

به نام خداوند بخشنده مهربان



آمار و روش تحقیق

مجموعه:

(روانشناسی)

مؤلف:

دکتر جعفر مبینی

ارشد و دکتری

آمادگی آزمون ارشد و دکتری

دکتر حبیبی، جعفر
آمار و روش تحقیق / دکتر جعفر حبیبی
تهران - مشاوران صعود ماهان: ۱۴۰۴
۱۸۱ص: جدول، نمودار، (آمادگی آزمون ارشد و دکتری)
شابک
ISBN/N: 978-600-458-700-6
وضعیت فهرست نویسی: فیپا مختصر
فارسی - چاپ چهارم
۱- آمار و روش تحقیق
۲- آزمونها و تمرینها
۲- آزمون دوره های تحصیلات تکمیلی
۴- دانشگاه ها و مدارس عالی - ایران - آزمونها
ج - عنوان



انتشارات مشاوران صعود ماهان



آمار و روش تحقیق	☐ نام کتاب:
مجید و هادی سیاری	☐ مدیران مسئول:
دکتر جعفر حبیبی	☐ مولف:
سمیه بیگی	☐ مدیر برنامه ریزی و تولید محتوا:
مشاوران صعود ماهان	☐ ناشر:
چهارم/۱۴۰۴	☐ نوبت و تاریخ چاپ:
۱۰۰۰ نسخه	☐ تیراژ:
۵/۵۰۰/۰۰۰ ریال	☐ قیمت:
ISBN ۹۷۸-۶۰۰-۴۵۸-۷۰۰-۶	☐ شابک:

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

تلفن: ۸۸۱۰۰۱۱۳ و ۸۴۰۱۳۱۳

کلیه حقوق مادی و معنوی این اثر متعلق به موسسه آموزش عالی آزاد ماهان می باشد. و هرگونه اقتباس و

کپی برداری از این اثر بدون اخذ مجوز پیگرد قانونی دارد.

بنام خدا

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

شاد باشید و دلی را شاد کنید

برادران سیاری

	مقدمه	
	پیشگفتار	
۷	بخش اول آمار	
۷	فصل اول شاخص‌های توزیع فراوانی و نمودارها	
۱۶	فصل دوم شاخص‌های مرکزی	
۲۴	فصل سوم شاخص‌های پراکندگی	
۳۱	فصل چهارم نمرات استاندارد و منحنی طبیعی	
۳۷	فصل پنجم همبستگی	
۴۳	فصل ششم پیش‌بینی	
۴۸	فصل هفتم برآورد (تعمیم)	
۵۳	فصل هشتم آزمون فرضیه	
۵۸	فصل نهم آزمون‌های آماری ناپارامتریک (غیر مشروط)	
۶۶	فصل دهم آزمون‌های پارامتریک (مشروط)	
۷۲	فصل یازدهم آزمون مقایسه چندگروهه (مقایسه چند میانگین)	
۸۶	فصل دوازدهم همبستگی چندمتغیره و انواع رگرسیون	
۹۲	فصل سیزدهم احتمالات	
۹۹	فصل چهاردهم پیوست	
۱۱۸	بخش دوم روش تحقیق	
۱۱۹	فصل اول متغیرها	
۱۲۳	فصل دوم تحقیقات و منابع شناخت	
۱۳۱	فصل سوم موضوع، سؤال و فرضیه‌های تحقیق	
۱۳۶	فصل چهارم ابزارهای اندازه‌گیری و طیف‌های سنجش	
۱۴۵	فصل پنجم نمونه‌گیری	
۱۴۹	فصل ششم تحقیقات کیفی	
۱۵۶	فصل هفتم روش‌های تحقیق کیفی	
۱۶۴	فصل هشتم روش‌های تحقیق کمی [از نوع توصیفی]	
۱۷۳	فصل نهم روش‌های تحقیق کمی [از نوع تجربی]	

آمار و روش تحقیق یکی از دروس تأثیرگذار برای کسب موفقیت در مقاطع ارشد و دکتری بسیاری از رشته‌ها از جمله روان‌شناسی و علوم تربیتی است. این درس دارای سرفصل‌هایی مشترک در تمامی گرایش‌های روان‌شناسی و علوم تربیتی بوده و در چند سال اخیر، سطح سؤالات کنکورهای ارشد و دکتری نیز یکسان و سؤالات از مباحثی مشترک طراحی شده است.

در کتاب حاضر سعی بر آن شده است که مطالب آماری و پژوهشی به‌طور جامع، با بیانی روان و با استفاده از تجربه سال‌ها تدریس در دروس ریاضی، آمار و روش‌های تحقیق، خدمت شما عزیزان ارائه شود که هم قابلیت استفاده برای آمادگی کنکورهای ارشد و دکتری را داشته باشد و هم بتواند به‌عنوان کتابی راهگشا در کنار منابع اصلی در دوره‌های ارشد و دکتری رشته‌های روان‌شناسی و علوم تربیتی مورد استفاده قرار گیرد. آشنایی با روش‌های مختلف آماری و انواع روش‌های تحقیق، تنها به مرحله آمادگی کنکور محدود نشده، بلکه دانشجویان عزیز در دروس دانشگاهی در مقاطع مختلف، پایان‌نامه ارشد، آزمون جامع دکتری و همچنین در رساله دکتری، نیاز به تسلط بر آمار و تکنیک‌های تحقیق خواهند داشت. اما غالباً به دلایل مختلفی مانند ضعف پایه ریاضی، بیان نامفهوم کتاب‌ها، عدم تسلط کافی برخی از اساتید و ... یادگیری کامل، مفهومی و کاربردی این درس به مشکلی بزرگ برای اغلب دانشجویان علوم انسانی تبدیل شده است. لذا تلاش شده است نکات آماری و پژوهشی، به‌صورت مفهومی و با رویکردی کاربردی خدمت شما عزیزان ارائه شود.

ترتیب ارائه مطالب در کتاب حاضر، کمی متفاوت با کتب دیگر (منابع اصلی و کمک‌آموزشی) بوده و سعی بر آن شده تا بهترین نظم ممکن را برای یادگیری فراهم سازد. به‌طور خاص در قسمت روش‌های تحقیق، مباحث به‌گونه‌ای مطرح شده است که دانشجو بعد از آشنایی با مقدمات تحقیق، با دسته‌بندی کلی پژوهش‌های کمی و کیفی در ابتدا آشنایی کامل و عمیق پیدا کرده و سپس به ترتیب از ساده‌ترین تا پیچیده‌ترین روش تحقیق مورد بحث قرار بگیرد. به‌عنوان مثال روش تحقیق آزمایشی که به دقیق‌ترین روش پژوهشی معروف است، در آخرین فصل روش‌های تحقیق مورد بحث قرار گرفته است.

اگر بخواهیم مطابقت با محتوای کنکورهای ارشد و دکتری را مدنظر قرار دهیم، این کتاب دارای پوشش حداقل ۸۰ درصدی بوده و در کنار درسنامه، کتاب تست با پاسخنامه تشریحی کنکورهای ارشد و دکتری سال‌های اخیر نیز تدارک دیده شده تا نیاز عزیزان را در مسیر آمادگی این آزمون‌ها تا حدود زیادی به‌طرف سازد. لذا توصیه می‌شود در کنار این دو مجموعه کتاب و تست، اگر دانشجویان عزیز تسلط کافی پیدا کردند و علاقه به مطالعه گسترده‌تر داشتند، صرفاً از منابع اصلی (و نه کمک‌آموزشی) مانند

کتاب دلاور، فرگوسن، شیولسون، میرز و ... استفاده کنند.

در پایان شکر خالق را به‌جا آورده و از خداوند حکیم مسئلت دارم تا این توفیق را نصیب بنده کرده تا بتوانم برای دانش‌پژوهان عزیز، معلمی سالم، آگاه و راهنما باشم.
جعفر حبیبی - مدرس آمار، روش‌های تحقیق و نرم‌افزارهای آماری

بخش اول

آمار

فصل اول شاخص‌های توزیع فراوانی و نمودارها

فصل دوم شاخص‌های مرکزی

فصل سوم شاخص‌های پراکندگی

فصل چهارم نمرات استاندارد و منحنی طبیعی

فصل پنجم همبستگی

فصل ششم پیش‌بینی

فصل هفتم برآورد (تعمیم)

فصل هشتم آزمون فرضیه

فصل نهم آزمون‌های آماری ناپارامتریک (غیر مشروط)

فصل دهم آزمون‌های پارامتریک (مشروط)

فصل یازدهم آزمون مقایسه چندگروهه (مقایسه چند میانگین)

فصل دوازدهم همبستگی چندمتغیره و انواع رگرسیون

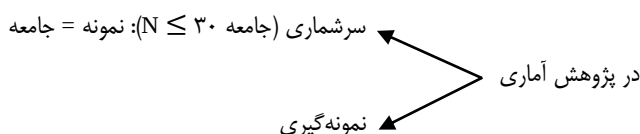
فصل سیزدهم احتمالات

فصل چهاردهم پیوست

فصل اول

شاخص‌های توزیع فراوانی و نمودارها

تعریف آمارهای توصیفی و استنباطی



- در پژوهش‌های سرشماری آمار توصیفی کفایت می‌کند.
- آمار توصیفی نمونه را توصیف می‌کند.
- در تحقیقاتی که نمونه‌گیری انجام می‌شود، بعد از توصیف نمونه نتایج آن را به جامعه تعمیم می‌دهیم و از آمار استنباطی استفاده می‌کنیم.

آمار توصیفی: طبقه‌بندی، تلخیص و توصیف اطلاعات حاصل از نمونه

۱. جدول توزیع فراوانی
 ۲. نمودار
 ۳. شاخص‌های مرکزی (میان / میانگین)
 ۴. شاخص‌های پراکندگی
 ۵. نمرات استاندارد و منحنی طبیعی
 ۶. همبستگی و رگرسیون
- مهم‌ترین کارکرد آمار توصیفی چیست؟ تلخیص (خلاصه کردن)
 - چگونه به هدف می‌رسد؟ داده‌های عددی را در شاخص‌های آماری مثل میانگین و واریانس خلاصه می‌کند.
 - هدف آمار توصیفی: توصیف اطلاعات حاصل از نمونه

آمار استنباطی: به تحلیل و تفسیر و تعمیم نتایج حاصل از تنظیم و محاسبه مقدماتی آماری تکیه دارد.

هدف آمار استنباطی: تعمیم اصول و یافته‌ها برای پیش‌بینی حوادث

آمار استنباطی (تحلیل آماری) به دو دسته آزمون فرضیه و تعمیم تقسیم می‌شود:

۱. آزمون فرضیه

آزمون‌های آماری مانند تحلیل واریانس وظیفه آزمودن فرضیه و بررسی معناداری را دارند؛ درواقع تأیید یا رد فرضیه‌ها را مشخص می‌کنند. به دلایل مختلفی، آزمون‌ها از گستردگی و تعدد بالایی برخوردارند. از مهم‌ترین این دلایل می‌توان به انواع فرضیه (همبستگی، مقایسه‌ای و علی) و انواع سطح سنجش متغیرها (اسمی، ترتیبی، فاصله‌ای و نسبی) اشاره کرد.

۲. تعمیم: برآورد ویژگی‌های جامعه از روی ویژگی‌های نمونه

حروف یونانی حروف لاتین

به ویژگی‌های عددی جامعه، پارامتر و به ویژگی‌های عددی نمونه، آماره گفته می‌شود.



علائم مربوط به پارامتر و آماره

ویژگی یا شاخص	میانگین	واریانس	انحراف یا معیار	نسبت	همبستگی	تعداد مشاهدات
آماره	$\bar{X}$ یا M	S^2	S	P	r	n
پارامتر	μ (مو)	σ^2	σ (زیگما)	π	ρ (رو)	N

وجه مشترک آمار توصیفی و استنباطی: شاخص نمونه (از شاخص آماری نمونه جهت برآورد پارامتر جامعه استفاده می کنند)

برآوردکننده: ویژگی نمونه‌ای که به صورت تصادفی از جامعه انتخاب شده باشد.

در واقع به کمک آماره ها ابتدا نمونه‌ای انتخاب شده از جامعه توصیف شده و سپس برآورد و فرض آزمایی پیرامون ویژگی های جامعه صورت می گیرد.

حدود واقعی اعداد

طبق قرارداد علم آمار، در جهت سنجش و اندازه گیری دقیق تر، می بایست هر داده کمی را به صورت طیفی پیوسته از حد پایین تا حد بالا در نظر بگیریم (نیم واحد پایین تر تا نیم واحد بالاتر).

۱۲	۱۲/۷۴۰	۱۲/۷۰
۱۱/۵ ۱۲/۵	۱۲/۷۴۵ ۱۲/۷۳۵	۱۲/۷۵ ۱۲/۶۵

جدول توزیع فراوانی

۱. جدول غیر طبقه بندی

نمره خام	فراوانی
X	f
۱۵	۲
۱۲	۳
۵	۱

۲. جدول طبقه بندی (دارای خطای طبقه بندی)^۱

X	f
۵-۰	۲
۱۰-۵	۴
۱۵-۱۰	۳

زمانی که تعداد داده ها و همچنین پراکندگی آن ها زیاد باشد از جدول طبقه بندی استفاده می کنیم.

استفاده از جداول طبقه بندی، باعث ایجاد خطایی در محاسبه شاخص های آماری خواهد شد که به آن خطای طبقه بندی گفته می شود.

(خطای طبقه بندی عبارت است از اختلاف بین شاخص آماری محاسبه شده در استفاده از توزیع فراوانی طبقه بندی شده و توزیع فراوانی طبقه بندی نشده)

نکته: خطای طبقه‌بندی با طول طبقات رابطه مستقیم و با تعداد طبقات رابطه عکس دارد.
به عنوان مثال، در جداول زیر که برای طبقه‌بندی داده‌های ۰ تا ۱۰۰ استفاده شده‌اند، جدول B مقدار خطای بیشتری دارد.

خطای بیشتر B		A	
۲۵-۰		۱۰-۰	
۵۰-۲۵		۲۰-۱۰	
.		.	
.		.	
.		۱۰۰-۹۰	

معرفی شاخص‌های توزیع فراوانی^۱

$$R = \text{Max} - \text{Min} + 1$$

↓
به خاطر حدود واقعی

۶،۹،۱۰،۱۵،۲۶

$$R = ۲۶/۵ - ۵/۵ = ۲۱$$

$$\text{یا } R = ۲۶ - ۶ + ۱ = ۲۱$$

$$I = x_H - x_L + 1$$

↓ ↓ ↓
بزرگ‌ترین داده طبقه کوچک‌ترین داده طبقه به خاطر حدود واقعی

x	f
۸ - ۱۳	۶
۱۴ - ۱۹	

$$i = ۱۳ - ۸ + ۱ = ۶$$

$$K = \frac{R}{I}$$

K تعداد طبقات

به کمک حاصل تقسیم دامنه تغییرات بر طول طبقات می‌توان تعداد طبقات جدول را محاسبه کرد.
به عنوان مثال، اگر دامنه تغییرات برابر ۱۰۰ و طول طبقات مساوی ۵ باشد، تعداد ۲۰ طبقه خواهیم داشت.



نحوه ساختن توزیع فراوانی طبقه‌بندی شده

اولین مرحله: محاسبه دامنه تغییرات

دومین مرحله: تعیین تعداد طبقات

معمولاً تعداد طبقات بین ۱۰ تا ۲۰ طبقه اختیار می‌شود، اما می‌توان از فرمول زیر که به قانون استرژ معروف است نیز استفاده کرد:

$$K = 1 + \sqrt[3]{\log(n)}$$

در فرمول بالا n و K به ترتیب عبارت‌اند از تعداد اعداد، و طبقه‌ها، پایه لگاریتم بر مبنای ۱۰ است. به طور مثال، چنانچه تعداد کل اعداد جمع‌آوری شده مساوی ۶۰ باشد، تعداد طبقه‌ها فرمول مورد بحث، برابر است با:

$$(\log^{10} 60 \text{ در پایه } 10 \text{ مساوی } 1/7782 \text{ است})$$

$$K = 1 + \sqrt[3]{\log n}$$

$$K = 1 + \sqrt[3]{\log 60}$$

$$K = 1 + \sqrt[3]{(1/7782)}$$

$$K = 6/86806 \approx 7$$

$$I = \frac{R}{K}$$

سومین مرحله: تعیین طول طبقات

نمونه‌ای از یک جدول توزیع فراوانی:

X	f (فراوانی مطلق)	Cf (فراوانی تراکمی/تجمعی)	$\frac{f}{n}$ فراوانی نسبی	$\frac{f}{n} \times 100$ فراوانی درصدی	$\frac{Cf}{n}$ فراوانی تراکمی نسبی	$\frac{Cf}{n} \times 100$ فراوانی تراکمی درصدی
۵	۲	۲	۰/۲	۲۰٪	۰/۲	۲۰٪
۸	۱	۳	۰/۱	۱۰٪	۰/۳	۳۰٪
۹	۳	۶	۰/۳	۳۰٪	۰/۶	۶۰٪
۱۳	۱	۷	۰/۱	۱۰٪	۰/۷	۷۰٪
۱۶	۳	10 (=n)	۰/۳	۳۰٪	۱	۱۰۰٪

$$\Sigma f = n$$

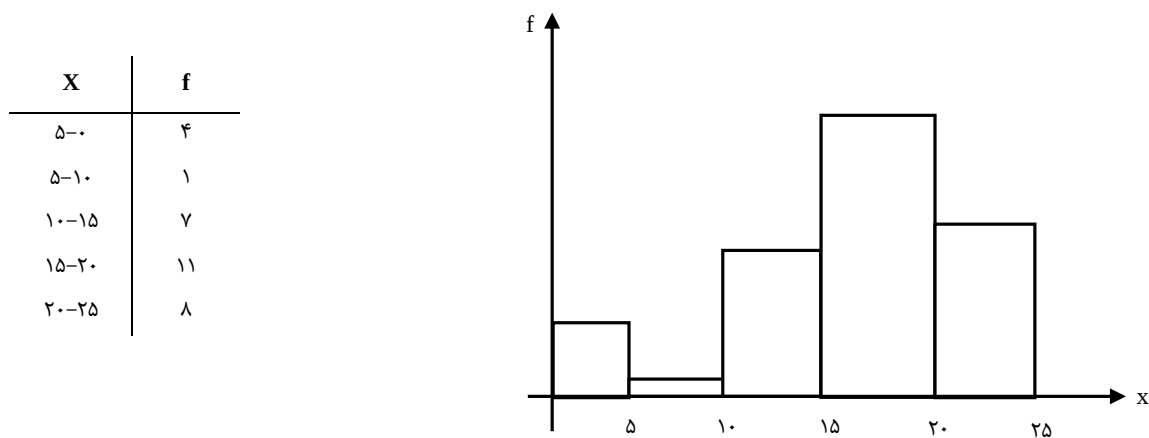
$$\Sigma \frac{f}{n} = 1$$

$$\Sigma \%f = \%100$$

نکته مهم در تنظیم جدول توزیع فراوانی این است که اگر داده‌ها یا طبقات از بزرگ به کوچک مرتب شوند، فراوانی تراکمی یا تجمعی از سمت پایین به بالای جدول محاسبه خواهد شد؛ درواقع فراوانی تراکمی همواره از سمت عدد یا طبقه کوچک شروع می‌شود.

« نمودارها » ← نمایش گرافیکی داده‌ها (ساده‌سازی ارائه‌ی داده‌ها)

۱. نمودار هیستوگرام^۱

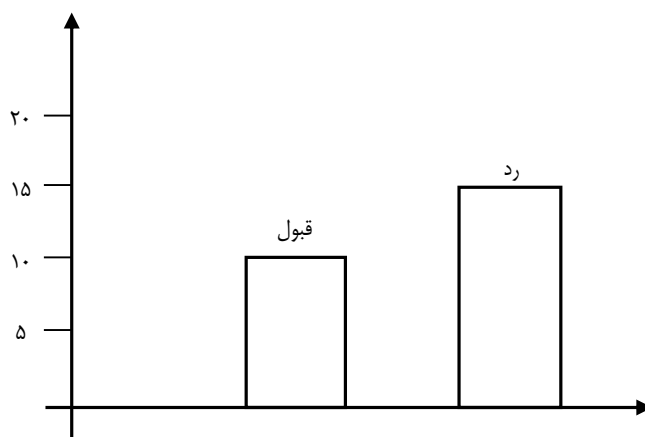


این نمودار در مقیاس‌های فاصله‌ای و نسبی (کمی و پیوسته) استفاده شده و کاربرد مهم آن در آمار، بررسی تقارن توزیع است.

۲. نمودار ستونی^۲ (میله‌ای)

جدول ۲-۱ توزیع فراوانی نمره‌های ۲۵ نفر در یک آزمون فیزیک

درجه‌بندی	فراوانی (f)	درصد (p)
قبول	۱۰	۴۰
رد	۱۵	۶۰



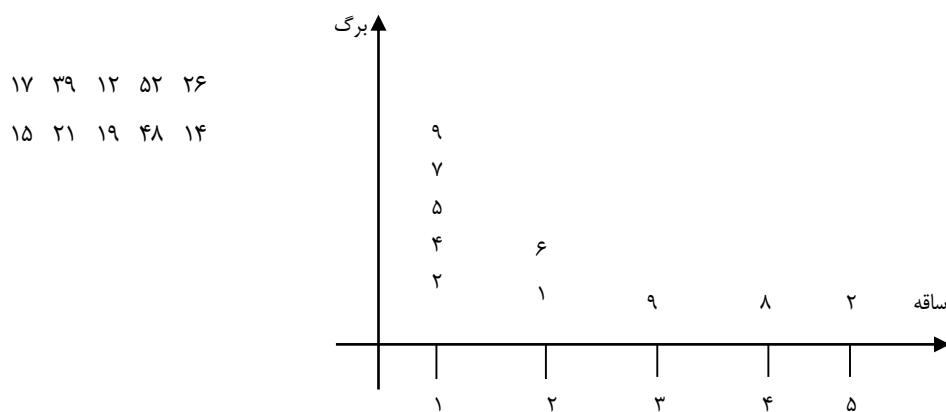
کاربرد نمودار ستونی همانند هیستوگرام بوده (بررسی تقارن توزیع) با این تفاوت که دارای مقیاس اسمی و گسسته است.
(حتی ترتیب طبقات در این نمودار اهمیتی ندارد)

1 - histogram

2 - bargraph

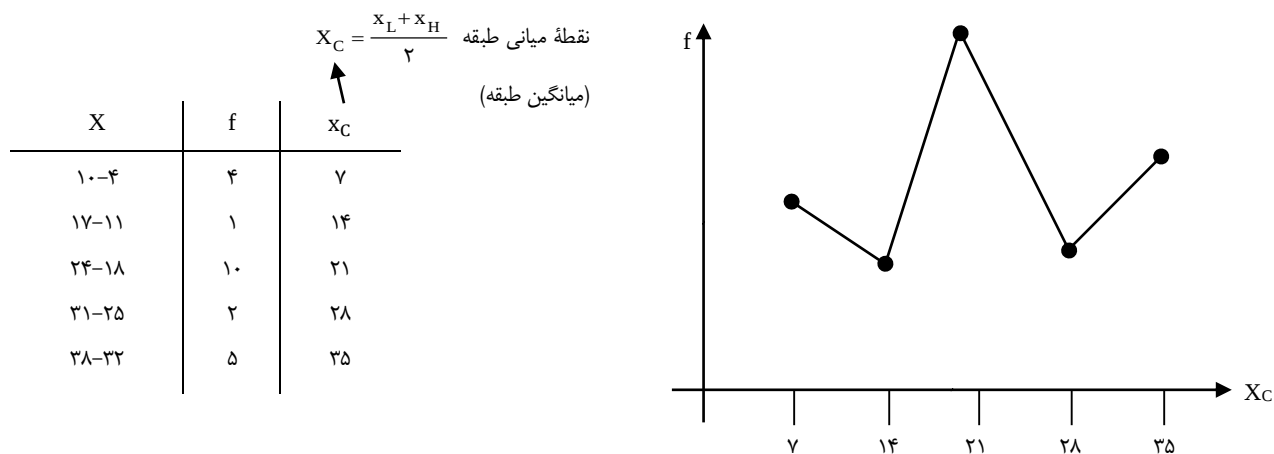


۳. نمودار ساقه و برگ (شاخه و برگ)^۱



از نظر مقیاس و کاربرد، این نمودار کاملاً همانند هیستوگرام است (کاربرد: بررسی تقارن - سطح سنجش: فاصله ای یا نسبی)، با این تفاوت که اگر تعداد داده‌ها کم باشد اطلاعات جزئی‌تری نسبت به هیستوگرام در اختیار قرار می‌دهد.

۴. نمودار چندضلعی^۲



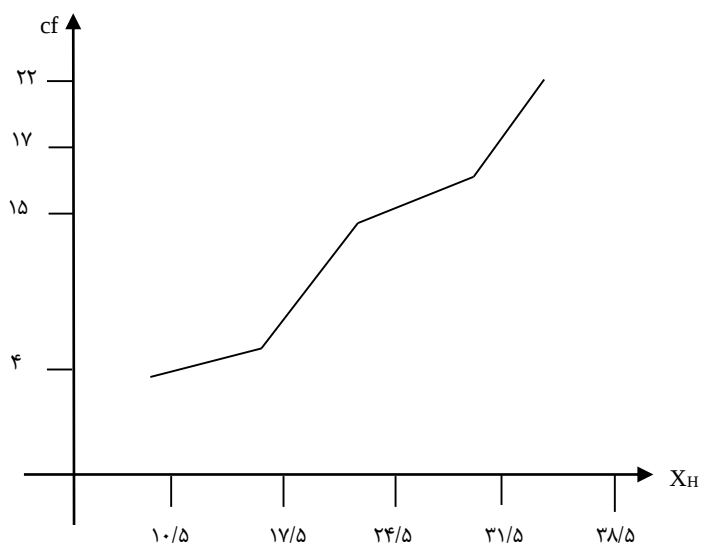
این نمودار نیز برای داده‌های کمی و پیوسته استفاده می‌شود و پرکاربردترین نمودار به شمار می‌رود، زیرا جهت مقایسه دو یا چند توزیع از آن استفاده می‌شود.

1 - Stem and leaf plot

2 - polygon

۵. نمودار چندضلعی تراکمی (اجایو)^۱

X	f	cf	X _H
۱۰-۴	۴	۴	۱۰/۵
۱۷-۱۱	۱	۵	۱۷/۵
۲۴-۱۸	۱۰	۱۵	۲۴/۵
۳۱-۲۵	۲	۱۷	۳۱/۵
۳۸-۳۲	۵	۲۲	۳۸/۵



این نمودار کمی و پیوسته است و جهت نمایش وضعیت یک داده نسبت به سایر داده‌ها (گروه) استفاده می‌شود.

۶. نمودار دایره‌ای^۲

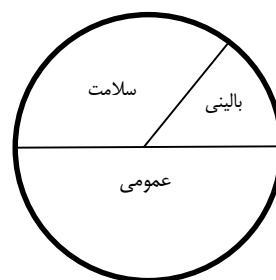
x	f
بالینی	۲۰
عمومی	۶۰
سلامت	۴۰

$$\frac{f}{n} \times ۳۶۰ = \text{اندازه قطاع}$$

$$\frac{۲۰}{۱۲۰} \times ۳۶۰ = ۶۰ \text{ بالینی}$$

$$\frac{۶۰}{۱۲۰} \times ۳۶۰ = ۱۸۰ \text{ عمومی}$$

$$\frac{۴۰}{۱۲۰} \times ۳۶۰ = ۱۲۰ \text{ سلامت}$$



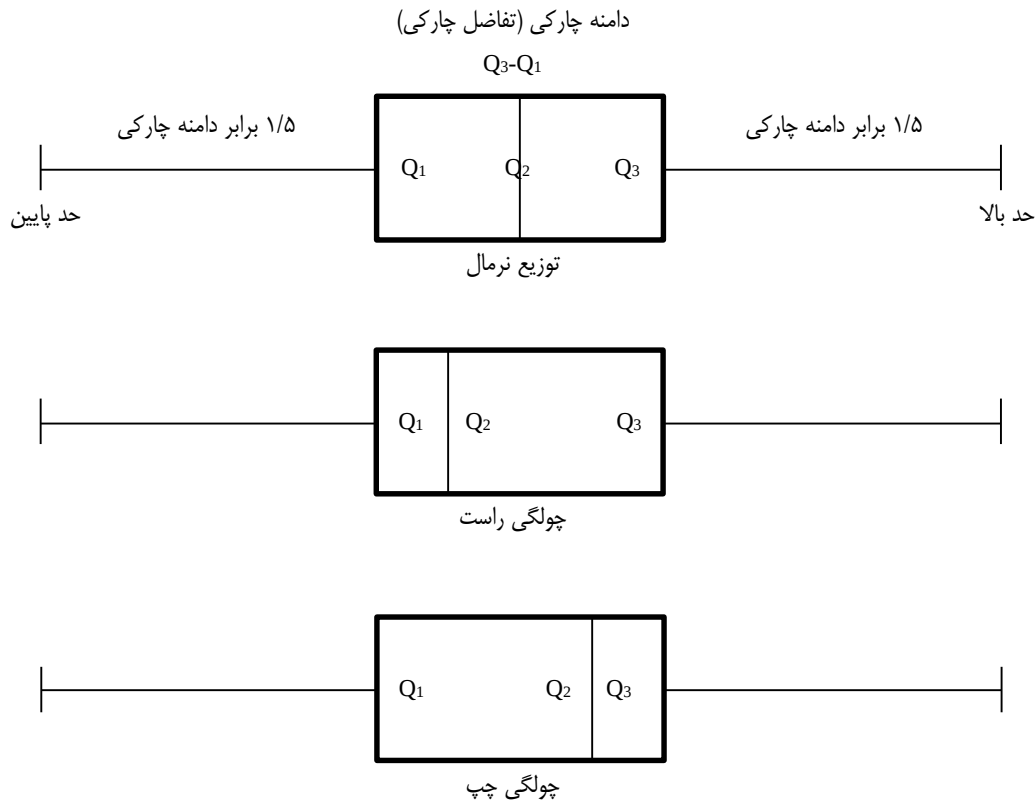
نمودار دایره‌ای دارای مقیاس اسمی - ترتیبی است و جهت نمایش نسبت جزء به کل استفاده می‌شود.

1 - cumulative (Ogive)

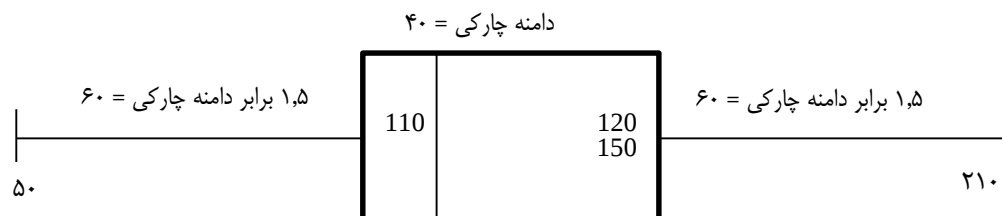
2 - Pie chart

۷. نمودار جعبه‌ای^۱

این نمودار کمی و پیوسته بوده و مهم‌ترین کاربرد آن، یافتن داده پرت است. به‌طوری‌که $1/5$ برابر دامنه چارکی محاسبه و از مقدار چارک اول کاسته و به مقدار چارک سوم افزوده می‌شود. دامنه به‌دست‌آمده، دامنه نرمال بوده و اعداد بیرون از این دامنه، داده پرت به شمار می‌روند:



- به عنوان مثال اگر در توزیع داده های یک پژوهش، مقدار چارک اول برابر ۱۱۰، چارک دوم ۱۲۰ و چارک سوم ۱۵۰ باشد، خواهیم داشت:



بدین ترتیب در این پژوهش، داده های مابین ۵۰ و ۲۱۰ داده پرت نبوده اما چنانچه بیرون از این بازه، عدد یا داده ای داشته باشیم، به عنوان داده پرت شناسایی و معرفی خواهد شد. ضمناً از آنجایی که مقدار چارک دوم به چارک اول نزدیکتر است، می توان نتیجه گرفت که بیشتر تراکم توزیع سمت چپ بوده و با این حساب، توزیع دارای چولگی راست یا + است.

فصل دوم

شاخص‌های مرکزی



۱. مد^۱ یا نما M_0

مد شاخصی مبنی بر تکرار است و به طبقه‌ای اطلاق می‌شود که دارای بیشترین فراوانی باشد. این شاخص ضعیف‌ترین و بی‌ثبات‌ترین شاخص مرکزی به شمار می‌رود، زیرا حتی ویژگی ترتیب برای آن مهم نیست (مقیاس اسمی) و لزوماً در مرکز توزیع قرار نمی‌گیرد. از مد می‌توان در تمامی مقیاس‌ها استفاده کرد و کاربرد آن زمانی است که قصد محاسبه سریع و آسان یک شاخص مرکزی را داشته باشیم.

مثال : $M_0 = 7$: 23 , 12 , 7 , 10 , 20 , 7 , 3 , 15 , 18 , 9 , 7 , 25 , 8 , 13 , 21 , 11 , 24

نکته: زمانی که دو عدد یا دو طبقه متوالی، دارای بیشترین فراوانی و به مقدار یکسان باشند، میانگین این دو عدد یا طبقه، به عنوان مد یا نما معرفی خواهد شد.

طرز محاسبه مد در داده‌های طبقه‌بندی شده (پیوسته):

برای محاسبه مد مراحل زیر را به ترتیب انجام می‌دهیم:

۱- ابتدا طبقه‌ای که بیشترین فراوانی مطلق را دارد انتخاب می‌کنیم.

۲- از فرمول مقابل مد را محاسبه می‌کنیم:

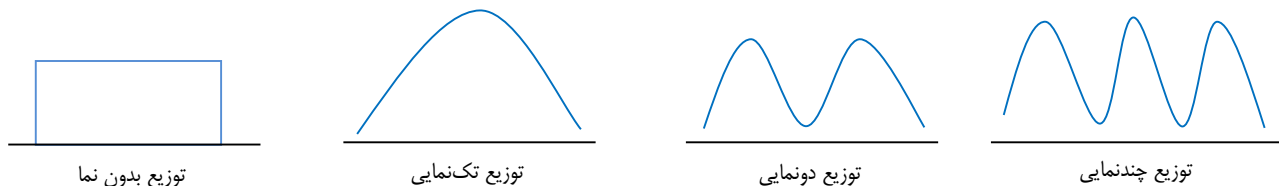
$$MO = L + \frac{f_a}{f_a + f_b} \times i$$

فراوانی مطلق طبقه بالایی نما $\nearrow$

حد پایین طبقه $\nwarrow$

فراوانی مطلق طبقه پایینی نما $\nwarrow$

نکته: توزیع‌های آماری مورد توجه در این کتاب، توزیع‌های تک‌نمایی هستند؛ اما لازم به ذکر است که توزیع فراوانی داده‌ها می‌تواند دو نمایی یا چند نمایی باشد:



۲. میانه^۲ M_d

میانه نقطه وسط یا پنجاه درصدی داده‌ها است.

این شاخص مبتنی بر ترتیب و دارای مقیاس ترتیبی است؛ پس در مقیاس‌هایی می‌توان از آن استفاده کرد که ویژگی ترتیب را داشته باشند؛ یعنی مقیاس‌های ترتیبی، فاصله‌ای و نسبی.

میانه برخلاف میانگین شاخصی کیفی است که تحت تأثیر ارزش عددی یا حجم داده‌ها قرار نمی‌گیرد. لذا میانه به داده‌های پرت کوچک یا بزرگ حساس نیست؛ بنابراین می‌توان گفت اگر توزیع داده‌ها چولگی یا داده پرت داشته باشد، بهترین و مناسب‌ترین شاخص مرکزی میانه است.

$$\sum |x - M_d| \leq \sum |x - A|$$

همواره و در هر توزیعی مجموع قدر مطلق انحراف داده‌ها از میانه حداقل است.

1 - mode

2 - median

«محاسبه میانه»

۱. وقتی تعداد داده‌ها فرد باشد.

۵، ۹، ۱۰، ۱۵، ۱۸

۲. وقتی تعداد داده‌ها زوج باشد.

۵، ۸، ۱۰، ۱۴، ۱۷، ۲۰

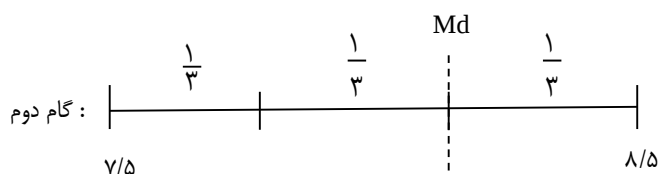
از این دو میانگین می‌گیریم (۱۲)

۳. زمانی که در میانه توزیع داده تکراری باشیم.

Md

۵، ۷، ۸، ۸، ۱۲، ۱۵، ۲۰

مکان میانه: گام اول $\frac{n}{2} = \frac{8}{2} = 4$



$$Md = 7/5 + \frac{2}{3} = 8/16$$

تعداد تکرارهای قبل از میانه + حدپایین عدد تکراری = **Md** : فرمول محاسبه میانه
تعداد کل تکرارها

x	f
۱۵	۱
۱۹	۳
۲۶	۴
۳۲	۲

Md

۱۵ ۱۹ ۱۹ ۱۹ ۲۶ ۲۶ ۲۶ ۳۲ ۳۲

مکان میانه: $\frac{10}{2} = 5$

$$Md = 25/5 + \frac{1}{4} = 25/75$$

Md

۱۲، ۱۵، ۱۸، ۱۸، ۱۸، ۱۸، ۲۰

مکان میانه: $\frac{n}{2} = \frac{7}{2} = 3/5$

$$Md = 17/5 + \frac{1/5}{4} = 17/87$$



۴. زمانی که تعداد داده‌ها زیاد باشد یا جدول طبقه‌بندی شده داشته باشیم:

$$Md = L + \frac{\frac{n}{2} - cf}{f} (i)$$

L = پایین حد واقعی طبقه‌ای که میانه در آن قرار دارد

cf = فراوانی تراکمی طبقه قبل از طبقه میانه

f = فراوانی مطلق طبقه‌ای که میانه در آن قرار دارد

i = فاصله طبقات

n = تعداد نمره‌ها

به‌طور مثال:

جدول توزیع فراوانی نمره‌های یک آزمون

طبقات	f	Cf
۴۸-۵۰	۲	۵۰
۴۵-۴۷	۳	۴۸
۴۲-۴۴	۴	۴۵
۳۹-۴۱	۶	۴۱
۳۶-۳۸	۸	۳۵
۳۳-۳۵	۸	۲۷
۳۰-۳۲	۷	۱۹
۲۷-۲۹	۶	۱۲
۲۴-۲۶	۳	۶
۲۱-۲۳	۲	۳
۱۸-۲۰	۱	۱

محاسبه میانه برای توزیع فراوانی بالا، شامل پنج مرحله به شرح زیر است:

$$۵۰ \div ۲ = ۲۵$$

۱. تقسیم تعداد کل نمره‌ها بر ۲

۲. مراجعه به ستون فراوانی تراکمی و پیدا کردن اولین طبقه‌ای که فراوانی تراکمی آن مساوی یا بزرگ‌تر از $\frac{N}{2}$ یعنی ۲۵ است. با مراجعه به جدول اولین طبقه‌ای

است که فراوانی تراکمی آن بزرگ‌تر از ۲۵ است (برای پیدا کردن این طبقه از پایین‌ترین طبقه شروع می‌کنیم).

$$L = ۳۲/۵$$

۳. حد پایین واقعی این طبقه یا L را تعیین می‌کنیم.

۴. مقادیر Cf و f را از جدول توزیع فراوانی استخراج می‌کنیم (به ترتیب ۱۹ و ۸).

۵. مقادیر تعیین شده به شرح فوق را در فرمول زیر جایگزین می‌کنیم.

$$Md = ۳۲/۵ + \frac{\frac{۵۰}{۲} - ۱۹}{۸} (۳) = ۳۴/۷۵$$

۳. میانگین^۱ [حسابی] M یا $\bar{x}$

ویژگی‌های میانگین

۱. میانگین دارای مقیاس فاصله‌ای است و در مقیاس‌های فاصله‌ای و نسبی می‌توان از آن استفاده کرد.
۲. اگر توزیع داده‌ها نرمال باشد بهترین و مناسب‌ترین شاخص مرکزی میانگین است.
۳. همواره و در هر توزیعی مجموع انحراف داده‌ها از میانگین برابر صفر است، به همین دلیل است که میانگین را مرکز ثقل یا نقطه تعادل توزیع می‌نامند.

$$X = 2 \quad 3 \quad 5 \quad 14 \quad \bar{X} = 6 \quad X - \bar{X} : \quad -4 \quad -3 \quad -1 \quad 8$$

$$\begin{array}{c} \text{---} \\ -8 \quad +8 \end{array}$$

همواره و در هر توزیعی: $\sum (x - \bar{x}) = 0$

۴. همواره و در هر توزیعی مجموع مجذور انحراف داده‌ها از میانگین حداقل است:

$$\sum (x - \bar{x})^2 \leq \sum (x - a)^2$$

هر عدد دلخواهی: a

محاسبه میانگین:

برای داده‌های خام: $\bar{x} = \frac{\sum x}{n}$

در جدول غیر طبقاتی: $\bar{x} = \frac{\sum f x}{n}$

X	f
۱۵	۳
۱۶	۴
۱۰	۲

برای جدول طبقاتی: $\bar{x} = \frac{\sum f x_c}{n}$

X	f	X_c
۱۵-۲۳	۳	۱۹
۲۴-۳۲	۱	۲۸
۳۳-۴۱	۴	۳۷



مثال: کلاس ۱۰ نفره‌ای دارای میانگین ۱۵ و کلاس ۴۰ نفره‌ای دارای میانگین ۱۰

$$\bar{X} = 15 \quad \bar{X} = 10$$

$$n = 10 \quad n = 40$$

$$\bar{X} = \frac{n_1 \bar{X}_1 + n_2 \bar{X}_2 + \dots}{n_1 + n_2 + \dots}$$

میانگین وزنی (میانگین کل، میانگین مرکب، میانگین میانگین‌ها)

$$\text{میانگین کل} = \frac{10 \times 40 + 15 \times 10}{50} = \frac{550}{50} = 11$$

فرمول تجربی پیرسون (رابطه میان شاخص‌های مرکزی):

$$m_1 = 3m_d - 2\bar{X}$$

هرگاه بخواهیم به کمک دو شاخص مرکزی، شاخص سوم را به دست بیاوریم، می‌توان با تقریب نسبتاً مناسبی از فرمول تجربی پیرسون استفاده کرد.

به عنوان مثال: در توزیعی با میانگین ۲۰ و نمای ۵، میانه چند است؟

$$5 = 3Md - 2 \times 20 : Md = 15$$

میانگین هندسی^۱

برای محاسبه میانگین درصدها، نسبت‌ها یا نرخ‌ها استفاده می‌شود.

$$G = \sqrt[n]{x_1 x_2 \dots x_n}$$

مثال: میانگین هندسی اعداد ۸ ۲۷ ۱۲۵

$$G = \sqrt[3]{8 \times 27 \times 125} = 30$$

میانگین هارمونیک^۲ یا همساز

برای محاسبه میانگین داده‌هایی با واحد ترکیبی (مانند $\frac{m}{s}$) استفاده می‌شود.

$$H \text{ یا } HM = \left(\frac{n}{\frac{1}{x_1} + \frac{1}{x_2} + \dots + \frac{1}{x_n}} \right)$$

نکته: در صورت محاسبه میانگین‌ها برای یک توزیع فراوانی، میانگین حسابی از نظر عددی بزرگ‌تر از میانگین هندسی و هارمونیک است:

$$\bar{X} > G > HM$$

فرمول محاسبه میانگین همساز از طریق میانگین‌های هندسی و حسابی.

$$HM = \frac{G^2}{\bar{X}}$$

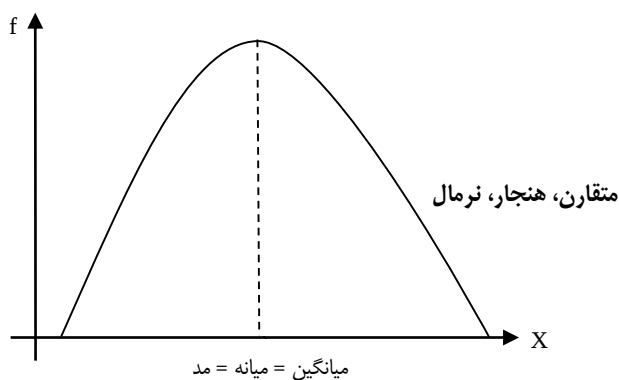
1 - Geometric mean

2 - Harmonic mean

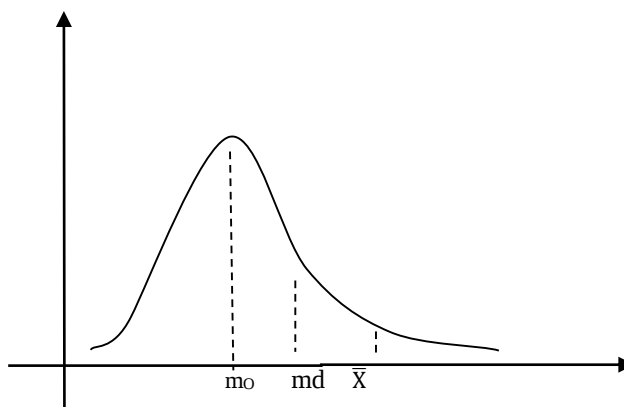
چولگی^۱ (کجی)

شاخص چولگی تحت تأثیر داده پرت است و جهت بررسی تقارن توزیع استفاده می‌شود.

X	f
۰	۲
۳	۴
۵	۷
۸	۹
۱۱	۱۵
۱۴	۸
۱۷	۷
۲۰	۱



X	f
۰	۴
۲	۱۷
۳	۱۲
۵	۸
۷	۵
۸	۵
۱۹	۱
۲۰	۱



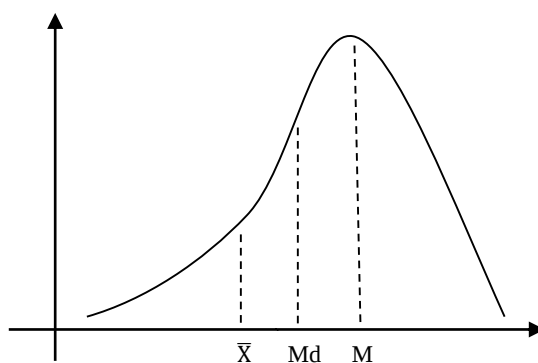
چولگی راست یا مثبت:

باینکه اکثریت نمره کم دارند، به خاطر چند تا نمره بزرگ میانگین بزرگ‌تر از نما و میانه شده است.

داده پرت بزرگ ← میانگین را بزرگ می‌کند ← چولگی راست

مثال برای چولگی راست: توزیع درآمد در تهران

X	f
۰	۲
۲	۱
۱۳	۴
۱۴	۶
۱۵	۶
۱۶	۱۰
۱۷	۱۵
۱۸	۶
۲۰	۴



چولگی چپ یا منفی

داده پرت کوچک ← میانگین را کوچک می‌کند ← چولگی چپ

فصل سوم

شاخص‌های پراکندگی



شاخص‌های پراکندگی^۱ (فاصله‌ای)

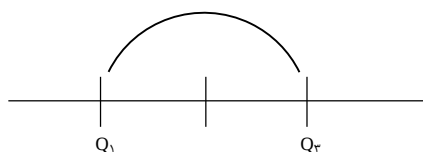


۱. دامنه تغییرات^۲ $R = \max - \min + ۱$

پراکنده‌تر: $R = ۱۸$: کلاس A و ۱۹ و و ۳ و ۲
همگن‌تر: $R = ۵$: کلاس B و ۱۸ و و ۱۴ و ۱۴

نه تنها دامنه تغییرات بلکه تمامی شاخص‌های پراکندگی به دنبال فاصله میان داده‌ها و دارای مقیاس فاصله‌ای هستند. دامنه تغییرات ضعیف‌ترین و بی‌ثبات‌ترین شاخص پراکندگی به شمار می‌رود، در محاسبه آن فقط از دو عدد ابتدایی و انتهایی استفاده شده و با تغییر یکی از این اعداد مقدار R نیز شدیداً تغییر می‌کند؛ لذا در میان شاخص‌های پراکندگی، دامنه تغییرات بیشترین حساسیت را به داده پرت دارد. پس زمانی در آمار از این مقیاس استفاده می‌شود که بدون داشتن داده پرت قصد محاسبه سریع و آسان یک شاخص پراکندگی را داشته باشیم.

۲. انحراف چارکی^۳



$$Q = \frac{Q_3 - Q_1}{۲} \quad (\text{چارک متوسط})$$

انحراف چارکی نشان می‌دهد داده‌ها به‌طور متوسط چقدر از میانه یا مرکز توزیع انحراف دارند. درواقع انحراف چارکی میزان پراکندگی را در ۵۰٪ وسط توزیع بررسی کرده لذا به داده‌های پرت کوچک یا بزرگ حساس نیست. پس می‌توان گفت زمانی که توزیع داده‌ها چولگی قابل ملاحظه‌ای داشته باشد، بهترین و مناسب‌ترین شاخص پراکندگی انحراف چارکی خواهد بود.

۳. انحراف متوسط^۴ $AD = \frac{\sum |x - \bar{x}|}{n}$ یا MD

انحراف متوسط عبارت است از متوسط قدر مطلق انحراف داده‌ها از میانگین. این شاخص نشان می‌دهد که داده‌ها به‌طور متوسط چقدر از میانگین انحراف دارند. انحراف متوسط یک ایراد اساسی دارد و آن اینکه به کمک قراردادی به نام قدر مطلق علائم جبری را از بین برده و در عملیات ریاضی و آماری نمی‌توان از آن استفاده کرد.

مثال: در یک توزیع ۳۰ تایی چنان چه مجموع انحرافات منفی از میانگین برابر ۶۰ شده باشد، مقدار انحراف متوسط چند است؟
از آنجایی که میانگین مرکز ثقل بوده و مقدار انحرافات منفی و مثبت از آن به یک اندازه است، خواهیم داشت:

$$AD = \frac{60 + 60}{30} = \frac{120}{30} = 4$$

1 - variability

2 - range

3 - Quartile deviation

4 - Mean deviation

$$S^2 = \frac{\sum (x - \bar{x})^2}{n} \quad \text{۴. واریانس}^1$$

واریانس عبارت است از متوسط مجذور انحراف داده‌ها از میانگین.

$$SS \text{ مجموع مجذورات } = \sum (x - \bar{x})^2 \longrightarrow S^2 \quad \text{دلیل نام‌گذاری واریانس:}$$

واریانس پرکاربردترین شاخص پراکندگی است، اما یک ایراد اساسی دارد و آن اینکه به دلیل مجذور کردن انحراف‌ها، اختلاف واحد اندازه‌گیری ایجاد می‌کند. مثلاً $m^2 \longrightarrow m$

برای رفع این ایراد است که از واریانس، جذر یا رادیکال گرفته شده و شاخصی به نام انحراف استاندارد ساخته شده است.

تفاوت واریانس‌های نمونه و جامعه: در آمار استنباطی (قضیه حد مرکزی) به طور تجربی اثبات شده است که واریانس نمونه تصادفی از واریانس کل جامعه کوچک‌تر است (داشتن اریب یا سوگیری). به همین دلیل برای اینکه در محاسبه واریانس نمونه این اریب را از بین ببریم لازم است در مخرج فرمول از $n - 1$ استفاده کنیم:

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

تست: واریانس برآورد (تعمیم) داده‌های مقابل چقدر است؟

واریانس نمونه

$$X = 2 \quad 3 \quad 6 \quad 8 \quad 16$$

نمونه است و می‌خواهیم تعمیم دهیم.

$$\bar{X} = \frac{35}{5} = 7$$

$(x - \bar{x})$	$(x - \bar{x})^2$	$\sum (x - \bar{x})^2 = 124$
$2 - 7 = -5$	25	$\left. \begin{array}{l} -10 \\ +10 \end{array} \right\}$
$3 - 7 = -4$	16	
$6 - 7 = -1$	1	
$8 - 7 = 1$	1	
$16 - 7 = 9$	81	

$$S^2 = \frac{\sum (x - \bar{x})^2}{n - 1} = \frac{124}{5 - 1} = 31 \quad \text{نمونه (برآورد، تعمیم)}$$

$$S^2 = \frac{\sum (x - \bar{x})^2}{n} = \frac{124}{5} = 24.8 \quad \text{جامعه (واریانس واقعی)}$$

واریانس

$$S^2 = \frac{\sum x^2 - \frac{(\sum x)^2}{n}}{n - 1} = \frac{369 - \frac{1225}{5}}{4} = \frac{124}{4} = 31 \quad \text{محاسبه واریانس به کمک اعداد خام:}$$



محاسبه واریانس مرکب

همان طور که در فصل قبلی به محاسبه میانگین کل (میانگین مرکب) برای دو یا چند میانگین (دو یا چند گروه) اشاره شد، زمانی که محقق علاقمند باشد واریانس مرکب دو یا چند گروه را محاسبه کند، می تواند از فرمول زیر که به فرمول مک نیمار^۱ معروف است، استفاده کند:

$$S_T^2 = \frac{n_A(\bar{X}_A^2 + S_A^2) + n_B(\bar{X}_B^2 + S_B^2)}{n_A + n_B} - \bar{X}_T^2$$

S_T^2 : واریانس مرکب

$\bar{X}_T^2$: میانگین مرکب

نکته: در صورتی که بیشتر از دو گروه داشته باشیم، عنا صر گروه های دیگر به صورت و مخرج فرمول فوق اضافه می شوند. ضمناً اگر قصد محاسبه انحراف استاندارد مرکب را داشته باشیم، از تمام فرمول مک نیمار، جذر یا رادیکال گرفته و بدین ترتیب انحراف استاندارد مرکب محاسبه خواهد شد.

۵. انحراف استاندارد^۲ (انحراف معیار)

$$S = \sqrt{S^2}$$

انحراف معیار برابر است با جذر یا رادیکال واریانس یعنی جذر یا ریشه دوم میانگین مجذور انحراف داده ها از میانگین. زمانی که توزیع داده ها نرمال باشد، یعنی داده پرت یا چولگی شدیدی نداشته باشیم، بهترین و مناسب ترین شاخص پراکندگی انحراف استاندارد خواهد بود.

سوال: در توزیعی با کجی $SK = -0.73$ چنانچه بهترین شاخص گرایش به مرکز، میانگین انتخاب شده باشد. بهترین شاخص پراکندگی کدام خواهد بود؟
پاسخ: S: زمانی که میانگین بهترین شاخص مرکزی است، توزیع نرمال بوده و بهترین شاخص پراکندگی انحراف معیار خواهد بود:

	پراکندگی	مرکزی
نرمال	S	$\bar{X}$
غیرنرمال	Q	Md

مقدار SK یا چولگی نیز در این سوال کمتر از یک بوده و نشان از نرمال بودن توزیع دارد.

نکته: بین انحراف معیار (S)، انحراف متوسط (MD) و انحراف چارکی همواره رابطه روبه رو برقرار است: $Q < MD < S$

(مقدار انحراف چارکی از سایر شاخص های پراکندگی کمتر بوده و مقدار انحراف معیار (S) نیز همواره کمی بیشتر از انحراف متوسط (MD) است) در واقع، در یک توزیع طبیعی، تقریباً ۵۰ درصد داده ها در فاصله $\pm 1Q$ ، ۵۸ درصد در فاصله $\pm 1MD$ و ۶۸ درصد داده ها در فاصله های $\pm 1S$ از میانگین قرار دارند.

$$S = \sqrt{S^2 - \frac{i^2}{12}} \quad \text{تصحیح شپرد}^3$$

در جدول طبقاتی، زمانی که میانگین را محاسبه می کنیم به دلیل استفاده از x_c مقدار خطا بسیار ناچیز شده و از نظر آماری اهمیتی ندارد اما در محاسبه انحراف استاندارد به دلیل مجذور شدن انحراف ها، این خطا افزایش می یابد و لازم است به کمک تصحیح شپرد مقدار آن تعدیل شود. ضمناً این تصحیح را زمانی استفاده می کنیم که تعداد طبقات کمتر از ۱۲ باشد: $k < 12$

۱ - Mc Nemar

2 - Standard deviation

3 - Sheppard's correction

سؤال: در جدول طبقاتی با تعداد ۹ طبقه و طول طبقات ۶ چنانچه واریانس ۲۸ به دست آمده باشد مقدار انحراف استاندارد چقدر است؟

$$S = \sqrt{28 - \frac{36}{12}} = \sqrt{25} = 5$$

$$i = 6$$

$$S^2 = 28$$

$$9 < 12$$

اگر $K > 12$ بود دیگر نیازی به تصحیح نداشت: $S = \sqrt{28}$

ضریب تغییرات^۱ (ضریب پراکندگی) $V = \frac{S}{\bar{X}}$ یا $c.v$

برای مقایسه پراکندگی یک ویژگی در گروه‌های مختلف، می‌توان از S یا S^2 استفاده کرد؛ اما برای مقایسه پراکندگی دو یا چند ویژگی مختلف با واحدهای متفاوت می‌بایست از شاخصی به نام ضریب تغییرات که مستقل از واحد اندازه‌گیری است استفاده کرد.

برای مثال:

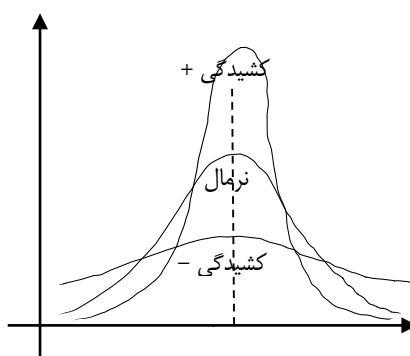
دانش‌آموزان یک کلاس	وزن	$S = 20$ $\bar{X} = 50$	$V = \frac{s}{\bar{x}} = \frac{20}{50} = 40\%$ یا 0.4 پراکنده‌تر
	قد	$S = 30$ $\bar{X} = 100$	$V = \frac{s}{\bar{x}} = \frac{30}{100} = 30\%$ یا 0.3 همگن‌تر

کشیدگی^۲

تحت تأثیر انحراف معیار

↓
کشیدگی Kurtosis

↓
بررسی تراکم



نمودار کشیدگی

تحت تأثیر داده پرت

↓
چولگی Skewness

↓
بررسی تقارن

کشیدگی، تحت تأثیر مقدار انحراف استاندارد بوده و جهت بررسی تراکم توزیع از آن استفاده می‌شود؛ به طوری که مقدار کشیدگی، با میزان تراکم توزیع رابطه مستقیم اما با میزان پراکندگی (S) رابطه عکس دارد. به طوری که در کشیدگی منفی یا پخ (هموار)، مقدار انحراف معیار زیاد بوده (تراکم کم) در حالی که در کشیدگی مثبت یا قله ای، مقدار انحراف معیار کم اما تراکم زیادی در توزیع داده ها وجود دارد.

1 - Coefficient of variation

2 - kurtosis



نکته: توزیع نرمال به توزیعی گفته می‌شود که نه چولگی شدید و نه کشیدگی شدید داشته باشد.

گشتاور^۱:

$$\frac{\sum(x-\bar{x})}{n} = 0 : m_1$$

گشتاور مرتبه اول

$$\frac{\sum(x-\bar{x})^2}{n} = \frac{n-1}{n} S^2 : m_2$$

گشتاور مرتبه دوم

$$\frac{\sum(x-\bar{x})^3}{n} : m_3$$

گشتاور مرتبه سوم

و به همین ترتیب تا مرتبه n ام.

$$Ku = \frac{m_3}{S^3} - 3 : \text{کشیدگی}$$

$$Sk = \frac{m_3}{S^3} : \text{چولگی}$$

فرمول های چولگی و کشیدگی بر حسب گشتاور:

همان طور که ملاحظه می‌شود، چولگی (کجی)، تحت تأثیر گشتاور مرتبه سوم بوده و کشیدگی، از گشتاور مرتبه چهارم تأثیر می‌پذیرد.

تأثیر اعمال ریاضی بر شاخص‌های مرکزی و پراکندگی:

به‌طور کلی شاخص‌های مرکزی به تمامی اعمال ریاضی حساس بوده درحالی‌که شاخص‌های پراکندگی فقط به ضرب و تقسیم حساس هستند. چنانچه تمامی داده‌های یک توزیع را در عدد ثابتی ضرب یا تقسیم کنیم، شاخص‌های پراکندگی نیز در همان عدد ضرب یا تقسیم خواهند شد به‌جز واریانس که در مجذور آن عدد ضرب یا تقسیم می‌شود.

تست: اگر به بزرگ‌ترین داده یک توزیع ۱۰ نمره اضافه شود، کدام یک از شاخص‌های زیر بدون تغییر باقی می‌ماند؟

- (۱) میانگین (۲) انحراف استاندارد (۳*) میانه (۴) انحراف متوسط

سؤال: تمامی داده‌های یک توزیع دو برابر شده، سپس به اضافه ۵ شده و میانگین و واریانس حاصل ۲۵ و ۳۶ به دست آمده است. میانگین و انحراف استاندارد اولیه چه قدر بوده است؟

$$2X + 5$$

$$2\bar{X} + 5 = 25 \quad 2\bar{X} = 20 \quad \bar{X} = 10$$

$$2^2 s^2 = 36 \rightarrow s^2 = 9 : S = 3$$

سؤال: اگر مقدار انحراف معیار X_n و ... و X_2 و X_1 برابر ۴ باشد، مقدار واریانس $10 + 5X_n$ و ... و $10 + 5X_1$ چقدر است؟

$$\begin{array}{c} S = 4 \\ \downarrow \times 5 \\ S_x = 4 \times 5 = 20 \quad S = 20 : S^2 = 400 \end{array}$$

مثال: فراوانی تمامی داده‌ها را در عدد ۳ ضرب کرده‌ایم، کدام تغییر درباره میانگین اتفاق می‌افتد؟

- (۱) افزایش (۲) کاهش (۳) تغییر نمی‌کند (۴) تغییر می‌کند

نکته مهم: در صورتی که فراوانی تمام داده‌ها به شکل یکسانی تغییر کند هیچ شاخص آماری تغییر نمی‌کند.

سؤال: اگر داده‌های یک توزیع، همگی با عدد ثابت مثبتی جمع شوند، مقدار ضریب تغییرات چه تغییری خواهد کرد؟

پاسخ: کاهش می‌یابد.

از آنجا که تنها شاخص‌های مرکزی به جمع و تفریق حساس هستند، پس در صورت اضافه شدن یک عدد به تمامی داده‌ها، در فرمول ضریب تغییرات مقدار مخرج (میانگین) افزایش پیدا می‌کند، در حالی که انحراف معیار بدون تغییر باقی می‌ماند، به همین دلیل مقدار ضریب تغییرات کاهش پیدا خواهد کرد.