی (**Ratio data**)

بدین ترتیب خواهیم داشت:

کیفی	اسمی رتبه ای	انواع داده ها
کمی	فاصله ای نسبتی	

انواع مقیاس اندازه گیری متغیرها Measurement Scales

یک متغیر را می توان در سطوح مختلف اندازه گیری کرد. انتخاب سطح مناسب برای اندازه گیری متغیر مورد مطالعه باعث می شود که داده های مورد گردآوری گویای واقعیت مورد مطالعه باشند. به طور کلی چهار سطح یا مقیاس برای اندازه گیری متغیرها می توان منظور داشت: مقیاس اسمی، مقیاس رتبه ای، مقیاس فاصله ای و مقیاس نسبتی.

۱- مقیاس اسمی Nomial Scale

مقیاس اسمی برای اندازه گیری متغیرهای کیفی به کار می رود. این مقیاس شامل حداقل دو مقوله متمایز است که هیچگونه تقدم یا تأخر در آن وجود ندارد. به عبارت دیگر میان مقوله های مقیاس اسمی نمی توان ترتیب خاصی در نظر گرفت. مثلاً برای متغیر جنسیت، دو مقوله نمی توان ترتیب ویژه ای منظور داشت. داده هایی که بر حسب نام خود یا سنجش پاره ای از خصوصیات براساس منتسب نمودن آن شی یا فرد مورد مطالعه به یک گروه مشخص می شوند. به عنوان مثال در مطالعه گروه خونی یا جنسیت که با A, B, AB, O یا زن و مرد مشخص می گردد.

۲- مقیاس رتبه ای Ordinal scale

مقیاس رتبه ای برای اندازه گیری متغیرهایی بکار می رود که پیوسته بوده و تفاوت حالت های مختلف صفت متغیر فقط از نظر سلسله مراتب وضع افراد قابل نمایان ساختن باشد. مثلاً برای تعیین رتبه های کنکور می توان مقیاس رتبه ای را بکار برد و افراد را در مراتب: رتبه اول، دوم تا رتبه آخر دسته بندی کرد. از جمله متغیرهایی که مقیاس رتبه ای برای آنها بکار می رود متغیر نگرش می باشد.

۳- مقیاس فاصله ای Interval Scale

مقیاس فاصله ای مقیاسی است که به وسیله آن می توان متغیرهای کمی را که دارای مبدأ سنجش قراردادی هستند اندازه گیری کرد. به وسیله این مقیاس نه تنها می توان افراد را رتبه بندی کرد بلکه تفاوت آنها، از نظر صفت متغیر مورد

مطالعه، را نیز می‌توان معین کرد. اما این مقیاس دارای صفر مطلق نمی‌باشد. مثلاً برای اندازه‌گیری هوش، ارزیابی عملکرد از مقیاس فاصله‌ای استفاده می‌شود. چه متغیر هوش دارای مقدار صفر نمی‌باشد. در اندازه‌گیری درجه حرارت نیز بدلیل نبود صفر واقعی نمی‌توان مقدار حرارت جسمی که درجه حرارتش ۸۰ است دو برابر حرارت جسمی که درجه حرارتش ۴۰ است دانست.

۴- مقیاس نسبتی Ratio Scale

مقیاس نسبتی، که بالاترین سطح اندازه‌گیری است، مقیاسی است که دارای مبدأ صفر مطلق بوده و از فاصله‌های مساوی برخوردار است. برای هر دو مقدار این مقیاس می‌توان نسبتی را تعیین کرد که حاکی از بیشی مقدار صفت متغیر در یک فرد نسبت به فرد دیگر مورد مطالعه باشد. مثلاً در اندازه‌گیری سالهای خدمت کارکنان، می‌توان فردی یافت که تازه استخدام باشد و سابقه خدمتش صفر باشد. همچنین، مثلاً می‌توان نسبت سابقه خدمت دو نفر که یکی دارای سابقه خدمت ۱۰ سال و دیگری ۵ سال است را حساب کرد. در این مثال، نفر اول سابقه خدمتش دو برابر نفر دوم است و نسبت سابقه‌شان ۲ می‌باشد که مساوی نسبت سابقه خدمت دونفر دیگر از کارکنان است که سابقه خدمت آنان به ترتیب ۶ سال و ۳ سال می‌باشد. به عبارت دیگر مقیاس نسبتی علاوه بر دارا بودن ویژگیهای مقیاس فاصله‌ای دارای مبدأ واقعی (صفر مطلق) نیز می‌باشد.

داده های کمی را به صورت دیگر نیز می توان تقسیم بندی کرد:

الف - داده های کمی پیوسته (Continuous quantitative data)

ب - داده های کمی گسسته (Discrete quantitative data)

داده های کمی پیوسته: داده هایی که بتواند بین مقادیر خود همه اعداد حقیقی ممکن را اختیار کند. مانند وزن، طول قد و فشار خون. بین دو وزن ۸۵ و ۸۶ کیلوگرم میتوان وزن ۸۵/۵ کیلوگرم هم یافت. داده هایی که متغیر وزن را در پنج نفر بر حسب کیلوگرم توصیف می کنند، عبارتند از:

۶۵، ۸۲، ۳۰، ۴۳ و ۵۳

چون این داده ها عدد هستند پس داده های ما کمی اند و چون در میان وزن فردی که ۶۵ کیلوگرم وزن دارد با فرد دیگری که فقط یک کیلوگرم وزن بیشتری دارد (یعنی ۶۶ کیلوگرم وزن دارد) می توان افراد بسیاری با وزن هایی همچون ۶۵/۱، ۶۵/۲۲، ۶۵/۴۴۵ و ... را انتخاب نمود پس این داده ها کمی پیوسته اند.

داده های کمی گسسته: داده هایی که تنها بتواند مقادیر به خصوصی اختیار کند مانند تعداد فرزندان یک خانواده، تعداد دندانهای فاسد و تعداد تصادفات جاده‌ای در روز

در مثال دیگر داده هایی که متغیر تعداد ضربان قلب را در پنج نفر بر حسب تعداد در دقیقه توصیف می کنند عبارتند از: ۷۲، ۹۳، ۶۶، ۱۱۰ و ۵۹

چون این داده ها عدد هستند پس داده های ما کمی اند و چون در میان تعداد ضربان قلب فردی که ۷۲ ضربه در دقیقه را

داراست با فرد دیگری که فقط یک ضربه در دقیقه تعداد ضربان بالاتری دارد. یعنی (۷۳ ضربه در دقیقه ضربان قلب دارد) نمی توان فرد دیگری انتخاب نمود پس این داده های کمی گسسته اند.

نامگذاری متغیرها

در پژوهشهایی که با مشاهده متغیرها و تولید داده ها سروکار داریم باید متغیرها را نامگذاری کرد: این امر در آزمودن فرضیه ها، به ویژه در پژوهشهای آزمایشی، باید رعایت شود. که به دو گروه متغیر اشاره می کنیم.

– متغیر مستقل (Independent variable)

– متغیر وابسته (Dependent variable)

– انواع متغیرهای پژوهش

متغیرهای تحقیق براساس نوع مقادیری که می توانند اختیار کنند به دو دسته متغیرهای کمی و کیفی تقسیم می شوند اما انواع متغیرها براساس نقش آنها در تحقیق عبارتند از:

۱- **متغیر مستقل** : متغیر مستقل متغیری است که در پژوهش های تجربی به وسیله پژوهشگر دستکاری می شود تا تاثیر (یا رابطه) آن بر روی پدیده دیگری بررسی شود.

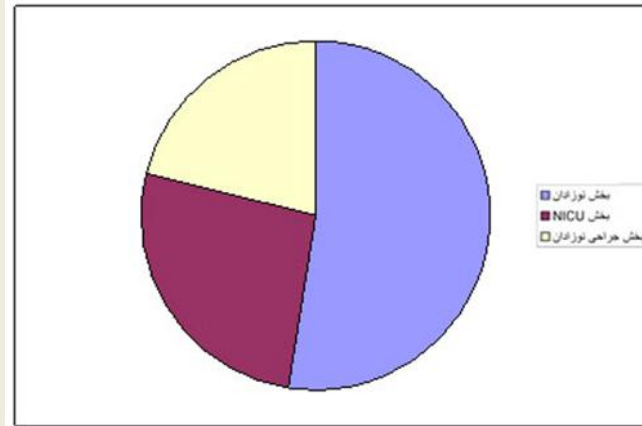
۲- **متغیر وابسته** : متغیر وابسته، متغیری است که تأثیر (یا رابطه) متغیر مستقل بر آن مورد بررسی قرار می گیرد. به عبارت دیگر پژوهشگر با دستکاری متغیر مستقل درصدد آن است که تغییرات حاصل را بر متغیر وابسته مطالعه نماید.

نمودارها (Diagrams)

برای نمایش تصویری داده ها از نمودارها استفاده می کنیم. نمودارها انواع زیادی دارند که در این قسمت به مهم ترین و شایع ترین آنها اشاره می کنیم

۱- نمودار دایره ای (Pie diagram)

نمودارهایی هستند که برای نمایش داده های کیفی به کار می روند. به عنوان مثال اگر بخواهیم فراوانی نوزادان بستری در یک بیمارستان را بر حسب بخش بستری نمایش دهیم در اینجا متغیر، بخش بستری است و اگر داده های بخش را بر حسب بخش نوزادان، بخش NICU و بخش جراحی نوزادان نامگذاری نمائیم داده های کیفی (از نوع اسمی) هستند. در این حالت برای نمایش تصویری داده ها می توان از نمودار دایره ای کمک گرفت و هر قطاع از دایره فراوانی یک بخش بستری را نشان خواهند داد. در چنین وضعیتی نمودار می تواند به صورت زیر باشد.



نمودار ۱ - فراوانی نوزادان بستری در یک بیمارستان بر حسب بخش بستری

۲- نمودار میله ای (Bar diagram)

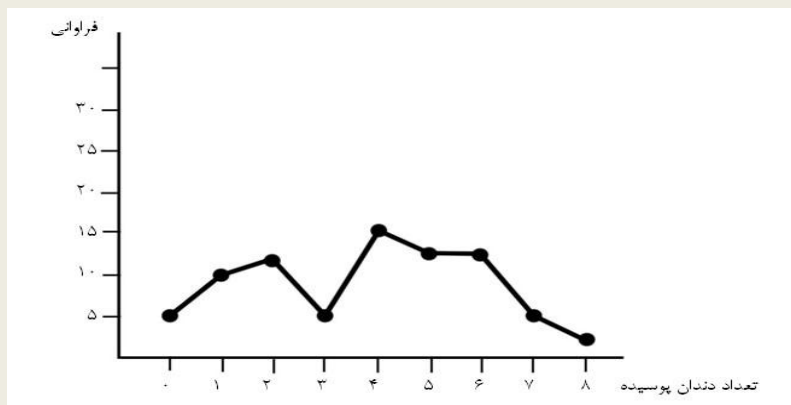
این نمودارها نیز برای نمایش داده های کیفی به کار می روند. به عبارت دیگر برای نمایش داده های کیفی بر حسب سلیقه از هر دو نوع نمودارهای دایره ای و میله ای می توان استفاده نمود. نمودار میله ای دو بعدی است و روی محور عرضی آن مقادیر فراوانی و روی محور طولی مقادیر صفت قرار می گیرند. به عنوان مثال اگر بخواهیم فراوانی عفونت های بیمارستانی را در یک بیمارستان بر حسب محل عفونت نشان دهیم در اینجا متغیر ما محل عفونت است و اگر داده های آن را بر حسب مایع مغزی - نخاعی، چشم، زخم جراحی، ... نامگذاری کنیم داده های ما کیفی (از نوع اسمی) هستند. در این حالت برای نمایش تصویری داده ها می توان از یک نمودار دو بعدی کمک گرفت و سپس فراوانی هر گروه از داده ها را به صورت یک میله رسم نمود.

نمودارها برای نمایش داده های کمی گسسته نیز به کار می روند.

۳- نمودار چندگوش

به عنوان مثال اگر بخواهیم فراوانی تعداد دندان پوسیده در دانش آموزان یک مدرسه را نشان دهیم در اینجا متغیر ما تعداد دندان پوسیده است می دانیم که داده های این متغیر، از نوع کمی گسسته هستند. به عبارت دیگر یک فرد یا دندان پوسیده ندارد یا یک دندان پوسیده، دو دندان پوسیده و ... یا حداکثر سی و دو دندان پوسیده دارد اما در فاصله دو واحد این داده ها هیچ مقدار دیگری قابل انتخاب نیست یعنی فردی را نداریم که مثلاً ۲/۳ دندان پوسیده داشته باشد.

در این حالت برای نمایش تصویری داده ها می توان از یک نمودار دو بعدی کمک گرفت که روی محور عرضی آن مقادیر فراوانی و روی محور طولی مقادیر صفت قرار می گیرند. سپس فراوانی متناظر هر گروه از داده ها را به صورت یک نقطه رسم نموده و با استفاده از خطوط نقاط را به یکدیگر وصل می کنیم در چنین وضعیتی نمودار می تواند به صورت زیر باشد.



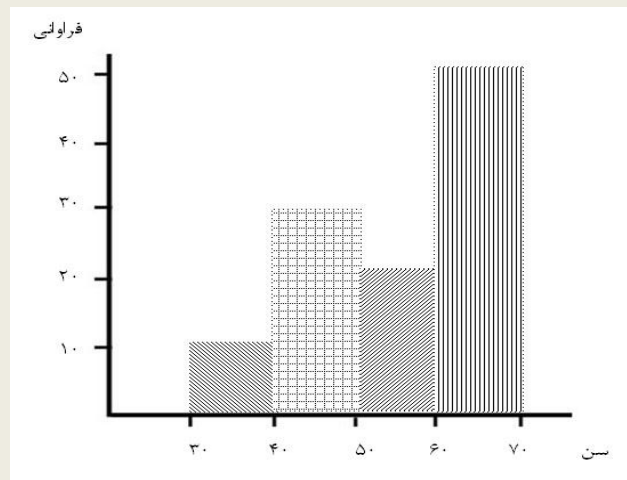
نمودار ۲- فراوانی دانش آموزان یک مدرسه بر حسب تعداد دندان پوسیده

۴- نمودار هیستوگرام (Histogram)

این نمودارها برای نمایش داده های کمی پیوسته به کار می روند به عنوان مثال اگر بخواهیم فراوانی سنی بیماران مبتلا به سرطان کولون در یک کلینیک شیمی درمانی را نشان دهیم در اینجا متغیر ما سن است می دانیم که داده های این متغیر از نوع کمی پیوسته هستند. به عبارت دیگر در فاصله ۲ واحد این داده ها مقادیر متنابهی از داده های سن قابل انتخاب هستند.

در این حالت برای نمایش تصویری داده ها مجدداً از یک نمودار دوبعدی کمک می گیریم که روی محور عرضی آن مقادیر فراوانی و روی محور طولی آن مقادیر صفت قرار می گیرند.

سپس فراوانی متناظر هر گروه سنی را به صورت یک جعبه (box) رسم می نماییم. در چنین وضعیتی نمودار می تواند به صورت زیر باشد.



نمودار ۳- فراوانی بیماران مبتلا به سرطان کولون در یک کلینیک شیمی درمانی بر حسب سن

به طور کلی نمودار داده ها را می توان به صورت زیر خلاصه کرد:

نمودارهای کیفی: دایره ای - میله ای

نمودارهای کمی: چند گوش - هیستوگرام

به جدول زیر توجه کنید:

سن	فراوانی
30-39	10
40-49	30
50-59	20
60-69	50

جدول یک - فراوانی بیماران مبتلا به سرطان کولون در یک کلینیک شیمی درمانی بر حسب سن

در اینجا چند اصطلاح را تعریف می کنیم:

۱- **فراوانی مطلق (Absolute frequency):** فراوانی مطلق همان فراوانی خام می باشد.

۲- **فراوانی نسبی (Relative frequency):** فراوانی نسبی یعنی فراوانی آن گروه از داده ها چه نسبتی از کل داده ها را شامل می شود. (معمولاً درصد از کل)

۳- **فراوانی تجمعی (Cumulative frequency):** فراوانی تجمعی یعنی فراوانی مطلق آن گروه از داده ها به علاوه فراوانی گروه های ماقبل

۴- **درصد تجمعی (Cumulative percent):** درصد تجمعی یعنی فراوانی تجمعی آن گروه از داده ها چه نسبتی از کل داده ها را شامل می شود (معمولاً درصد از کل)

بدین ترتیب برای نمایش کامل تر داده های نمودار شماره چهار جدول زیر را خواهیم داشت:

سن	فراوانی مطلق	فراوانی نسبی	فراوانی تجمعی	درصد تجمعی
30-39	10	9/09	10	9/09
40-49	30	27/27	40	36/36
50-59	20	18/18	60	54/54
60-69	50	45/45	110	100

جدول دو - فراوانی بیماران مبتلا به سرطان کولون در یک کلینیک شیمی درمانی بر حسب سن

به عنوان مثال در گروه بیماران ۴۹-۴۰ ساله فراوانی مطلق داده ها ۳۰ است زیرا تعداد بیماران ما ۳۰ نفر است. فراوانی نسبی برابر ۲۷/۲۷ است زیرا عدد ۳۰، ۲۷/۲۷٪ از کل داده ها (۱۱۰ بیمار) را شامل می شود. فراوانی تجمعی برابر ۴۰ است زیرا ۳۰ نفر در این گروه و ۱۰ نفر در گروه قبلی قرار دارند و درصد تجمعی برابر ۳۶/۳۶ است زیرا عدد ۴۰، ۳۶/۳۶٪ از کل داده ها (۱۱۰ بیمار) را شامل می شود.

۵- **صدک (Percentile):** عددی است که درصد توزیعی برابر با یا کمتر از آن عدد را نشان می دهد

۶- **چارک (Quartile):** چارک به هر یک از سه مقداری که یک مجموعه از داده های مرتب را به چهار بخش مساوی تقسیم می کند گفته می شود. به اینصورت هر کدام از آن بخش ها یک چهارم از نمونه یا جمعیت را به نمایش می گذارد.

Q1 چارک اول: نقطه ای که ۲۵ درصد نمره ها را از پایین جدا می کند. به آن نقطه ۲۵ درصد هم گفته می شود.

Q۲ چارک دوم: (میانۀ)، توزیع را به دو قسمت مساوی تقسیم می کند.

Q۳ چارک سوم: نقطه ای که سه چهارم نمره ها در زیر آن و بقیه در بالای آن واقع شده اند. به آن نقطه ۷۵ درصد هم گفته می شود.

دامنه تغییرات بین چارک ها: فاصله بین چارک اول و سوم

با استفاده از درصد تجمعی می توان صدک ها و چارک های مختلف را معین نمود.

به عنوان مثال قد دویست شرکت کننده مرد در مسابقه دو و میدانی به صورت زیر می باشد:

قد	فراوانی مطلق	فراوانی نسبی	فراوانی تجمعی	درصد تجمعی
<50	4	2	4	2
150-154	6	3	10	5
155-159	20	10	30	15
160-164	30	15	60	30
165-169	40	20	100	50
170-174	40	20	140	70
175-179	30	15	170	85
180-184	20	10	190	95
185-189	6	3	196	98
>190	4	2	200	100

جدول سه - توزیع فراوانی قد دویست شرکت کننده مرد در مسابقه دو و میدانی

در مثال فوق قد ۱۶۹ روی صدک پنجاهم و چارک دوم قرار دارد

برای توصیف هر چه بهتر متغیرهایی که داده های آنها کمی هستند، از دو گروه از شاخص های مرکزی و شاخص های پراکنده استفاده می کنیم. اگر بخواهیم نحوه تمرکز (مرکزیت) داده ها را توصیف نمائیم از شاخص های دسته اول و اگر چنانچه بخواهیم نحوه پراکندگی داده ها را توصیف نمائیم از شاخص های دسته دوم استفاده می کنیم. انواع شاخص های مرکزی به شرح زیر می باشند:

۱- میانگین (Mean): میانگین که آن را با $\bar{X}$ نمایش می دهند همان متوسط یا معدل داده ها می باشد و برای محاسبه آن خواهیم داشت:

$$R = X_{\max} - X_{\min} = 101 - 1 = 100$$

به عبارت دیگر باید تک تک داده ها را با هم جمع نمائیم و مجموع را بر تعداد داده ها تقسیم کنیم.

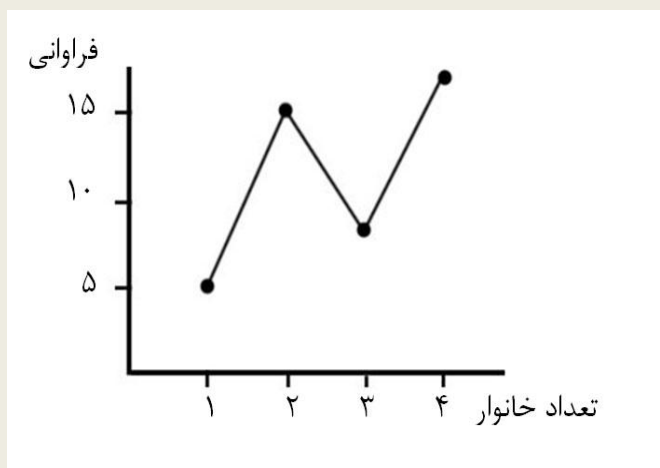
$\bar{X}$ میانگین جامعه است و در جایی که بخواهیم میانگین را برای یک نمونه محاسبه کنیم به جای آن از $\bar{X}$ استفاده می کنیم. به عنوان مثال میانگین اعداد ۴ و ۱۰۱ و ۲۳ و ۱۲ و ۱ برابر خواهد بود با:

$$\bar{X} = \frac{\sum X_i}{N} = \frac{1 + 12 + 23 + 101 + 4}{5} = 28 / 2$$

نکته: میانگین داده های گروه بندی شده از فرمول زیر محاسبه می شود:

$$\mu = \frac{\sum X_i N_i}{N}$$

به عنوان مثال اگر نمودار توزیع فراوانی بعد خانوار (تعداد خانوار) در یک محله به صورت زیر باشد:



میانگین بعد خانوار در این محله برابر خواهد بود با:

$$\bar{X} = \frac{\sum X_i N_i}{N} = \frac{1 \times 5 + 2 \times 15 + 3 \times 10 + 4 \times 20}{5 + 10 + 15 + 20} = 2 / 9$$

۲- میان (Median):

میان آن داده ای است که نیمی از آن بزرگ تر و نیمی دیگر از آن کوچک تر هستند. به عبارت دیگر تمامی داده ها را به دو قسمت مساوی تقسیم می کند.

برای محاسبه میان اعداد، ابتدا تمامی داده ها را به صورت صعودی مرتب می کنیم. سپس محل میان را با استفاده از فرمول

$$\frac{N+1}{2}$$

به دست می آوریم. به اندازه محل میان از کوچک ترین داده به سمت بزرگ ترین داده ها حرکت می کنیم به داده ای که می رسیم آن داده، میان ما خواهد بود.

به عنوان مثال اگر اعداد ۱ و ۴ و ۱۰۱ و ۲۳ و ۱۲ و ۱ را داشته باشیم و میان این اعداد را بخواهیم به دست بیاوریم ابتدا باید آنها را به صورت صعودی مرتب کنیم در چنین حالتی رشته اعداد ما بدین صورت خواهند بود:

۱ و ۴ و ۱۲ و ۲۳ و ۱۰۱

$$\text{محل میان} = \frac{N+1}{2} = \frac{5+1}{2} = 3$$

سپس محل میان بدین صورت به دست خواهد آمد:

پس میان ما سومین عدد در میان این رشته مرتب شده اعداد خواهد بود لذا در اینجا میان ۱۲ است.

اگر تعداد داده های ما (N) فرد باشد مشکلی نداریم اما اگر زوج باشد محل میان ما اعشاری خواهد بود در چنین وضعیتی باید دو عدد وسط رشته اعداد مرتب شده را با هم جمع و تقسیم بر دو کنیم. حاصل آن میان ما خواهد بود.

به عنوان مثال اگر اعداد ۱، ۴، ۱۰۱، ۲۳، ۱۲، ۱ را داشته باشیم و میان این اعداد را بخواهیم ابتدا آنها را به صورت صعودی مرتب کنیم در چنین حالتی رشته اعداد ما بدین صورت هستند:

۱، ۴، ۱۲، ۲۳، ۵۱، ۱۰۱

سپس محل میان خواهد بود:

$$\text{محل میان} = \frac{N+1}{2} = \frac{6+1}{2} = 3.5$$

$$\frac{23 + 51}{2} = 37$$

در اینجا دو عدد وسط رشته اعداد ما ۲۳ و ۵۱ هستند لذا میان ما برابر است با

۳- نما (Mode)

نما داده ای است که بیشترین فراوانی را داراست پس در جایی که تعدادی عدد داریم باید ببینیم کدام عدد بیش از سایرین

موجود است به عنوان مثال اگر اعداد ۴، ۳، ۷، ۲، ۳، ۳، ۶، ۲ را داشته باشیم نما برابر با ۳ خواهد بود زیرا در میان این هشت عدد، عدد ۳ است که بیشتر از سایر اعداد تکرار شده است.

انواع شاخص های پراکندگی به شرح زیر می باشد:

۱- دامنه تغییرات (Range)

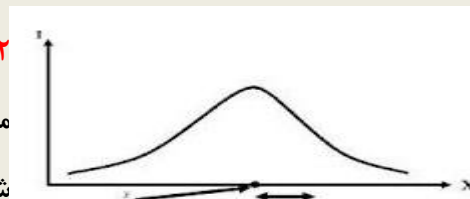
دامنه تغییرات یا دامنه که آن را با **R** نشان می دهند به صورت زیر محاسبه می شود:

$$R = X_{\max} - X_{\min}$$

به عبارت دیگر برای محاسبه دامنه تغییرات باید بزرگ ترین داده را انتخاب و منهای کوچک ترین داده نمائیم.

۲- میانگین انحرافات (Mean deviation)

میانگین انحرافات که آن را با **M.D.** نشان می دهند به صورت زیر محاسبه می شود.



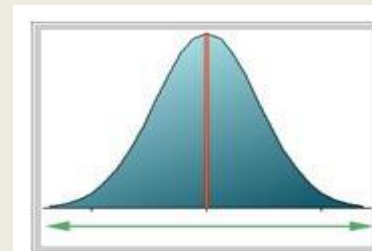
$$M.D. = \frac{\sum |X_i - \mu|}{N}$$

میانگین انحرافات یعنی میانگین انحراف هر یک از داده ها نسبت به مقدار متوسط داده ها

پس میانگین انحرافات برای سه عدد ۳۰۰، ۲۰۰ و ۱۰۰ برابر خواهد بود با:

$$M.D. = \frac{\sum |X_i - \bar{X}|}{N} = \frac{|100 - 200| + |200 - 200| + |300 - 200|}{3} = \frac{200}{3} = 66.66$$

۲- واریانس (Variance)



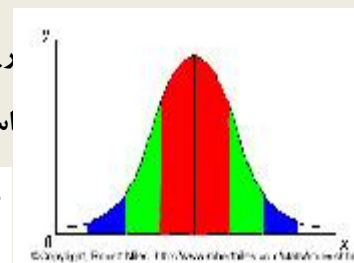
مقدار واریانس با میانگین گیری از مربع فاصله مقدار محتمل و یا مشاهده شده با مقدار مورد انتظار محاسبه می شود. در مقایسه با میانگین می توان گفت که میانگین مکان توزیع را نشان می دهد، در حالی که واریانس مقیاسی است که نشان می دهد که داده ها حول میانگین چگونه پخش شده اند. واریانس را با S^2 نشان می دهند به صورت زیر محاسبه می شود:

$$\sigma^2 = \frac{\sum (X_i - \mu)^2}{N}$$

۳- انحراف معیار (Standard deviation)

ریشه دوم واریانس که انحراف معیار نامیده می شود دارای واحدی یکسان با متغیر اولیه است. نحوه محاسبه آن به شرح ذیل است:

$$\sigma = S.D. = \sqrt{\frac{\sum (X_i - \mu)^2}{N}}$$



واریانس و انحراف معیار بهترین شاخص ها جهت بررسی پراکندگی داده های کمی هستند. لازم به یادآوری است در جایی که تعداد داده های ما ۳۰ و یا کمتر از ۳۰ باشند در مخرج کسر باید به جای N از $N-1$ استفاده کنیم و فرمول واریانس و انحراف معیار ما به صورت زیر خواهند بود:

$$S^2 = \frac{\sum (X_i - \bar{X})^2}{N - 1}$$

$$S = S.D. = \sqrt{\frac{\sum (X_i - \bar{X})^2}{N - 1}}$$

و در این حالت به جای استفاده از S^2 از S (برآورد واریانس) استفاده می کنیم. پس واریانس و انحراف معیار برای ۳ عدد ۳۰۰ و ۲۰۰ و ۱۰۰ برابر خواهد بود با:

$$S^2 = \frac{\sum (X_i - \bar{X})^2}{N - 1} = \frac{(100 - 200)^2 + (200 - 200)^2 + (300 - 200)^2}{3 - 1} = \frac{10/000 + 10/000}{2} = 10/000$$

و $S = 100$

نکته: اگر تمامی داده ها را با مقدار ثابتی جمع و یا از مقدار ثابتی کم کنیم میانگین داده های جدید به همان مقدار افزایش و یا کاهش می یابند ولی انحراف معیار و واریانس تغییر نمی کند.

به عنوان مثال اگر میانگین نمرات شرکت کنندگان در آزمون دستیاری امسال ۴۰۰ و انحراف معیار نمرات آنها ۱۰ و واریانس آن ۱۰۰ باشد و سپس به تمام شرکت کنندگان ۱۰ نمره اضافه نمایند میانگین نمرات جدید ۴۱۰ خواهد بود اما انحراف معیار و واریانس همان مقادیر قبلی را خواهند داشت.

نکته: اگر تمامی داده ها را در مقدار ثابتی ضرب و یا بر مقدار ثابتی تقسیم نمائیم میانگین و انحراف معیار داده های جدید به همان مقدار بزرگ و یا کوچک می شود اما واریانس با مجذور عدد ثابت تغییر می کند. به عنوان مثال اگر متوسط طول قد دانش آموزان یک کلاس ۱/۷۰ متر و انحراف معیار قد آنها از متر باشد و بخواهیم واحد اندازه گیری طول را ۰/۱ متر به سانتی متر تغییر دهیم میانگین داده های جدید ۱۷۰ (۱۰۰ × ۱/۷۰)، انحراف معیار آنها ۱۰ (۱۰۰ × ۰/۱) و واریانس آنها ۱۰۰ (۱۰۰ × ۰/۱) = ۱۰۰۰۰ خواهد بود.

۴- ضریب تغییرات (Coefficient of variation)

شاخص های قبلی جهت بررسی پراکندگی یک صفت بودند اما ضریب تغییرات شاخصی است که برای مقایسه پراکندگی دو یا چند صفت به کار می رود ضریب تغییرات را با **C.V** نمایش می دهند و به صورت زیر محاسبه می شود:

$$C.V. = \frac{\sigma}{\mu} \times 100$$

به عنوان مثال اگر متوسط تعداد ضربان قلب تعدادی از افراد ۸۰ و انحراف معیار تعداد ضربان قلب آنها ۴ باشد در این حالت ضریب تغییرات تعداد ضربان قلب این افراد برابر خواهد بود با:

$$C.V. = \frac{\sigma}{\mu} \times 100 = \frac{4}{80} \times 100 = 5$$

حال اگر متوسط درجه حرارت بدن همین افراد ۳۷ و انحراف معیار آنها ۰/۷۴ باشد در این حالت ضریب تغییرات درجه حرارت این افراد برابر است با:

$$C.V. = \frac{0.74}{37} \times 100 = 2$$

در این حالت می توانیم نتیجه بگیریم که داده های تعداد ضربان قلب ۲/۵ برابر $(\frac{5}{2} = 2.5)$ پراکنده تر از داده های درجه حرارت بدن هستند.

جامعه و نمونه

در مدل استنباط آماری، فرض بر این است که می خواهیم در مورد یک مجموعه خیلی بزرگ (شاید نامحدود)، اطلاعات کسب کنیم (مثلا نمره پیشرفت تحصیلی درس ریاضی دانش آموزان کلاس پنجم دبستان در سراسر کشور). به این مجموعه، جامعه گفته می شود. گاه حجم جامعه آن قدر بزرگ است که نمی توان تمام آن را مطالعه نمود، لذا از کل مجموعه، یک زیرمجموعه به عنوان نمونه کل مشاهدات ممکن برای مطالعه انتخاب می شود. به این زیرمجموعه که شامل تعداد محدودی از اعضای جامعه است "نمونه" گفته می شود. اما جهت استنباط خصوصیات جامعه از روی خصوصیات نمونه، مدل آماری ایجاب می کند که اعضای گروه نمونه به صورت تصادفی انتخاب شوند. نمونه تصادفی به نمونه ای گفته می شود که همه اعضای جامعه به یک اندازه شانس شرکت و انتخاب شدن در آن را داشته باشند. همچنین انتخاب هر فرد مستقل از افراد دیگر صورت گیرد.

مفاهیم و روشهای نمونه گیری

از آنجا که در سرشماری تمام واحدهای جامعه باید شمارش شود این کار پرهزینه و وقت گیر خواهد بود. برای صرفه جویی در وقت و هزینه مجبوریم روش دیگری را بکار ببریم. در اینجا است که اهمیت روش نمونه گیری آشکار می شود. در نمونه گیری معمولاً نمونه کوچکی از جامعه را بررسی می کنیم و آن را برای کل جامعه تعمیم می دهیم. هر وقت تصمیم بگیریم که بوسیله

بررسیهای نمونه‌ای اطلاعاتی را تهیه کنیم، فوراً با دو مطلب مواجه می‌شویم: تعریف دقیق جامعه‌ای که علاقمند به مطالعه آن هستیم، و گزینش مشخصه یا مشخصه‌هایی که باید ثبت شوند. مفاهیم کلی برای نمونه‌گیری از قبیل جامعه، نمونه، سرشماری و... را برای ارائه دید کلی از روش نمونه‌گیری و مزایای آن در انجام بررسیهای آماری ضروری است معرفی شوند.

انواع بررسیهای نمونه‌ای

در یک تحقیق همیشه میسر نیست که کل آزمودنی‌های یک جامعه را مورد بررسی قرار دهیم چرا که ممکن است از نظر زمان، منابع مالی، منابع انسانی و... محدودیت‌هایی موجود باشد و امکان بررسی همه جامعه موجود نباشد؛ لذا باید جزئی از آزمودنی‌ها را به گونه‌ای که نماینده معقولی برای کل آزمودنی‌های آن جامعه باشند، انتخاب نموده و نتایج بررسی روی آنها را به کل جامعه تعمیم داد. البته با توجه به وجود خطا در این عمل، لازم است میزان خطا را نیز محاسبه و اعلام نماییم.

ویژگی‌های جامعه از قبیل μ (میانگین جامعه)، σ (انحراف معیار جامعه) و σ^2 (واریانس جامعه) به عنوان پارامترهای جامعه مطرح هستند. به عبارت دیگر هر خصوصیت عددی یک جامعه را پارامتر آن جامعه گویند. و آماره ویژگی و خصوصیت نمونه است. آماره‌های تمایل به مرکز، پراکندگی و سایر آماره‌های نمونه مورد بررسی، به عنوان تقریبی از پارامترهای تمایل به مرکز و پراکندگی و سایر پارامترهای جامعه محسوب می‌شود. بدین لحاظ یافته‌های نمونه به جامعه تعمیم داده می‌شود. به عبارت دیگر، آماره‌های نمونه به عنوان برآوردهایی از پارامترهای جامعه مورد استفاده قرار می‌گیرد. شکل زیر رابطه بین جامعه و نمونه را نشان می‌دهد.

مزایای نمونه‌گیری

دلایل استفاده از نمونه‌گیری به جای سرشماری جامعه نسبتاً واضح و روشن است. در بررسی تحقیقاتی که درگیر صدها یا حتی هزاران عضو هستند، جمع‌آوری اطلاعات یا آزمودن هر عضو می‌تواند به طور علمی غیرممکن باشد. حتی اگر موارد ذکر شده ممکن باشد، زمان، هزینه و سایر منابع انسانی مانع عمده‌ای بر سر راه پژوهش هستند. احتمالاً بررسی یک نمونه، به جای کل جامعه گاهی اوقات منجر به نتایجی می‌شود که از روانی بالاتری برخوردار است زیرا اساساً در این حالت خستگی کمتری وجود خواهد داشت و بنابراین در جمع‌آوری اطلاعات خطاهای کمتری صورت می‌گیرد، خصوصاً زمانی که اعضای جامعه زیادند. در برخی موارد، استفاده از کل جامعه برای آگاهی یا آزمودن چیزی ناممکن است. برای مثال در آزمون عمر یک دسته از لامپ‌ها، اگر مجبور باشیم هر لامپ تولیدی را بسوزانیم، چیزی برای فروش نخواهد ماند. به طور خلاصه می‌توان گفت نمونه‌گیری دارای مزیت‌های زیر می‌باشد:

۱- هزینه کمتر

۲- زمان کمتر، سرعت بیشتر

۳- دقت بیشتر در اجرا

۴- به علت محفوظ ماندن واحدهای جامعه یا تخریب نشدن آنها

طرح‌های نمونه‌گیری

بسته به نوع جامعه آماری می‌توان از روش‌های مختلفی برای نمونه‌گیری استفاده نمود. به طور کلی دو نوع طرح نمونه‌گیری وجود دارد: نمونه‌گیری تصادفی و غیر تصادفی. در **نمونه‌گیری تصادفی**، اعضای جامعه به عنوان آزمودنی‌های نمونه منتخب، از شانس و احتمال یکسانی برخوردارند. در **نمونه‌گیری غیر تصادفی**، اعضای جامعه به عنوان آزمودنی‌های نمونه منتخب، از شانس و احتمال یکسانی برای انتخاب شدن برخوردار نیستند. **طرح‌های نمونه‌گیری تصادفی وقتی استفاده می‌شوند که نماینده (معرف) بودن نمونه به خاطر اهداف تعمیم‌پذیری حائز اهمیت است. وقتی که زمان یا سایر عوامل، نسبت به تعمیم‌پذیری از اهمیت بسیار زیادتری برخوردار می‌شوند، عموماً نمونه‌گیری غیر تصادفی استفاده می‌شود.** هر کدام از این دو طرح اصلی، استراتژی‌های نمونه‌گیری مختلفی دارند. بسته به میزان تعمیم‌پذیری مورد نظر، فراهم بودن زمان و سایر منابع و هدفی که پژوهش به دنبال آن است، انواع مختلفی از طرح‌های نمونه‌گیری تصادفی و غیر تصادفی انتخاب خواهند شد. در ادامه این جلسه انواع مختلف این طرح‌ها معرفی خواهند شد:

۱- نمونه‌گیری تصادفی

نمونه‌گیری تصادفی می‌تواند ماهیتاً نامحدود (یا نمونه‌گیری تصادفی ساده) یا محدود (نمونه‌گیری تصادفی پیچیده) باشند.

۱-۱- نمونه‌گیری تصادفی ساده یا نامحدود

در طرح نمونه‌گیری تصادفی نامحدود، که عموماً به نمونه‌گیری تصادفی ساده مشهور است، هر عضوی در جامعه برای انتخاب شدن به عنوان یک آزمودنی از شانس مساوی و معین برخوردار است به طوری که انتخاب هیچ عضوی در انتخاب عضو بعدی مؤثر نباشد. به عبارت دیگر هر آزمودنی به طور مستقل از سایر آزمودنی‌ها، شانس انتخاب شدن برابر دارد. احتمال انتخاب شدن هر عضو عبارت است که:

برای انتخاب نمونه به شیوه تصادفی ساده چندین روش وجود دارد: قرعه‌کشی نام اعضا، انتخاب رندومی آزمودنی‌ها از روی شماره آنها در لیست، استفاده از جدول اعداد تصادفی، استفاده از ترتیب مضرب‌ها.

از جمله مزیت‌های این طرح این است که از حداقل پراکندگی برخوردار است و دارای بیشترین قدرت و تعمیم‌پذیری است. این فرآیند معایبی نیز دارد، پُرحمت و هزینه‌بر است، و علاوه بر این ممکن است همیشه یک فهرست تازه و به‌روز از اعضای جامعه آماری در دست نباشد.

۱-۲- نمونه‌گیری تصادفی پیچیده یا محدود

به جای طرح نمونه‌گیری تصادفی ساده، چندین طرح نمونه‌گیری تصادفی پیچیده می‌تواند مورد استفاده قرار گیرد. این شیوه نمونه‌گیری تصادفی نوعی طرح جایگزین پایا و بعضاً کارآمدتری نسبت به طرح نمونه‌گیری نامحدود که مورد بحث قرار گرفت، ارائه می‌دهند. چون اطلاعاتی که می‌تواند برای یک حجم نمونه‌گیری معین با استفاده از برخی شیوه‌های نمونه‌گیری تصادفی پیچیده به دست آید نسبت به استفاده و به کارگیری طرح نمونه‌گیری تصادفی ساده بیشتر است، لذا کارایی آن بیشتر می‌باشد. اکنون پنج طرح از طرح‌های نمونه‌گیری تصادفی پیچیده که بیشتر متداول اند را مورد بحث قرار می‌دهیم:

۱-۲-۱- نمونه‌گیری سیستماتیک

هر گاه تعداد آزمودنی‌های یک جامعه آماری بزرگ باشد، انجام نمونه‌گیری تصادفی ساده ممکن است با مشکل مواجه شود. در چنین مواقعی برای انتخاب نمونه از روش نمونه‌گیری تصادفی سیستماتیک یا نظام‌دار استفاده می‌شود.

فرض کنید حجم جامعه N باشد. واحدهای جامعه از ۱ تا N شماره گذاری می‌کنیم. انتخاب نمونه ای منظم به حجم n از جامعه ای به حجم N به ترتیب زیر است: ابتدا جامعه را به n قسمت تقسیم می‌کنیم. $(\frac{n}{N} = K)$ که K فاصله هر قسمت است. سپس از واحدهای شماره ۱ تا K یک واحد به طور تصادفی (مثلاً واحد شماره m) را انتخاب کرده و آن را واحد اول نمونه قرار می‌دهیم. سپس به شماره m مضارب صحیح K را اضافه می‌کنیم تا شماره واحدهای بعدی نمونه مشخص شوند. این عمل را تا زمانی ادامه می‌دهیم که شماره ای بزرگتر از N به دست می‌آید. K را فاصله نمونه‌گیری می‌نامند. نمونه حاصل به صورت زیر است:

$$y_m, y_{m+k}, y_{m+2k}, \dots, y_{m+(n-1)k}$$

مسئله‌ای که محقق باید در طرح نمونه‌گیری تصادفی از آن آگاه باشد، احتمال ایجاد نوعی تورش (پراکندگی) در نمونه است. برای مثال فرض کنید که پنجمین خانه اتفاقاً در نبش ساختمان باشد و واحدهایی که بعد از شمارش هفت تا هفت تا انتخاب می‌شوند نیز در نبش ساختمان باشند. اگر کانون پژوهش در این مطالعه کنترل آلودگی صوتی برای ساکنانی باشد که عایق-بندی مناسب را به کار می‌برند، در این صورت ساکنان خانه‌های نبشی نسبت به خانه‌هایی که بین دو ساختمان قرار دارند احتمالاً آن قدر در معرض سروصدا نخواهند بود. بنابراین وقتی از ساکنان واحدهای نبشی آپارتمان، اطلاعاتی در مورد سطح

می‌آوریم، محقق ممکن است اطلاعاتی را که جمع‌آوری کرده دچار نوعی تورش باشد. احتمال استخراج یافته‌های ناصحیح از چنین اطلاعاتی بسیار بالاست. بنابراین، در نمونه‌گیری سیستماتیک زمینه برای ایجاد تورش وجود دارد و محقق باید طرح‌های خود را به دقت مورد توجه قرار دهد و مطمئن شود که طرح نمونه‌گیری سیستماتیک برای بررسی مناسب است و این کار را قبل از اخذ تصمیم در مورد استفاده از آن انجام دهد.

۱-۲-۲- نمونه‌گیری طبقه‌ای شده

اگرچه نمونه‌گیری به برآورد پارامترهای جامعه کمک می‌کند، ولی ممکن است گروه‌های فرعی معینی از اعضاء در جامعه وجود داشته باشند که انتظار رود در مورد یک مطالعه مورد توجه محقق، پارامترهای مختلفی داشته باشند. به عبارت دیگر هر گاه آزمودنی‌های جامعه از نظر برخی خصوصیات همگون نباشند و بتوان آنها را در طبقات مختلف گروه‌بندی کرد، نمونه‌گیری باید طوری انجام شود که به نسبت افراد هر طبقه در جامعه، همان نسبت در نمونه نیز انتخاب شود.

برای مثال، برای پژوهشی که مدیر بخش منابع انسانی علاقه‌مند به ارزیابی میزان نیازهای آموزشی کارکنان است، کل سازمان، جامعه پژوهش را تشکیل خواهد داد. اما میزان کیفیت و فشردگی آموزش مورد نیاز مدیران ارشد، مدیران سطح پایین‌تر، سرپرستان، تحلیل‌گران رایانه، کارکنان دفتری و غیره برای هر گروه، متفاوت خواهد بود. آگاهی از انواع تفاوت‌هایی که وجود دارد در ایجاد برنامه‌های آموزشی سودمند و مفید در سازمان کمک خواهد کرد. به این دلیل، اطلاعات باید به شیوه‌ای جمع‌آوری شود که به ارزیابی نیازهای هر سطح از گروه‌های فرعی در جامعه بتواند کمک کند. در این صورت واحد تجزیه و تحلیل می‌تواند سطح هر گروه باشد.

فرآیند نمونه‌گیری تصادفی طبقه‌ای همان‌طور که از نامش پیداست، متضمن فرآیندی از تشکیل طبقات و تجزیه آنهاست که به وسیله انتخاب تصادفی آزمودنی از هر طبقه دنبال می‌شود. در این نمونه‌گیری ابتدا جامعه به گروه‌های ناسازگاری که در بطن پژوهش مناسب، مرتبط و بامعنا هستند، تقسیم می‌شوند. بدون پیش گرفتن شیوه نمونه‌گیری تصادفی طبقه‌ای، یافتن تفاوت‌ها در پارامترهای گروه‌های فرعی درون یک جامعه نمی‌تواند ممکن باشد.

طبقه‌بندی نوعی طرح نمونه‌گیری کارآمد در پژوهش است؛ یعنی، طبقه‌بندی، اطلاعات بیشتری در مورد حجم نمونه معینی فراهم می‌کند. طبقه‌بندی باید خطوطی را دنبال کند که با پرسش‌های پژوهش متناسب باشند. برای مثال اگر ترجیحات مشتری را برای یک محصول مورد بررسی قرار می‌دهیم، طبقه‌بندی جامعه می‌تواند بر اساس حوزه‌های جغرافیایی، بخش‌های بازار، سن مصرف‌کننده، جنسیت مصرف‌کننده، یا ترکیب‌های مختلفی از این موارد صورت گیرد. روش طبقه‌بندی، میزان همگنی را در هر طبقه و ناهمگنی را بین طبقات نشان می‌دهد. به عبارت دیگر، اختلاف‌ها و تفاوت‌های بین گروه‌ها بیشتر از تفاوت‌های درون گروه خواهد بود.

وقتی جامعه به شیوه‌های صحیحی طبقه‌بندی گردید،

تصمیم‌گیری در مورد نمونه‌گیری که دارای تناسب خاصی نیست وقتی گرفته می‌شود که برخی از طبقات بسیار کوچک و برخی دیگر بسیار بزرگند، یا وقتی که نسبت به تغییرپذیری ویژگی‌ها درون یک طبقه خاصی که تعداد بیشتری در نمونه جای خواهند گرفت، نوعی شک و تردید وجود داشته باشد و یا محدودیت‌های زمان و هزینه وجود داشته باشد.

۱-۲-۳- نمونه‌گیری تصادفی خوشه‌ای

در نمونه‌گیری خوشه‌ای، گروه‌ها یا خوشه‌هایی از اعضا، از گروه‌هایی که بین اعضای آنها نوعی عدم‌تجانس وجود دارد،

تقسیم جمعیت به خوشه‌ها (افراد یا اشیا مجاور یکدیگر از نظر مکانی) یا گروه‌های مجاور از نظر مکانی و سپسTM انتخاب چند خوشه یا گروه بطور تصادفی از بین خوشه‌ها می‌باشد. (نحوه انتخاب خوشه‌ها می‌تواند به صورت تصادفی ساده یا منظم باشد) تقسیم بندی خوشه‌ها بر اساس مناطق جغرافیایی یا فواصل مکانی متداول است. نمونه‌گیری خوشه‌ای می‌تواند یک یا چند مرحله‌ای باشد.

برای مطالعه انتخاب می‌شوند. این نمونه‌گیری به دو شیوه انجام می‌شود:

در مطالعه ارزیابی شیوع کم خونی در دانش آموزان ابتدایی شهر تهران یک شیوه نمونه‌گیری می‌تواند به صورت زیر باشد: ابتدا از بین ۲۲ منطقه شهر تعدادی از مناطق به صورت تصادفی ساده انتخاب می‌شوند (۵ منطقه)، لیستی از مدارس ابتدایی مناطق پنج گانه انتخاب شده تهیه می‌کنیم و سپس ۱۰ دبستان در هر منطقه به صورت تصادفی ساده انتخاب می‌گردد و در نهایت در هر دبستان انتخاب شده ۵۰ دانش آموز به صورت تصادفی منظم براساس لیست حروف الفبایی انتخاب می‌گردند که با کمی دقت در می‌یابیم این شیوه نمونه‌گیری، نمونه‌گیری خوشه‌ای سه مرحله‌ای می‌باشد ▪

یک تفاوت نمونه‌گیری خوشه‌ای با سایر روش‌های نمونه‌گیری تصادفی این است که در نمونه‌گیری خوشه‌ای نیاز به وجود چارچوب نمونه‌گیری در کلیه اعضای جامعه نمی‌باشد چرا که در خوشه‌هایی که انتخاب نمی‌شوند نیازی به وجود چارچوب نمونه‌گیری نبوده و فقط چارچوب نمونه‌گیری خوشه‌ای منتخب را نیاز داریم ▪

الف) نمونه‌گیری خوشه‌ای تک مرحله‌ای

جامعه به خوشه‌ها یا گروه‌های مناسب تقسیم می‌شود و آزمودنی‌های نمونه، به صورت تصادفی از تعدادی خوشه‌ها یا گروه‌های مورد نیاز، انتخاب می‌شوند و همه عناصر از خوشه‌ای که به صورت تصادفی انتخاب شده است مورد بررسی قرار می‌گیرند.

ب) نمونه‌گیری خوشه‌ای چند مرحله‌ای

ابتدا از خوشه‌ها یا گروه‌های جامعه، نمونه‌ای اولیه به شیوه تصادفی استخراج می‌کنیم؛ سپس از بین این خوشه‌ها بار دیگر نمونه‌ای ثانویه تصادفی انتخاب می‌کنیم و اینکار را تا زمانی که مرحله‌نهایی این سطح برسیم انجام می‌دهیم. آنگاه در آخر نمونه‌ای خواهیم داشت که از هر واحدی، عضوی در آن وجود دارد.

یکی از انتقادهای وارده بر نمونه‌گیری خوشه‌ای این است که ممکن است انتخاب مضاعف اعضاء در چندین خوشه، اتفاق بی‌افتد. اگرچه این شیوه کم هزینه است اما از نوعی کارائی در دقت و اطمینان در نتایج برخوردار نیست.

۲- نمونه‌گیری غیرتصادفی

در طرح‌های نمونه‌گیری تصادفی، عناصر یا اعضا درون جامعه از هیچ‌گونه احتمال مساوی انتخاب شدن به عنوان آزمودنی-های نمونه برخوردار نیستند. این بدین معناست که یافته‌های حاصل از بررسی نمونه، نمی‌تواند با اطمینان، به جامعه تعمیم داده شود. در عین حال، محققان ممکن است در برخی موارد نیاز به تعمیم‌پذیری نداشته باشند و فقط قصد بدست آوردن برخی اطلاعات مقدماتی به شیوه‌ای سریع و کم‌هزینه داشت باشند. در چنین مواردی آنها می‌توانند به نمونه‌گیری غیرتصادفی متوسل شوند. برخی از طرح‌های نمونه‌گیری غیرتصادفی نسبت به برخی دیگر قابلیت اعتماد بیشتری دارند و رهنمودهای مهمی برای کسب اطلاعات مفید در مورد جامعه می‌توانند ارائه دهند.

اهمیت طرح نمونه‌گیری و اندازه نمونه

اگر طرح نمونه‌گیری مناسب، مورد استفاده قرار نگیرد، صرف یک حجم نمونه بزرگ اجازه نخواهد داد که یافته‌ها به جامعه تعمیم داده شود. همین‌طور، اگر حجم نمونه برای سطح اطمینان و دقت مورد انتظار کافی نباشد، هیچ طرح نمونه‌گیری پیشرفته نمی‌تواند برای محقق در تحقق اهداف بررسی مفید باشد. بنابراین، در اتخاذ تصمیمات، برای مثال راجع به نمونه-گیری باید هم طرح نمونه‌گیری و هم حجم نمونه مد نظر قرار گیرند. در عین حال اگر حجم نمونه بسیار زیاد باشد، می‌تواند مشکل‌زا باشد و زمینه مساعدی برای دچار شدن به خطای نوع دوم را ایجاد کند. به عبارت دیگر، یافته‌های پژوهش خود را می‌پذیریم در حالی که باید آنها را رد کرده باشیم. به عبارت دیگر، وقتی نمونه بسیار بزرگ است و حتی روابط در سطوح معنی‌دار ضعیف است یقین پیدا می‌کنیم که این روابط معنی‌دار در نمونه در مورد جامعه واقعاً صادق است، در حالی که واقعیت دلالت بر آن دارد که چنین چیزی ممکن نیست. بنابراین نه اندازه بزرگ و نه اندازه کوچک به پروژه‌های پژوهشی کمک چندانی نمی‌کنند.

به طور کلی قواعد زیر را برای تعیین حجم نمونه می‌توان به کار برد:

۱. برای بیشتر تحقیقات، حجم نمونه بزرگتر از ۳۰ و کمتر از ۵۰۰ مناسب است.
۲. در جاهایی که نمونه‌ها باید به نمونه‌های فرعی تقسیم شوند، حداقل حجم نمونه برای هر طبقه ۳۰ عضو ضروری است.
۳. در تحقیقات چندمتغیره، حجم نمونه باید چند برابر (ترجیحاً ۱۰ برابر) و به بزرگی تعداد متغیرها در پژوهش باشند.
۴. برای پژوهش‌های تجربی ساده با کنترل‌های آزمایشی شدید (زوج‌های جفت شده) پژوهش موفق با نمونه‌های دارای حجم ۱۰ تا ۲۰ نیز میسر است.

مراحل اصلی در یک بررسی نمونه‌ای

اهداف بررسی: همواره باید حکمی روشن و صریح درباره هدفهای بررسی در دست باشد. در غیر این صورت با افزایش حجم کار و جزئیات دیگر نمونه گیری، تصمیم‌هایی اتخاذ می‌شوند که با اصل اهداف هماهنگی ندارند.

جامعه مورد نمونه گیری: جامعه‌ای که نمونه از آن می‌گیریم، باید دقیقاً تعریف شود. جامعه‌ای که از آن نمونه می‌گیریم باید منطبق بر جامعه هدف باشد یعنی مجموعه تمام عناصر مورد مطالعه.

جمع آوری داده‌ها: لازم است تحقیق کنیم که تمام داده‌ها به اهداف بررسی مربوط‌اند و هیچ داده اساسی از قلم نیفتاده است.

درجه دقت مطلوب: نتایج یک بررسی نمونه‌ای همیشه با عدم حتمیت همراه است، زیرا اولاً نسبتی از جامعه مورد اندازه گیری قرار گرفته است و ثانیاً اندازه گیری‌ها همیشه با خطا همراه‌اند. میزان این عدم دقت را می‌توان با نمونه‌های بزرگتر و با استفاده از وسایل اندازه گیری دقیق‌تر تقلیل داد.

روش اندازه گیری: در جامعه، برای اندازه گیری واحدهای نمونه، انتخاب ابزار اندازه گیری و روش اندازه گیری واجد اهمیت است.

چارچوب: قبل از انتخاب نمونه جامعه را باید به بخش‌هایی تقسیم کرد. این بخش‌ها را واحدهای نمونه گیری یا فقط واحدها می‌نامند. فهرست کامل واحدهای نمونه گیری را چارچوب می‌نامند.

انتخاب نمونه: حال طرح‌های متعددی وجود دارند که می‌توان با آنها نمونه را انتخاب کرد. برای هر طرحی و با توجه به درجه دقت مورد نیاز در برآوردها باید حجم خاصی از نمونه را مشخص نمود.

پیش آزمون: تجربه نشان داده است که قبل از انجام نمونه گیری نهایی، امتحان کارایی پرسشنامه و یا روشهای مورد نظر با مقیاسی کوچک بسیار مفید است.

آموزش آمارگران: در بررسیهای جامع نمونه‌ای، اغلب با مسائل خاص حرفه‌ای مواجهیم. لذا آمارگران باید قبلاً درباره هدف نمونه گیری و روشهای نمونه گیری و جمع آوری داده‌ها و سایر خط‌مشی‌ها آموزش ببینند.

تلخیص و تحلیل داده‌ها: اولین مرحله، آماده کردن پرسشنامه‌های تکمیل شده برای انتقال داده‌ها به نرم افزار مربوطه است.

اطلاعات حاصل برای بررسیهای آتی: هر نمونه‌ای که از جامعه گرفته می‌شود بالقوه راهنمایی برای اصلاح نمونه گیریهای بعدی است.

چه روش نمونه گیری را باید بکار برد؟

تعیین طرحی از نمونه گیری که باید به کار برد و انتخاب کردن حجمهای نمونه‌ای، از موضوعهای کلیدی در طرح ریزی یک بررسی هستند. انتخاب یک روش نمونه گیری مناسب مبتنی بر عاملهایی از قبیل ساختار جامعه، نوع اطلاع مورد جستجو، و تسهیلات اداری و پرسنل موجود برای اجرای بررسی است. در رابطه با انتخاب روش نمونه گیری مناسب، حجم نمونه مورد نیاز با مشخص کردن یک درجه دقت مطلوب برای برآوردها تعیین می‌شود. آنگاه باید این موضوع را هم تحقیق کرد که آیا بودجه‌ای که به بررسی اختصاص داده شده است، امکان تهیه این حجم نمونه را می‌دهد.

آزمون فرض

فرض آماری، ادعایی در مورد یک یا چند جامعه مورد بررسی است که ممکن است درست یا نادرست باشد.

به عبارت دیگر فرض آماری، یک ادعا یا گزاره‌ای در مورد توزیع یک جمعیت یا پارامتر توزیع یک متغیر تصادفی است. فرضیه آماری، نقطه آغاز آزمون فرض است و اصولاً بدون داشتن فرضیه آماری امکان انجام یک آزمون دشوار است. فرضیه آماری به دو نوع فرض صفر (H_0) و فرض خلاف (H_A) بیان می‌شود. اغلب فرضیه بیانگر این مطلب است که یک ارتباط علیتی بین دو متغیر وجود دارد به شکلی که میزان یکی (متغیر مستقل یا **Independent**) تا حدودی تعیین کننده دیگری متغیر وابسته یا (**Dependent**) است. فرضیه‌ای که در آزمون‌های آماری مورد آزمون قرار می‌گیرد فرضیه صفر است که همیشه حاکی از عدم وجود تفاوت می‌باشد. اما فرض خلاف همان فرضیه پژوهشی است که می‌تواند جهت‌دار یا غیر جهت‌دار باشد. البته انتخاب فرضیه جهت‌دار دلخواه و تصادفی نیست، بلکه در صورتی فرضیه پژوهشی را می‌توان جهت‌دار تدوین کرد که تئوری یا تحقیقات قبلی شواهدی برای آن ارائه کنند.

انواع خطا در استنباط آماری

پس از انجام آزمون‌های آماری، محقق در مورد رد یا عدم رد فرضیه صفر تصمیم می‌گیرد. اگر نتایج آزمون به گونه‌ای باشد که نتوان آن را رد کرد، جایی برای اثبات یا تأیید فرضیه پژوهشی باقی نمی‌ماند، اما اگر فرضیه صفر رد شود، به‌طور غیرمستقیم فرضیه پژوهشی تأیید می‌شود. اگر فرضیه صفر در واقع صحیح باشد ولی محقق تصمیم به رد آن بگیرد خطای نوع اول رخ داده است. بر عکس اگر فرضیه صفری در واقع فرضیه‌ای غیر صحیح باشد ولی محقق آن را تأیید کند، دچار خطای نوع دوم شده است.

p-value

تقریباً همه نرم افزارهای آماری به جای اینکه آزمون‌ها را با توجه به عددی به نام (**p-value**) معروف به سطح معنی داری (که در جدول‌های خروجی نرم افزار با **Significant Level** مشاهده می‌کنید) را محاسبه می‌کند. اگر فرض کنیم شما می‌خواهید آزمونی را در سطح ۹۸ درصد انجام دهید (یعنی آلفا ۲ درصد باشد)،

اگر **p-value** کمتر از ۲ درصد باشد، فرض صفر رد می‌شود و اگر **p-value** بیشتر باشد فرض صفر رد نمی‌شود. اگر می‌خواهید آزمونی را در سطح ۹۵ درصد انجام بدهید (یعنی آلفا ۵ درصد باشد)، اگر **p-value** کمتر از ۵ درصد باشد، فرض صفر رد می‌شود و اگر بیشتر باشد فرض صفر رد نمی‌شود.

گروه: یک متغیر می‌تواند به لحاظ بررسی یک ویژگی خاص در یک گروه و یا دو و یا بیشتر مورد بررسی قرار گیرد.

نکته ۱: دو گروه می تواند وابسته و یا مستقل باشد. دو گروه وابسته است اگر ویژگی یک مجموعه افراد قبل و بعد از وقوع یک عامل سنجیده شود. مثلاً میزان رضایت شغلی کارکنان قبل و بعد از پرداخت پاداش و همچنین اگر در مطالعات تجربی افراد از نظر برخی ویژگی ها در یک گروه با گروه دیگر همسان شود.

جامعه نرمال: جامعه ای است که از توزیع نرمال تبعیت می کند.

توزیع نرمال: یکی از مهمترین توزیع ها در نظریه احتمال است. و کاربردهای بسیاری در علوم دارد.

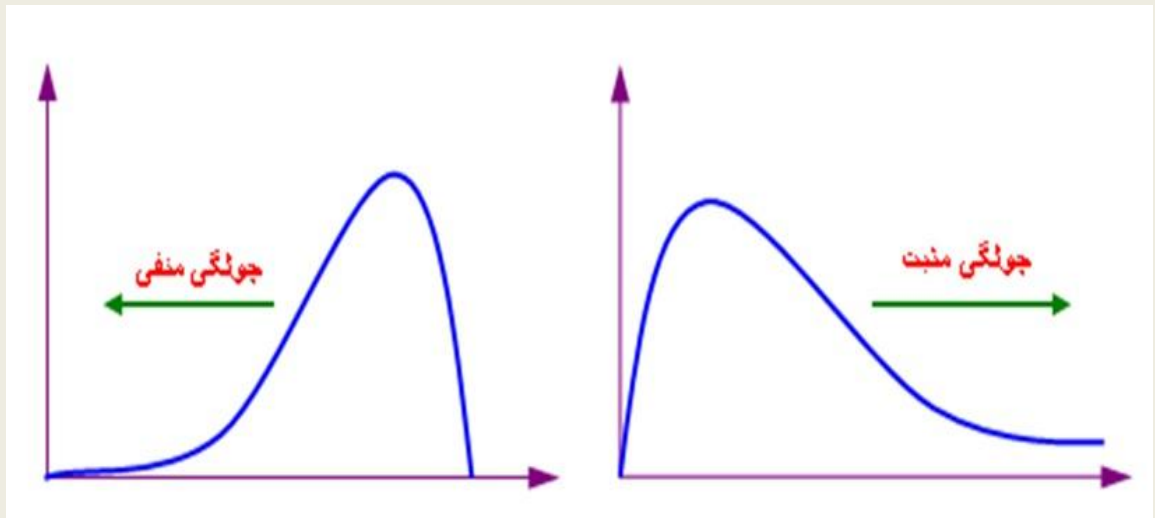
فرمول این توزیع بر حسب دو پارامتر امید ریاضی و واریانس بیان می شود. منحنی رفتار این تابع تا حد زیادی شبیه به زنگ های کلیسا می باشد. این منحنی دارای خواص بسیار جالبی است برای مثال نسبت به محور عمودی متقارن می باشد، نیمی از مساحت زیر منحنی بالای مقدار متوسط و نیمه دیگر در پایین مقدار متوسط قرار دارد و اینکه هرچه از طرفین به مرکز مختصات نزدیک می شویم احتمال وقوع بیشتر می شود. سطح زیر منحنی نرمال برای مقادیر متفاوت مقدار میانگین و واریانس فراگیری این رفتار آنقدر زیاد است که دانشمندان اغلب برای مدل کردن متغیرهای تصادفی که با رفتار آنها آشنایی ندارند، از این تابع استفاده می کنند. به عنوان مثال در یک امتحان درسی نمرات دانش آموزان اغلب اطراف میانگین بیشتر می باشد و هر چه به سمت نمرات بالا یا پایین پیش برویم تعداد افرادی که این نمرات را گرفته اند کمتر می شود. این رفتار را بسهولت می توان با یک توزیع نرمال مدل کرد.

اگر یک توزیع نرمال باشد مطابق قضیه چسبی بی شف ۲۶,۶۸٪ مشاهدات در فاصله میانگین، مثبت و منفی یک انحراف معیار قرار دارد. و ۴۴,۹۵٪ مشاهدات در فاصله میانگین، مثبت و منفی دو انحراف معیار قرار دارد. و ۷۳,۹۹٪ مشاهدات در فاصله میانگین، مثبت و منفی سه انحراف معیار قرار دارد.

نکته ۱: واضح است که داده های رتبه ای دارای توزیع نرمال نمی باشند.

نکته ۲: وقتی داده ها کمی هستند و تعداد نمونه نیز کم است تشخیص نرمال بودن داده ها توسط آزمون کولموگروف – اسمیرنوف مشکل خواهد شد.

چولگی و کشیدگی: علاوه بر گرایش مرکزی تقارن یک مشخصه مهم توزیع است.



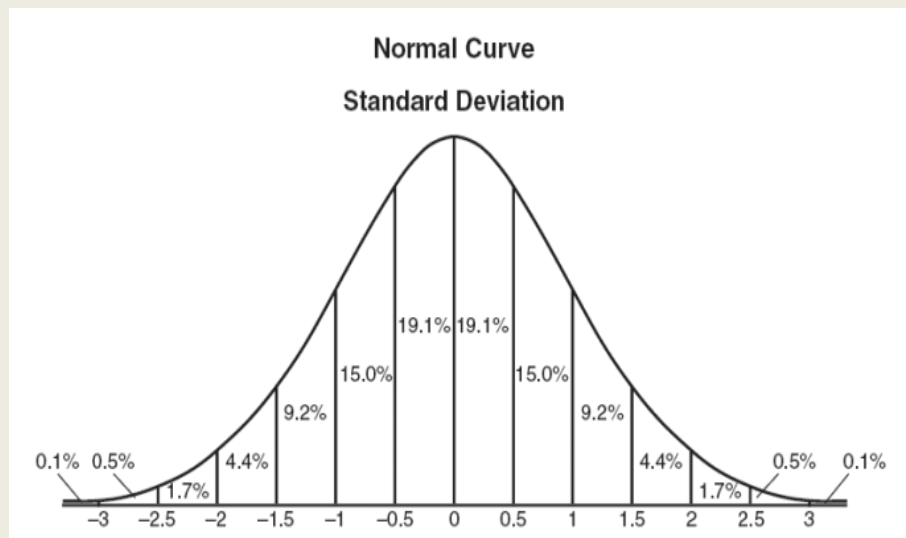
چولگی معیاری از تقارن یا عدم تقارن تابع توزیع می‌باشد. میانگین یک متغیر دارای چولگی در مرکز قرار ندارد.

برای یک توزیع کاملاً متقارن چولگی صفر و برای یک توزیع نامتقارن با کشیدگی به سمت مقادیر بالاتر چولگی مثبت و برای توزیع نامتقارن با کشیدگی به سمت مقادیر کوچکتر مقدار چولگی منفی است. کشیدگی یا **kurtosis** نشان دهنده ارتفاع یک توزیع است. به عبارت دیگر کشیدگی معیاری از بلندی منحنی در نقطه ماکزیمم است و مقدار کشیدگی برای توزیع نرمال برابر ۳ می‌باشد. کشیدگی مثبت یعنی قله توزیع مورد نظر از توزیع نرمال بالاتر و کشیدگی منفی نشانه پایین تر بودن قله از توزیع نرمال است. برای مثال در توزیع t که پراکندگی داده‌ها بیشتر از توزیع نرمال است، ارتفاع منحنی کوتاه تر از منحنی نرمال است.

قبل از هر گونه آزمونی که با فرض نرمال بودن داده‌ها صورت می‌گیرد باید آزمون نرمال بودن صورت گیرد. در حالت کلی چنانچه چولگی و کشیدگی در بازه $(-2, 2)$ نباشند داده‌ها از توزیع نرمال برخوردار نیستند.

توزیع نرمال:

این توزیع دارای دو پارامتر است که یکی تعیین کننده مکان (μ) و دیگری تعیین کننده مقیاس (σ) توزیع هستند. همچنین میانگین توزیع با پارامتر مکان و پراکندگی آن با پارامتر مقیاس برابر است. منحنی تابع احتمال حول میانگین توزیع متقارن است. در حالت خاص اگر $\mu = 0$ و $\sigma = 1$ باشد توزیع، نرمال استاندارد نامیده می‌شود



آزمون کولموگروف-اسمیرنوف

پس از بررسی نرمال بودن کشیدگی و یا چولگی توزیع داده‌ها، از آزمون کولموگروف-اسمیرنوف استفاده می‌شود تا از نرمال بودن داده‌ها اطمینان حاصل گردد. هنگام بررسی نرمال بودن داده‌ها ما فرض صفر مبتنی بر اینکه توزیع داده‌ها نرمال است را در سطح خطای ۵٪ تست می‌کنیم. بنابراین اگر آماره آزمون بزرگتر مساوی ۰٫۰۵ بدست آید، در این صورت دلیلی برای رد فرض صفر مبتنی بر اینکه داده نرمال است، وجود نخواهد داشت. به عبارت دیگر توزیع داده‌ها نرمال خواهد بود. برای آزمون نرمالیت فرض‌های آماری به صورت زیر تنظیم می‌شود:

H₀ : توزیع داده‌های مربوط به هر یک از متغیرها نرمال است

H₁ : توزیع داده‌های مربوط به هر یک از متغیرها نرمال نیست .

چنانچه سطح معناداری در آزمون کولموگروف-اسمیرنوف که در این جدول با **sig** نمایش داده می‌شود بیشتر از ۰٫۰۵ باشد می‌توان داده‌ها را با اطمینان بالایی نرمال فرض کرد. در غیر این صورت نمی‌توان گفت که داده‌ها توزیع‌شان نرمال است.

ماهیت آمار استنباطی

نقش آمار توصیفی در واقع، جمع‌آوری، خلاصه کردن و توصیف اطلاعات کمی به دست آمده از نمونه‌ها یا جامعه‌ها است. اما محقق معمولاً کار خود را با توصیف اطلاعات پایان نمی‌دهد، بلکه سعی می‌کند آنچه را که از بررسی گروه نمونه به دست آورده است به گروه‌های مشابه بزرگتر تعمیم دهد. تئوری‌های روان‌شناسی از طریق تعمیم نتایج یک یا چند مطالعه به آنچه که ممکن است در مورد کل افراد جامعه صادق باشد به وجود می‌آیند. از طرف دیگر در اغلب موارد مطالعه تمام اعضای یک

جامعه ناممکن است. از اینرو محقق به شیوه‌هایی احتیاج دارد که بتواند با استفاده از آنها نتایج به دست آمده از مطالعه گروه‌های کوچک را به گروه‌های بزرگتر تعمیم دهد. به شیوه‌هایی که از طریق آنها ویژگی‌های گروه‌های بزرگ براساس اندازه‌گیری همان ویژگی‌ها در گروه‌های کوچک استنباط می‌شود آمار استنباطی گفته می‌شود. استنباط آماری که در واقع یک نوع نتیجه‌گیری کلی از جزء به کل است، در معرض آزمایش و خطاست. یک جنبه از استنباط آماری محاسبه برآوردهایی (Estimates) از پارمترهای جامعه مانند میانگین جامعه از طریق آماره‌های نمونه ها مانند میانگین نمونه است..

آزمون‌های آمار استنباطی

آزمون‌های آماری مورد استفاده جهت تجزیه و تحلیل اطلاعات به دست آمده از یک گروه کوچک (نمونه) و تعمیم آن به جامعه مورد نظر با توجه به مقیاس اندازه‌گیری متغیرها، به دو گروه "پارامتریک" و "ناپارامتریک" تقسیم می‌شوند..

فنون آمار پارامتری شدیداً تحت تاثیر مقیاس سنجش متغیرها و توزیع آماری جامعه است اگر متغیرها از نوع فاصله ای و نسبی باشند در صورتیکه فرض شود توزیع آماری جامعه نرمال یا بهنجار است از روشهای پارامتریک استفاده می شود، که حداقل شاخص آماری آنها میانگین و واریانس است در غیراینصورت از روشهای ناپارامتریک استفاده می شود..

اگر متغیرها از نوع اسمی و ترتیبی بوده حتماً از روشهای ناپارامتری استفاده می شود، که شاخص آماری آنها میانه و نما است.

آزمون پارامتری: آزمون‌های پارامتریک، آزمون‌های هستند که توان آماری بالا و قدرت پرداختن به داده های جمع آوری شده در طرح های پیچیده را دارند. در این آزمون ها داده ها توزیع نرمال دارند.

به روش‌های آماری کلاسیک نظیر آزمون ANOVA، آزمون تی، تحلیل کواریانس، ضریب هم‌بستگی پیرسون، آزمون‌های پارامتریک گفته می‌شود. زیرا فرضیه‌های مربوط به پارامتر جامعه را آزمایش می‌کنند. آزمون‌های پارامتری را می‌توان مؤثرترین آزمون‌ها دانست، اما شرط استفاده از این آزمون‌ها آن است که پیش فرض‌های اساسی آنها مراعات شود. این پیش فرض‌ها بر چگونگی توزیع جامعه و بر روش استفاده از مقیاسی که برای به‌کمیت در آوردن داده‌ها به‌کار می‌رود،

آزمون‌های پارامتریک به پیش فرض‌های زیر متکی اند:

۱- سطوح سنجش متغیرها فاصله ای یا نسبتی است (مانند طیف لیکرت)

۲- توزیع نمرة ها نرمال باشد.

۳- واریانس نمونه ها برابر باشد.

۴- موارد مشاهده شده مستقل از هم باشند

آزمون های غیر پارامتری: آزمون هائی می باشند که داده ها توزیع **غیر نرمال** داشته و در مقایسه با آزمون های پارامتری از توان تشخیصی کمتری برخوردارند. (مانند آزمون من - ویتنی و آزمون کروسکال و والیس)

نکته ۳: اگر جامعه نرمال باشد از آزمون های پارامتری و چنانچه غیر نرمال باشد از آزمون های غیر پارامتری استفاده می نمائیم.

نکته ۴: اگر نمونه بزرگ باشد، طبق قضیه حد مرکزی اگر جامعه نرمال نباشد می توان از آزمون های پارامتری استفاده نمود.

استفاده از آزمونهای ناپارامتریک در موارد زیر امکانپذیر است:

۱- وقتی که سطوح سنجش متغیرها اسمی یا ترتیبی باشد.

۲- زمانی که واریانس نمونه ها برابر نباشد.

۳- هنگامی که توزیع نمره ها نرمال نباشد.

قابل ذکر است قبل از ورود به الگوریتم انتخاب آزمون آماری بهتر است به سوالات زیر پاسخ دهیم:

۱- آیا اختلافی بین میانگین (نسبت) یک ویژگی در دو یا چند گروه وجود دارد؟

۲- آیا دو متغیر ارتباط دارند؟

۳- چگونه می توان یک متغیر را با استفاده از متغیر های دیگر پیش بینی کرد؟

۴- چه چیزی می توان با استفاده از نمونه در مورد جامعه گفت؟

پس از انتخاب آزمون آماری مناسب حال می توان با هر یک از آزمون ها به صورت تخصصی برخورد کرد:

خلاصه آزمون های پارامتریک

از پرکاربردترین آزمون های پارامتریک می توان به آزمون t و آزمون تحلیل واریانس اشاره کرد.

آزمون t :

آزمون تی یک آزمون پارامتریک است که به منظور تعیین معناداری تفاوت بین دو میانگین بکار می رود. در حقیقت خانواده‌ای از توزیع‌ها است که با استفاده از آنها فرضیه‌هایی که درباره نمونه در شرایط جامعه ناشناخته است، آزمون می‌شود. اهمیت این آزمون (توزیع) در آن است که پژوهشگر را قادر می‌سازد با نمونه‌های کوچکتر (حداقل ۲ نفر) اطلاعاتی درباره جامعه به دست آورد. آزمون t شامل خانواده‌ای از توزیع‌ها است (برخلاف آزمون Z) و این‌طور فرض می‌کند که هر نمونه‌ای دارای توزیع مخصوص به خود است و شکل این توزیع از طریق محاسبه درجات آزادی مشخص می‌شود. به عبارت دیگر توزیع t تابع درجات آزادی است و هرچه درجات آزادی افزایش پیدا کند به توزیع طبیعی نزدیکتر می‌شود. از سوی دیگر هرچه درجات آزادی کاهش یابد، پراکندگی بیشتر می‌شود. خود درجات آزادی نیز تابعی از اندازه نمونه انتخابی هستند. هرچه تعداد نمونه بیشتر باشد بهتر است. از آزمون t می‌توان برای تجزیه و تحلیل میانگین در پژوهش‌های تک‌متغیری یک‌گروهی و دوگروهی و چند متغیری دوگروهی استفاده کرد.

این آزمون برای ارزیابی میزان همقواری یا یکسان بودن و نبودن میانگین نمونه‌ای با میانگین جامعه در حالتی به کار می‌رود که انحراف معیار جامعه مجهول باشد. چون توزیع t در مورد نمونه‌های کوچک (کمتر از ۳۰) با استفاده از درجات آزادی تعدیل می‌شود، می‌توان از این آزمون برای نمونه‌های بسیار کوچک استفاده نمود. همچنین این آزمون مواقعی که خطای استاندارد جامعه نامعلوم و خطای استاندارد نمونه معلوم باشد، کاربرد دارد.

برای به کار بردن این آزمون، متغیر مورد مطالعه باید در مقیاس **فاصله‌ای** باشد، شکل توزیع آن **نرمال** و تعداد نمونه **کمتر از ۳۰** باشد.

انواع آزمون تی وجود دارد:

۱- آزمون تی یک گروهی (یک نمونه‌ای)

۲- آزمون تی مستقل

۳- آزمون تی وابسته یا جفت شده

این توزیع توسط ویلیام گاست مطرح شد. در پژوهش هایی که انحراف استاندارد جامعه مورد مطالعه ناشناخته است از آزمون تی استفاده می کنیم.

اهمیت این توزیع ها در آن است که می توان با نمونه های کوچک (حداقل ۲ نفر) اطلاعاتی درباره جامعه بدست آورد.

در واقع توزیع t بر اساس درجات آزادی نامگذاری می شود و تابع درجات آزادی هر نمونه است.

درجات آزادی

درجات آزادی به تعداد ارزشهایی گفته می شود که پس از اعمال برخی محدودیت ها می توانند آزادانه تغییر کنند.

درجات آزادی را با $d.f$ نشان میدهند. درجات آزادی برای آزمون فرضیه اختلاف بین میانگین نمونه و میانگین فرضی جامعه برابر با $n-1$ است، مثلاً اگر بخواهیم سه عدد را با این محدودیت که مجموع آنها ۲۰ باشد بنویسیم درجات آزادی برابر است با

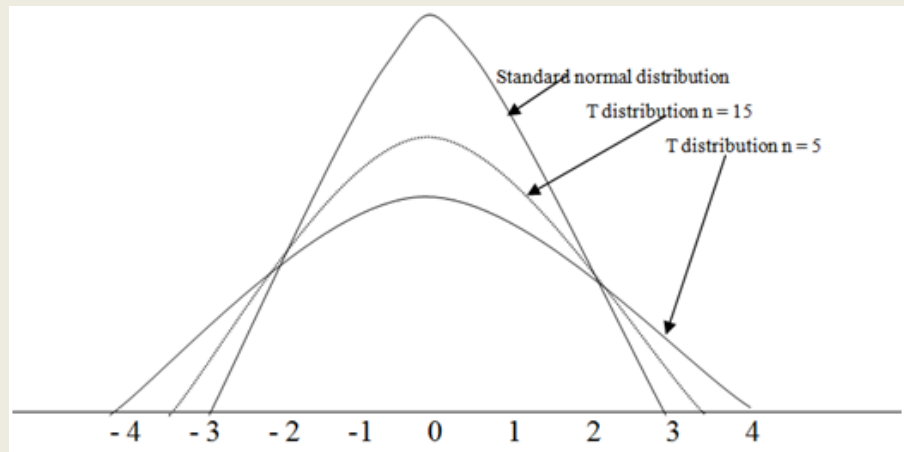
$d.f = n-1 = 3-1 = 2$. برای آزمون معنادار بودن تفاوت بین میانگین های نمونه های مستقل به صورت $n_1 + n_2 - 2$ می باشد.

ویژگی های توزیع t استودنت

توزیع های تی شبیه توزیع نرمال استاندارد هستند.

میانگین آنها صفر و انحراف استاندارد آنها کمی بزرگتر از یک است.

با افزایش درجات آزادی انحراف استاندارد توزیع به یک نزدیک می شود. توزیع تی مانند توزیع نرمال به منظور تعیین احتمال بکار می رود. در توزیع های تی هرچه درجات آزادی کوچکتر باشند مقدار t جدول بزرگتر خواهد بود و بر عکس.



انواع آزمون t :

آزمون t یک گروهی: (One- Sample T Test)

برای آزمون فرض پیرامون میانگین یک جامعه استفاده می شود.

شرایط آزمون t تک نمونه ای:

این آزمون، زمانی مورد استفاده قرار می گیرد که یک نمونه از جامعه داریم و می خواهیم میانگین آن را با یک حالت معمول و رایج، استاندارد و یا حتی یک عدد فرضی و مورد انتظار مقایسه کنیم. هرچند در عمل اغلب ممکن نیست که بدانیم آیا جامعه ای که نمونه از آن بیرون آمده به واقع دارای توزیع نرمال هست یا نه، اما باید درستی این مفروضه را بپذیریم

مثال یک: در یک آزمون دو دامنه قصد داریم این فرضیه را که میانگین اطلاعات عمومی دانش آموزان کلاس چهارم دبستان در یک آزمون ۱۰۰ سؤالی ۹۰ است را آزمون کنیم. نمونه ای به حجم ۲۵ نفر را به صورت تصادفی انتخاب می کنیم چنانچه میانگین و انحراف استاندارد به ترتیب ۹۲ و ۸ باشد این فرضیه را که بین میانگین جامعه و نمونه تفاوت معناداری وجود ندارد آزمون می کنیم.

ابتدا خطای استاندارد میانگین را بدست می آوریم:

مقدار t را محاسبه می کنیم. $(n-1=25-1=24)$

حالا مقدار $t=1/25$ محاسبه شده با $T_{.975}=2/064$ جدول مقایسه می کنیم .

در آزمون دو دامنه جدول با سطح اطمینان $p>.05$ ، بنابراین چون t محاسبه شده از T کوچکتر است فرضیه صفر تأیید می شود.

مثال ۲:

مدیر دبیرستانی اعتقاد دارد دانش آموزان ورزشکار، از نظر مهارت کلامی پایین تر از میانگین هستند، اما بدون انجام ورزش همین دبیرستان ادعای کند که مهارت کلامی آنان به ترتیب ۱۰۴ و ۱۰۵ بدست می آید.

مربی ورزش تصمیم گرفته است با یک آزمون در سطح ۰/۰۵ این فرضیه را که بین دو میانگین جامعه و نمونه تفاوت معناداری وجود ندارد را از منکر کند. در اینجا نیز برای محاسبه t تک نمونه ای را اجرا می کنیم.

آزمون t برای دو گروه مستقل: (Independent – Samples)

جهت مقایسه میانگین دو جامعه استفاده می شود. در آزمون t برای دو نمونه مستقل فرض می شود واریانس دو جامعه برابر است. این آزمون زمانی استفاده می شود که میانگین دو گروه از پاسخگویان را با یکدیگر مقایسه می کند. از این آزمون برای محاسبه فاصله اطمینان و یا آزمون فرضیه تفاوت میانگین دو جمعیت (در زمان نامشخص بودن انحراف استاندارد و استقلال نمونه ها از یکدیگر) استفاده می شود به گروه ها یا نمونه هایی مستقل می گوئیم که انتخاب آزمونی ها یا موارد در یک نمونه در انتخاب موارد یا آزمونی ها ی گروه دیگر دخالت نداشته باشد.

مثلاً: بررسی تفاوت طول قد مردان و زنان که دو گروه کاملاً مستقل هستند.

شرایط آزمون t برای دو گروه مستقل:

۱- گروه ها را به صورت تصادفی و مستقل از یکدیگر بر می گزینیم

۲- توزیع هر یک از دو جامعه مورد نظر نرمال است

۳- تغییر پذیری (واریانس) نمره ها در هر دو جامعه یکسان است.

مثال ۱- پژوهشگری قصد دارد تأثیر دو روش مختلف در یادگیری درس ریاضی دانش آموزان کلاس سوم را بیازماید . بنابراین ابتدا دونه نمونه به صورت تصادفی از بین آنان انتخاب می کند و گروه اول را با روش A و گروه دوم را با روش B بکار می برد . (آموزش می دهد)

Ho عبارت است از اینکه اختلاف بین میانگین های دو جامعه در روش A و B صفر است (آزمون دو دامنه است با سطح معنا داری $(\alpha/2)$).

مثال ۲: پژوهشگری که یک متخصص اندازه گیری است مایل است اثر انگیزش دانشجویان را در نتایج شخصیت سنج های استاندارد بداند . او از میان دانشجویان دو گروه را به صورت تصادفی بر می گزیند از گروه اول می خواهد که پرسش نامه حساسیت اجتماعی ها رتمن را تکمیل کنند و به آن ها گفته می شود که با حداکثر تلاش پاسخ دهند و از گروه دوم خواسته می شود که بدون ذکر نام به پرسش نامه پاسخ دهند زیرا پژوهشگر می خواهد اطلاعاتی در مورد کل دانشجویان بدست آورد . و مایل است این فرضیه (صفر) را بیازماید که در نمره متوسط این شخصیت سنج برای هر دو جامعه هایی که گروه های نمونه معرف آن هاست نرمال است . بنابراین آزمون t برای دو گروه مستقل را مشخص های استنباطی مناسبی برای آزمون فرضیه می داند . او همچنین می بیند که توزیع نمره های دو نمونه ، تک نمایی و متقارن است. در نتیجه به استفاده از یک آزمون t برای دو گروه مستقل هدایت می شود .

آزمون t برای گروه های وابسته (زوجی): (Paired – Samples T Test)

برای آزمون فرض پیرامون دو میانگین از یک جامعه استفاده می شود. در طرح های پیش آزمون – پس آزمون (اندازه گیری مکرر) و طرح های گروه های جفت شده یا همتا شده بکار برد. برای مثال ، مقایسه وضعیت بیمار در زمان های قبل و بعد از استفاده از یک دارو،

شرایط t برای گروه های وابسته:

۱- از هر زوج مشاهده مستقل از هر زوج دیگر باشد

۳- متغیر مورد مطالعه فاصله ای باشد

۴- توزیع تفاوت ها در جامعه نرمال یا تقریباً نرمال است.

مثال: فرض کنید می خواهیم فشارخون بیماران را قبل و بعد از رژیم غذایی خاصی اندازه بگیریم. با احتمال ۹۵٪ آیا رژیم غذایی جدید بر فشار خون تاثیر داشته یا نه؟

فشارخون بعد از رژیم غذایی	فشار خون قبل از رژیم غذایی	شماره بیماران
۱۴۰	۱۷۰	۱
۱۶۰	۱۷۰	۲
۱۵۰	۱۴۰	۳
۱۶۰	۱۴۰	۴
۱۵۰	۱۷۰	۵

مثال: فرض کنید پژوهش گری علاقه مند است تاثیر نوعی روش تدریس را در به خاطر سپردن کلمات بی معنی آزمون کند. او ابتدا میزان توانایی آزمونی ها را به خاطر سپرده کلمات اندازه گیری می کند. سپس روش تدریس مورد نظر را اجرا می کند و مجدداً توانایی هر یک از آزمودنی ها را اندازه می گیرد و در انتها برای تعیین معناداری نتایج پیش آزمون و پس آزمون را با هم مقایسه می کند

طرح جفت های همتراز شده: در این طرح آزمودنی ها هر نمونه (آزمایش و کنترل) براساس یک یا چند متغیر که با متغیر وابسته رابطه دارند همتراز می شوند.

مثال: پژوهشگری روشی تهیه کرده است که می تواند دانش آموزان دبیرستان را که در خواندن ضعیف هستند تشویق به خواندن بیشتر کند. این روش مستلزم استفاده از کتابهای خاصی است که علاوه بر مواد خواندنی دارای تصاویری است. او معتقد است که استفاده از این روش در طول یک دوره چند ماهه موجب می شود پیشرفت خواندن شاگردانی که برای

جبرانی را طی کنند بهبود یابد. به منظور تایید این نظریه، او گروهی از شاگردان را که مشغول برنامه آموزشی جبرانی هستند به صورت تصادفی انتخاب می کند (۲۵ نفر). او ابتدا از آنها یک پیش آزمون می گیرد تا سطح اولیه پیشرفت خواندن آنها را بدست آورد سپس مواد جدید را برای یک دوره دو ماهه روی آزمودنی ها بکار می برد. و در پایان یک پس آزمون می گیرد تا سطح پیشرفت خواندن آنها را پس از اجرا تعیین کند. در این موقعیت برای آزمون این فرضیه صفر که مداخله تجربی پژوهشگر تاثیری در پیشرفت خواندن دانش آموزان برنامه جبرانی ندارد در برابر فرضیه مخالف که مداخله تجربی پیشرفت خواندن را بهبود می بخشد، استفاده از آزمون t برای دو گروه نمونه وابسته مناسب می باشد.

آنالیز واریانس (ANOVA): از این آزمون به منظور بررسی اختلاف میانگین چند جامعه آماری استفاده می شود.

واریانس شاخصی است که پراکندگی داده ها را پیرامون میانگین نشان میدهد. زمانی که پژوهشگری بخواهد بیش از دو میانگین (بیش از دو نمونه) را با هم مقایسه کند، باید از تحلیل واریانس استفاده کند. تحلیل واریانس روشی فراگیرتر از آزمون t است و برخی پژوهشگران حتی وقتی مقایسه میانگین های دو نمونه مورد نظر است نیز از این روش استفاده می کنند. طرح های متنوعی برای تحلیل واریانس وجود دارد و هر یک تحلیل آماری خاص خودش را طلب می کند. از جمله این طرح ها می توان به تحلیل یک عاملی واریانس (تحلیل واریانس یک راهه) و تحلیل عاملی متقاطع واریانس، تحلیل واریانس چندمتغیری، تحلیل کوواریانس یک متغیری و چندمتغیری و ... اشاره کرد

آنالیز واریانس یک شیوه آماری کارآمد برای مقایسه میانگین یک صفت کمی در سطوح یک متغیر مقوله ای است. به عنوان مثال وقتی می خواهیم سه روش آموزش ریاضیات را در یک آموزشگاه با هم مقایسه کنیم، باید میانگین نمرات سه گروه از دانشجویان را که هر گروه با یک روش مختلف تدریس شده اند، با هم مقایسه شوند. در این جا به متغیر روش آموزش، متغیر فاکتور یا متغیر مستقل می گویند. این متغیر به عنوان مثال شامل سه سطح باشد. امکان دستکاری سطوح متغیر مستقل در آنالیز واریانس وجود دارد. آن چیزی که باعث تفاوت بین آنالیز واریانس و رگرسیون است، همین موضوع است که در رگرسیون، متغیرهای مستقل قابل دستکاری نیستند ولی در آنالیز واریانس می توان سطوح متغیر مستقل را دستکاری کرد.

از این آزمون به منظور بررسی اختلاف میانگین چند جامعه آماری استفاده می شود. برای نمونه جهت بررسی معنی دار بودن تفاوت میانگین نمره نظرات پاسخ دهندگان بر اساس سن یا تحصیلات در خصوص هر یک از فرضیه های پژوهش استفاده می شود.

مثال: متوسط زمان بستری شدن بیماران در ماه شهریور برای یک بیماری خاص در ۵ بیمارستان به صورت زیر می باشد؛ بررسی کنید که آیا میان متوسط زمان بستری شدن (به روز) بیماران ۵ بیمارستان تفاوت معنی داری وجود دارد یا نه. در صورت وجود اختلاف نشان دهید که میان کدام بیمارستانها در این زمینه تفاوت وجود دارد. (تیمار = متوسط زمان بستری شدن بیماران)

تعداد تیمار	۱	۲	۳	۴	۵	۶	۷	۸	۹	۱۰
بیمارستان A	۷	۷	۸	۶	۷	۵	۸	۷	۶	۸
بیمارستان B	۵	۶	۶	۶	۷	۷	۸	۸	۸	۸
بیمارستان C	۵	۵	۴	۴	۷	۴	۵	۵	۵	۷
بیمارستان D	۸	۹	۹	۱۱	۶	۱۰	۱۱	۱۱	۱۰	۱۲
بیمارستان E	۴	۹	۶	۴	۴	۴	۵	۵	۴	۶

آزمون تحلیل واریانس یکطرفه: آزمونی که برای مقایسه میانگین یک صفت کمی در بیش از دو گروه استفاده می شود. آزمون تحلیل واریانس یکطرفه گسترش یافته آزمون t دو نمونه مستقل است. یعنی زمانی که بخواهیم به مقایسه میانگین های دو گروه بپردازیم از آزمون t استفاده می کنیم، ولی اگر مقادیر متغیر مستقل سه وجهی یا بیشتر باشد از آزمون تحلیل واریانس یکطرفه بهره میگیریم. استفاده از این آزمون نیازمند رعایت شرایط زیر است:

۱- توزیع داده ها نرمال باشد.

۲- داده های انتخاب شده مستقل از هم باشند .

۳- واریانس داده ها یکسان باشند .

پیش فرض این آزمون آن است که آیا بین میانگین های نمونه تحقیق تفاوت معنی داری وجود دارد یا نه؟ یعنی واریانس مشاهده شده بین میانگین ها را با واریانس ناشی از تصادف مقایسه می کند .

ضریب همبستگی گشتاوری پیرسون:

برای محاسبه همبستگی دو مجموعه داده استفاده می شود.

آزمون های ناپارامتریک آمار استنباطی

در پژوهش هایی که در سطح مقیاس های اسمی و رتبه ای اجرا می شوند، باید از آزمون های ناپارامتریک برای تجزیه و تحلیل اطلاعات استفاده شود. آزمون های زیادی برای این امر وجود دارد که براساس نوع تحلیل (نیکویی برازش، همسویی دو نمونه مستقل، همسویی دو نمونه وابسته، همسویی K نمونه مستقل و همسویی K نمونه وابسته) و مقیاس اندازه گیری می توان دست به انتخاب زد. از آزمون های مورد استفاده برای پژوهش ها در سطح اسمی می توان به آزمون X²، آزمون تغییر مک نمار، آزمون دقیق فیشر و آزمون کاکرن اشاره کرد. از آزمون های مورد استفاده برای پژوهش ها در سطح رتبه ای می توان به آزمون های کولموگروف – اسمیرونوف، آزمون تقارن توزیع، آزمون علامت، آزمون میانه، آزمون مان – ویتنی، آزمون تحلیل واریانس دو عاملی فریدمن و ... اشاره کرد.

اشاره به معادل های آزمون t در آمار ناپارامتریک

در آمار نا پارامتریک برای تعیین معناداری میانگین هابین دو گروه بجای آزمون t از آزمون های U مان و ویتنی و ویل کاکسون استفاده می شود.

۱- آزمون U مان ویتنی:

این آزمون معادل غیر پارامتریک آزمون t مستقل است و برای مقایسه داده هایی که از طرح های گروه های مستقل بدست می آیند مورد استفاده قرار می گیرد.

۲- آزمون ویل کاکسون :

این آزمون ها معادل غیر پارامتریک آزمون t جفت شده است و برای داده هایی که از اندازه گیری مکرر و طرح های زوج های همتا شده بدست می آیند بکار می رود .

در هر دو این آزمون ها رتبه بندی صورت می گیرد و مخاسبات روی رتبه ها انجام می گیرد .

شرایط بکار گیری این آزمون ها به جای آزمون t

۱- هنگامی که داده ها فقط ترتیبی هستند.

۲- هنگامی که داده ها فاصله ای یا نسبی هستند ولی دارای توزیع غیرنرمال می باشند.

۳- هنگامی که داده ها فاصله ای یا نسبی هستند اما واریانس های دو نمونه برابر نیستند.

همبستگی بین متغیرها:

یکی از تعاریف اساسی در علم آمار تعریف همبستگی و رابطه بین دو متغیر می باشد. بطور کلی شدت وابستگی دو متغیر به یکدیگر را همبستگی تعریف می کنیم. و ممکن علاوه بر شدت همبستگی جهت همبستگی نیز مورد نیاز پژوهشگر باشد. در آمار انواع زیادی از ضرایب همبستگی متفاوت وجود دارند که هر کدام همبستگی بین دو متغیر را با توجه به نوع داده ها و شرایط متغیرها اندازه گیری می کنند. لذا با توجه به اهمیت این موضوع که چه ضریب همبستگی را در چه زمانی مورد استفاده قراردهیم، به عنوان مثال مقدار هموگلوبین و تعداد گلبولهای قرمز خون وقتی از هم مستقلند که توزیع هموگلوبین به ازای مقادیر مختلف از تعداد گلبول قرمز و برعکس یکسان باشند . بعبارت دیگر با تغییر تعداد گلبول قرمز توزیع هموگلوبین و یا با تغییر مقدار هموگلوبین توزیع تعداد گلبول قرمز تغییر نکند. در صورتی که شرط فوق برقرار نباشد دو صفت وابسته اند. همانطور که دیدیم مطالعه ارتباط بین دو صفت در یک جامعه هنگامی مطرح می شود که افراد به طور تصادفی از این جامعه انتخاب شوند و برای هر یک از این افراد هر دو صفت اندازه گیری شده کمیت تصادفی باشند. در اینجا قصد داریم به تعریف انواع همبستگی بپردازیم.

انواع ضرایب همبستگی

محاسبه ضرایب همبستگی تا حدود زیادی متأثر از مقیاس اندازه گیری متغیر ها است، بعنوان مثال برای متغیرهای اسمی جهت رابطه اصلا معنی ندارد، بین جنس و معدل تنها می توان گفت که شدت وابستگی چه مقدار است اما افزایش یا کاهش جنس معنی ندارد.

با توجه به نوع متغیر ها ضریب همبستگی می تواند یکی از حالت های زیر را داشته باشد.

۱- دو متغیر اسمی

۲- دو متغیر رتبه ای

۳- دو متغیر فاصله ای - نسبی

۴- متغیر اسمی و متغیر رتبه ای

۵- متغیر اسمی و متغیر فاصله ای - نسبی

۶- متغیر رتبه ای و متغیر فاصله ای - نسبی

برای هر کدام از حالت های بالا ضرایب همبستگی متفاوتی وجود دارند .

هرگاه رابطه بین دو متغیر را بررسی کردیم و بین آن دو رابطه معنی دار وجود داشت می توانیم ضریب همبستگی و میزان شدت همبستگی را نیز محاسبه کنیم .

ضرایب همبستگی در واقع وابستگی دو متغیر را مشخص می کنند:

اگر ضریب همبستگی بین ۰/۲۵ تا ۰/۳۵ ضریب بسیار پایین است - تنها ۴٪ تغییرات مشترک میان دو متغیر را نشان می دهد

اگر ضریب همبستگی بین ۰/۳۵ تا ۰/۶۵ ضریب متوسط است - حدود ۲۵٪ تغییرات مشترک میان دو متغیر را نشان می دهد

اگر ضریب همبستگی بین ۰/۶۵ تا ۰/۸۵ ضریب بالایی است - تا ۷۲٪ تغییرات مشترک میان دو متغیر را نشان می دهد

انواع ضرایب همبستگی:

ضریب همبستگی پیرسون ۲: این ضریب میزان همبستگی بین دو متغیر **فاصله ای یا نسبی** را محاسبه کرده مقدار آن بین ۱+ و ۱- می باشد اگر مقدار بدست آمده مثبت باشد به معنی این است که تغییرات دو متغیر به طور هم جهت اتفاق می افتد یعنی با افزایش در هر متغیر، متغیر دیگر نیز افزایش می یابد و برعکس اگر مقدار ۲ منفی شد یعنی اینکه دو متغیر در جهت

عکس هم عمل می کنند یعنی با افزایش مقدار یک متغیر مقادیر متغیر دیگر کاهش می یابد و برعکس. اگر مقدار بدست آمده صفر شد نشان می دهد که هیچ رابطه ای بین دو متغیر وجود ندارد و اگر +۱ شد همبستگی مثبت کامل و اگر -۱ شد همبستگی کامل و منفی است.

ضریب همبستگی رتبه ای اسپیرمن : آن را با علامت **P** نشان می دهند و همواره بین +۱ و -۱ می باشد از لحاظ سطح سنجش **ترتیبی** است. در صورتی که داده های ما به صورت فاصله ای و نسبی باشند می توانیم آنها را به رتبه تبدیل کنیم. مهم نیست کدام متغیر وابسته و کدام متغیر مستقل است.

نکته مهم: اگر ضریب همبستگی صفر باشد نشان دهنده عدم وجود همبستگی است

ضریب همبستگی همواره بین +۱ و -۱ در نوسان است

اگر ضریب همبستگی کمتر از صفر باشد همبستگی ناقص و منفی است یعنی با افزایش یک متغیر دیگری کاهش می یابد

اگر ضریب همبستگی بزرگتر از صفر باشد ناقص و مثبت است یعنی با افزایش یک متغیر، دیگری نیز افزایش می یابد

اگر صفر باشد نشان دهنده عدم وجود همبستگی است

به طور مثال رابطه میان فشارخون سیستولیک مردان ۶۰ سال به بالا و سن آنها یا در مطالعه تحولات اجتماعی به دنبال رابطه درآمد سرپرست خانواده و میزان تحصیل فرزندان می باشیم و مثالهایی از این دست.....

در جدول زیر میزان معدل دانش آموزان و میزان تحصیلات آنها آمده است. هدف تعیین میزان همبستگی و نوع ارتباط معدل با میزان تحصیلات مادر می باشد

تحصیلات معدل	(۱) زیر دیپلم	(۲) دیپلم	(۳) فوق دیپلم	(۴) لیسانس و بالاتر
(۱) زیر ۱۲	۱۰	۷	۶	۲
(۲) ۱۲-۱۵	۱۲	۷	۵	۵
(۳) ۱۵-۱۷	۵	۴	۶	۱۰
(۴) ۱۷-۲۰	۵	۷	۱۰	۱۲