

علم آمار

در آمار توصیفی مفهوم اولیه جامعه آماری می‌باشد. جامعه آماری دسته‌ای از اشیاء و یا افراد را گویند که لااقل در یک خاصیت مشترک باشند و ما بخواهیم موضوع یا موضوع‌هایی را در آن‌ها بررسی کنیم. جامعه آماری یا محدود است یا نامحدود و جامعه آماری محدود جوامعی هستند که اعضای آن قابل شمارش است و جامعه آماری نامحدود اعضای آن قابل شمارش نیستند. هر چه خاصیت مشترک بیشتر شود جامعه آماری محدودتر می‌شود و صفت مشترک را صفت مشخصه می‌گویند. اندازه جامعه آماری را با N نشان می‌دهند.

انواع صفت متغیر:

الف - متغیر کمی: که قابل اندازه‌گیری و قابل شمارش باشد مانند سن، قد، وزن و ...

ب - متغیر کیفی: که قابل اندازه‌گیری و شمارش نیست مانند شماره ملی، جنسیت، گروه خونی و ...

انواع متغیر کمی:

الف - متغیر کمی پیوسته: متغیری است که مجموعه مقادیر خود را از مجموعه اعداد حقیقی اختیار می‌کند. مانند وزن

ب - متغیر کمی ناپیوسته (گسسته): مجموعه مقادیر خود را از مجموعه اعداد طبیعی یا معادل هم ارز آن اختیار می‌کند.

انواع متغیر کیفی:

الف - متغیر کیفی اسمی: این متغیرها قابل شمارش نیستند مانند شماره شناسنامه که کیفی اسمی می‌باشد که فقط یک مشخصه برای فرد می‌باشد.

ب - متغیر کیفی ترتیبی: این متغیرها همگی کیفی اسمی‌اند ولی ویژگی‌های خاصی دارند که آن‌ها را در یک ترتیب قابل قبول قرار می‌دهد. مانند ایزوها، درجه نظامی، مقاطع تحصیلی و متغیری می‌تواند کیفی ترتیبی باشد که در قالب یک استاندارد باشد.

مراحل انجام یک کار آماری

۱- مشاهده یا جمع‌آوری اطلاعات

۲- پردازش یا مرتب کردن داده‌ها و تعیین شاخص‌های آماری

۳- تحلیل آماری

مشاهده:

۱- **سرشماری:** سنجش تک تک افراد جامعه آماری است. به دلایل مشکلاتی که در سرشماری وجود دارد همواره اقدام به این عمل نمی‌شود. این مشکلات شامل: هزینه بر بودن، زمان بر بودن و ... در برخی موارد موجب از بین رفتن جامعه آماری می‌شود.

۲- **نمونه‌گیری:** تعدادی از افراد جامعه را انتخاب و مورد بررسی قرار می‌دهیم. دو نکته مهم در نمونه‌گیری را می‌توان چنین گفت الف: به چه روشی انتخاب شود؟ ب: چند نمونه لازم است؟

جامعه آماری

جامعه آماری عبارتست از مجموعه تمام افراد، گروه‌ها، اشیاء و یا رویدادهایی که دارای یک یا چند ویژگی مشترک باشند. تعداد اعضای جامعه را حجم یا اندازه جامعه می‌نامند و با حرف بزرگ N نشان می‌دهند.

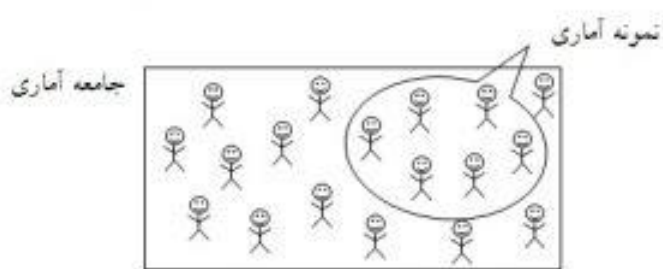
مثال: جامعه کارکنان شاغل در بانک تجارت شهر تهران

نمونه آماری

نمونه آماری گروه کوچکتري از جامعه است که طبق ضابطه‌ای معین برای مشاهده و تجزیه و تحلیل انتخاب می‌شود و باید معرف جامعه باشد. نتایج نمونه ای را که معرف جامعه نباشد نمی‌توان به جامعه تعمیم داد. تعداد اعضای نمونه را با حرف کوچک n نشان می‌دهند.

مثال: کارکنان شاغل در بانک تجارت منطقه ۵ شهر تهران

شکل ۱: جامعه آماری و نمونه آماری



شاخص های پراکندگی

شاخص های پراکندگی، میزان پراکندگی (تغییرات) مقادیر هر متغیر را در اطراف میانگین نشان می دهند. زمانی که در شاخصهای گرایش به مرکز، داده -ها را در یک اندازه واحد خلاصه کنیم، طبیعتاً بخشی از جزئیات و اطلاعات حذف خواهد شد. از این رو باید به دنبال شاخص -هایی برای اندازه -گیری تفاوت موردها در یک متغیر باشیم تا از نحوه و میزان پراکندگی و تغییر داده -ها در اطراف میانگین (یا میانه و مد) مطلع شویم.

به عنوان مثال نمرات درس های دو دانشجو را با یکدیگر مقایسه میکنیم. معدل (میانگین نمرات) هر دو دانشجو یکسان و برابر با ۱۵ است. نمرات دانشجوی الف بدین صورت است: ۱۷، ۱۶، ۱۵، ۱۴ و ۱۳. و نمرات دانشجوی ب بدین صورت است: ۲۰، ۱۸، ۱۶، ۱۱ و ۱۰ است.

نگاهی به نمرات دو دانشجو نشان می دهد که دامنه نمرات دانشجوی الف محدودتر است و داده -ها از نمره میانگین (۱۵) فاصله کمتری دارند. نتیجه می -گیریم در حالتی که میانگین، یک شاخص گرایش به مرکز است در هر دو دانشجو برابر با ۱۵ است اما پراکندگی نمرات در دانشجوی ب بیشتر از دانشجوی الف است.

شاخص های گرایش به مرکز اطلاعاتی از نحوه پراکندگی داده -ها و نحوه توزیع -شان به ما نمی -دهند و جهت اطلاع از نحوه پراکندگی داده -ها، باید از شاخص -های پراکندگی استفاده کنیم. از مهم -ترین شاخص -های پراکندگی می -توان به انحراف استاندارد، واریانس، ضریب تغییرات و دامنه تغییرات اشاره کرد.

دامنه تغییرات:

ساده ترین شاخص پراکندگی است و برابر با اختلاف بزرگترین و کوچکترین مشاهده است. یکی از عیوب دامنه تغییرات این است که تنها به بزرگترین و کوچکترین مشاهده توجه کرده و چگونگی قرار گرفتن مشاهدات بین این دو مقدار را در نظر نمی گیرد.

مثال: دامنه مشاهدات زیر را به دست آورید

6, 4, 5, 6, 5, 5, 4

R: 6-4=2

انحراف معیار:

انحراف معیار (Standard deviation) از دو واژه تشکیل یافته است. جزء اول یعنی انحراف به میزان دوری هر عضو یک مجموعه داده از مقدار میانگین گفته می شود. واژه معیار نیز به معنی استاندارد بودن این مقدار است. هر چه انحراف معیار مجموعه ای از داده ها عدد پایین تری باشد، نشانه آن است که داده ها به میانگین نزدیک هستند و پراکندگی اندکی دارند. در صورتی که انحراف معیار عدد بزرگی باشد، نشان می دهد که پراکندگی داده ها زیاد است. پس انحراف معیار، عددی برای نشان دادن میزان پراکندگی اعضای یک مجموعه از داده ها است. انحراف معیار نیز میزان پراکندگی داده ها از نقطه میانگین را نشان می دهد و از این رو مقیاسی دوبعدی برای برآورد توزیع داده ها در اختیار ما قرار می دهد.

رای مثال اگر یک معلم هستید، احتمالاً برایتان مهم است که بدانید دانش‌آموزان شما در امتحانی که اخیراً گرفته‌اید چه عملکردی داشته‌اند. اگر ۲۰ یا ۳۰ دانش‌آموز داشته باشید با نگاه کردن به تک‌تک نمرات شاید نتوانید برآورد صحیحی از عملکرد کل کلاس به دست آورید، ولی مسلماً در صورتی که میانگین نمره‌های همه دانش‌آموزان را محاسبه کنید، می‌توانید بدانید که وضعیت کل کلاس چگونه بوده است. برای مثال اگر میانگین نمره‌های کلاس برابر با ۱۲٫۵ باشد و میانگین محاسبه شده برای امتحان قبلی ۱۴ بوده باشد، نشان دهنده افت نمرات است و نیاز به چاره‌جویی وجود دارد.

شما به عنوان یک معلم باید با کدام دانش‌آموزان بیشتر کار کنید؟ مسلماً برای دانش‌آموزانی که عملکرد بهتری دارند نیاز چندانی به تلاش بیشتر وجود ندارد، اما به دانش‌آموزانی که عملکرد ضعیف‌تری دارند می‌بایست توجه ویژه‌ای صورت بگیرد. اما چگونه می‌توان فهمید که کدام دانش‌آموزان عملکرد بالاتر دارند، متوسط هستند یا عملکرد ضعیف‌تری دارند؟ پاسخ به این سؤال از طریق محاسبه انحراف معیار است. انحراف معیار مقیاسی به دست می‌دهد که با استفاده از آن می‌توانیم بدانیم میانگین اختلاف عملکرد دانش‌آموزان از نقطه میانگین کلاسی چقدر است.

برای مثال فرض کنید در کلاس شما انحراف معیار برابر با ۲٫۵ باشد. اگر توزیع نمرات دانش‌آموزان به صورت یک توزیع نرمال باشد (که در اغلب موارد در مورد چنین اندازه‌گیری‌هایی از توزیع نرمال پیروی می‌کند)، این عدد نشان می‌دهد که نمرات بیش از دو سوم یا ۶۸٫۲٪ از دانش‌آموزان شما در محدوده ۲٫۵ ± ۱۲٫۵ قرار دارد. این عدد طبق تعریف انحراف معیار به دست می‌آید. یک سوم دیگر از دانش‌آموزان یا نمراتی بالاتر از ۱۵ کسب کرده‌اند که طبعاً نیاز چندانی به تلاش بیشتر شما ندارند و یا نمراتی زیر ۱۰ کسب کرده‌اند که مسلماً نیازمند توجه ویژه هستند. بدین ترتیب با محاسبه انحراف معیار نمره‌های کلاسی می‌توانید دانش‌آموزان را به سه دسته ضعیف (کمتر از ۱۰)، متوسط (۱۰ تا ۱۵) و قوی (بالاتر از ۱۵) تقسیم‌بندی کنید.

واریانس:

واریانس یک مفهوم آماری است که نشان می‌دهد مقادیر یک متغیر تصادفی چگونه حول میانگین آن توزیع شده‌اند. برای محاسبه واریانس باید ابتدا میانگین ساده اعداد را پیدا کنید. در ادامه برای هر عدد باید مقدار میانگین را از آن کم کرده و نتیجه بدست آمده را به توان دو برسانید و در واقع مربع اختلاف را بدست آورید. سپس باید میانگین مربع اختلافاتی که بدست آورده‌اید را محاسبه نمایید تا واریانس بدست آید.

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

مثال: فرض کنید متصدی یک محل نگهداری از سگ‌ها می‌خواهد قد سگ‌های موجود را به منظور خاصی اندازه‌گیری کند. نتایج این اندازه‌گیری قد (از شانه) به شرح زیر است:

۳۰۰، ۴۳۰، ۱۷۰، ۴۷۰ و ۶۰۰ میلی‌متر

اینک می‌خواهیم میانگین، واریانس و انحراف معیار این داده‌ها را پیدا کنیم. گام اول یافتن میانگین است:

$$\text{میانگین} = \frac{600 + 470 + 170 + 430 + 300}{5} = \frac{1970}{5} = 394$$

اکنون اختلاف قد هر کدام از سگ‌ها را از مقدار میانگین حساب می‌کنیم و برای محاسبه واریانس، اختلاف تک‌تک داده‌ها را به توان دو رسانده و سپس میانگین می‌گیریم.

$$\sigma^2 = \frac{206^2 + 76^2 + (-224)^2 + 36^2 + (-94)^2}{5} = \frac{42436 + 5776 + 50176 + 1296 + 8836}{5} \\ = \frac{108520}{5} = 21704$$

پس، واریانس برابر است با: ۲۱۷۰۴

و انحراف معیار همان جذر واریانس است، پس:

$$\sigma = \sqrt{21704} = 147.32 \dots = 147 \text{ (نزدیک ترین داده)}$$

و اما نکته خوب در مورد انحراف معیار، سودمند بودن آن است. اکنون می‌توانیم بفهمیم قد کدام سگ‌ها در محدوده یک انحراف معیار میانگین (۱۴۷ میلی‌متر) قرار دارد. پس با استفاده از انحراف معیار، ما یک راه “استاندارد” برای یافتن محدوده مقادیر نرمال، مقادیر بیش از نرمال و مقادیر کمتر از نرمال در دست داریم.

متغیرهای استاندارد:

فرض کنید که X متغیر تصادفی دارای توزیع نرمال استاندارد باشد. احتمال اینکه X بین a و b قرار داشته باشد از رابطه زیر بدست می‌آید:

$$\int_a^b \frac{1}{\sqrt{2\pi}\sigma} e^{-(x-\mu)^2/(2\sigma^2)} dx$$

در صورتیکه در رابطه فوق به جای میانگین عدد صفر و به جای واریانس عدد یک قرار داده شود، یک حالت خاص از توزیع نرمال به نام توزیع نرمال استاندارد بدست می‌آید. توزیع نرمال با پارامترهای میانگین برابر با صفر و واریانس برابر با یک، توزیع نرمال استاندارد نامیده می‌شود. متغیر تصادفی دارای چنین توزیعی نیز متغیر تصادفی نرمال استاندارد گفته شده که معمولاً با Z یا برخی اوقات با N نشان داده می‌شود.

ضریب تغییرات:

ضریب تغییرات (پراکندگی) شاخصی است که برای اندازه‌گیری توزیع پراکندگی داده‌های آماری به کار می‌رود. از ضریب تغییرات برای مقایسه پراکندگی دو یا چند صفت (متغیر) استفاده می‌کنند و کاربرد اصلی آن مقایسه متغیرهایی است که واحدهای سنجش متفاوتی دارند. ضریب تغییرات میزان پراکندگی را به ازای یک واحد از میانگین بیان می‌کند.

کند. این شاخص تنها برای سطح سنجش نسبی کاربرد دارد. ضریب تغییرات را با علامت CV نشان می‌دهند. معمولاً ضریب تغییرات را در عدد ۱۰۰ ضرب می‌کنند تا عدد نهایی برحسب درصد به دست بیاید. ضریب تغییرات از تقسیم انحراف استاندارد بر میانگین به دست می‌آید و فرمول آن بدین صورت است:

$$\text{ضریب تغییرات} = \frac{\text{انحراف استاندارد}}{\text{میانگین}}$$

مثال: زمانی که بخواهیم میزان پراکندگی دو متغیر سن (برحسب سال) و قد (سانتی‌متر) گروهی از افراد را مقایسه کنیم از این شاخص استفاده می‌کنیم. فرض کنیم اطلاعات مربوط به سن و قد ۱۰۰ نفر را جمع‌آوری کرده‌ایم و این نتایج به دست آمده است: میانگین قد افراد برابر با ۱۷۰ سانتی‌متر با انحراف استاندارد ۱۷ سانتی‌متر و میانگین سن افراد برابر با ۵۰ سال با انحراف استاندارد ۵ سال است. ضریب تغییرات دو متغیر را به دست می‌آوریم:

$$\begin{aligned} \text{ضریب تغییرات قد} &= \frac{17}{170} = 0.10 \quad \text{--->} 0.10 \times 100 = 10 \\ \text{ضریب تغییرات سن افراد} &= \frac{5}{50} = 0.10 \quad \text{--->} 0.10 \times 100 = 10 \end{aligned}$$

احتمالات

در آمار احتمال اندازه امکان وقوع حادثه است.

مفهوم اولیه در احتمال حادثه و پیش آمد می باشد. هر کاری را که راجع به وقوع یا عدم وقوع آن صحبت کنیم به آن حادثه یا پیش آمد می گویند. حوادث را با حروف بزرگ نمایش می دهند. مثلاً احتمال حادثه A برابر است با: $p(A)$ به دنبال هر آزمایشی یک حادثه اتفاق می افتد.

آزمایش: برقرار کردن مجموعه شرطهای معین به دفعات دلخواه و تکراری است.

شروط معین: این شروط مثلاً در پرتاب یک مکعب باید دارای ویژگیهای: در زمین سفت و سخت باشد، مولکول سازنده آن یکسان باشد و همه رویه های آن یک اندازه باشد.

در آمار ۳ حادثه وجود دارد:

۱- یقین (احتمال حوادث یقین ۱ می باشد)

۲- غیرممکن (احتمال حوادث غیرممکن صفر می باشد)

۳- تصادفی (از قبل نمی توان گفت چه چیزی می آید ولی می توانیم پیش بینی کنیم و احتمال حوادث تصادفی بین ۱ و صفر اتفاق می افتد)

طبق تعریف کلاسیک احتمال: احتمال هر حادثه ای برابر است با تعداد حالات مطلوب به تعداد حالات ممکن.

$$P(A) = \frac{m_A}{n} \leftarrow \frac{\text{تعداد حالات مطلوب}}{\text{تعداد حالات ممکن}} = \text{احتمال}$$

هیچ وقت m_A نمی تواند بزرگ تر از n باشد و احتمال همیشه یک عدد مثبت بین ۱ و صفر می باشد و احتمال هر حادثه ای شبیه فراوانی نسبی می باشد.

خواص (قوانین) احتمال

۱- احتمال هر حادثه ای یک عدد مثبت است. $P(A) \geq 0$

۲- احتمال هر حادثه ای همواره بین ۱ و صفر است. $0 \leq P(A) \leq 1$

۳- مجموع احتمال حوادث همواره یک است. $\sum p_i = 1$

تمام قوانینی را که در مجموعه ها داریم در حوادث نیز صدق می کند.

۴- احتمال عکس هر حادثه برابر است با یک منهای احتمال آن حادثه $P(\bar{A}) = 1 - P(A)$ و منظور از $\bar{A}$ یعنی غیر از A (نه A)

۵- اگر دو حادثه ناسازگار باشند:

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

$$P(A+B) = P(A) + P(B)$$

$$P(A+B+C) = P(A) + P(B) + P(C)$$

حوادث ناسازگار: حوادثی هستند که عضو مشترک ندارند.

۶- اگر دو حادثه سازگار باشند:

$$P(A+B) = P(A) + P(B) - P(A.B)$$

$$P(A+B+C) = P(A) + P(B) + P(C) - P(AB) - P(AC) - P(BC) + P(ABC)$$

$$P(A \times B) = P(A) \times P(B|A) \text{ احتمال شرطی}$$

$$P(ABC) = P(A) \times P(B|A) \times P(C|AB)$$

مثال: در کیسه‌ای ۴ مهره سیاه و ۶ مهره قرمز وجود دارد. دو مهره به تصادف از آن کیسه خارج می‌کنیم مطلوب است احتمال اینکه: $(B=4)$, $(R=6)$

- الف) هر دو سیاه باشد. $P(B_1 B_2)$
 ب) هر دو قرمز باشد. $P(R_1 R_2)$
 ج) اولی سیاه و دومی قرمز باشد. $P(BR)$
 د) یکی سیاه و یکی قرمز باشد. $P(BR+RB)$

پاسخ:

الف) هر دو سیاه باشد.

$$P(B_1 B_2) = P(B_1) \times P(B_2 | B_1) \rightarrow \frac{4}{10} \times \frac{3}{9} = \frac{12}{90}$$

ب) هر دو قرمز باشد.

$$P(R_1 R_2) = P(R_1) \times P(R_2 | R_1) \rightarrow \frac{6}{10} \times \frac{5}{9} = \frac{30}{90}$$

ج) اولی سیاه و دومی قرمز باشد.

$$P(BR) = P(B) + P(R|B) \rightarrow \frac{4}{10} \times \frac{6}{9} = \frac{24}{90}$$

د) یکی سیاه و یکی قرمز باشد.

$$P(BR+RB) = P(BR) + P(RB) \rightarrow P(B) \times P(R|B) + P(R) \times P(B|R) \rightarrow \frac{4}{10} \times \frac{6}{9} + \frac{6}{10} \times \frac{4}{9} = \frac{48}{90}$$

مثال: با توجه به جدول زیر از این جامعه ۵۰ عضوی ۲ نفر به تصادف انتخاب می‌شوند. احتمال اینکه این دو نفر **مرد یا متأهل باشند**، چیست؟

جنسیت	متأهل	مجرد	جمع
زن	12	8	20
مرد	18	12	30
جمع	30	20	50

(جدول سازگار)

$$P(\text{مرد} + \text{متأهل}) = P(\text{مرد}) + P(\text{متأهل}) - P(\text{مرد متأهل}) = \frac{30}{50} + \frac{30}{50} - \frac{18}{50} = \frac{42}{50}$$

$$P(AB) = P(A) \times P(B|A) \rightarrow P(B|A) = \frac{P(BA)}{P(A)}$$

$$P(B|A) = \frac{P(BA)}{P(A)} \leftrightarrow \frac{\text{احتمال حاصل ضرب دو حادثه}}{\text{احتمال حادثه دوم}} = \text{احتمال شرطی}$$

اگر علاوه بر شروط معین که در آزمایش وجود دارد و یک حادثه اتفاق بی افتد یک شرط دیگر را برقرار کنیم به آن **احتمال شرطی** می‌گویند.

فضای نمونه:

مجموعه نتایج ممکن برای یک آزمایش تصادفی را «فضای نمونه» (Sample Space) می‌نامند. معمولاً فضای نمونه یک آزمایش تصادفی را با حرف Ω نمایش می‌دهند. برای مثال در پرتاب یک سکه به منظور مشاهده شیر (H) یا خط (T)، فضای نمونه برابر با $\Omega = \{H, T\}$ خواهد بود.

همچنین اگر دو سکه مستقل از یکدیگر پرتاب شوند (یا یک سکه دوبار پرتاب شود)، فضای نمونه به صورت $\Omega = \{HH, HT, TH, TT\}$ درخواهد آمد. منظور از TH مشاهده خط در پرتاب اول و شیر در پرتاب دوم است در حالیکه HT به معنی مشاهده شیر و سپس مشاهده خط در پرتاب دوم است.

همچنین در آزمون زبان چینی که دارای ۱۰ پرسش چهار گزینه‌ای است، اگر تعداد پاسخ‌های صحیح مورد نظر باشد، فضای نمونه برابر است با $\Omega = \{0, 1, 2, 3, 4, 5, 6, 7, 8, 9, 10\}$. قابل ذکر است که چنین آزمایش‌هایی دارای فضای نمونه متناهی هستند.

اگر سکه‌ای را تا مشاهده اولین شیر پرتاب کنیم و تعداد پرتاب‌های صورت گرفته تا رسیدن به چنین وضعیتی مورد نظر باشد، فضای نمونه برای این آزمایش تصادفی به صورت $\Omega = \{1, 2, \dots\}$ است، زیرا ممکن است این اتفاق در اولین پرتاب، دومین پرتاب یا ... اتفاق بیافتد. این فضای نمونه را نامتناهی شمارش پذیر می‌نامند.

در ورزش تیر و کمان اگر فاصله محل اصابت تا مرکز تخته ملاک باشد، فضای نمونه برای این آزمایش تصادفی به صورت $\Omega = [0, r]$ است، که منظور از r ، شعاع تخته هدف در این ورزش است. چنین فضای نمونه‌ای با توجه به اینکه یک فاصله از اعداد حقیقی است، نامتناهی و ناشمارا محسوب می‌شود.

پیشامد:

در حالتی که فضای نمونه متناهی باشد، می‌توان گفت هر زیر مجموعه‌ای از فضای نمونه یک پیشامد تلقی می‌شود. در این حالت با توجه به فضای نمونه‌ای مربوط به پرتاب دو سکه، می‌توان $A = \{TT, TH\}$ را پیشامد مشاهده خط در اولین پرتاب در نظر گرفت. مشخص است که $A \subset \Omega$.

از آنجایی که A یک پیشامد است، انتظار داریم مکمل آن یعنی A' نیز یک پیشامد باشد. از طرفی خود Ω نیز یک پیشامد خواهد بود زیرا $\Omega \subset \Omega$ ، پس انتظار داریم که $\emptyset$ نیز یک پیشامد باشد. به Ω **پیشامد حتمی** و به $\emptyset$ **پیشامد محال** می‌گویند.

در نتیجه اگر F را مجموعه همه زیر مجموعه‌های Ω در نظر بگیریم، به آن «فضای پیشامد» می‌گوییم. مجموعه F باید در شرایط زیر صدق کند:

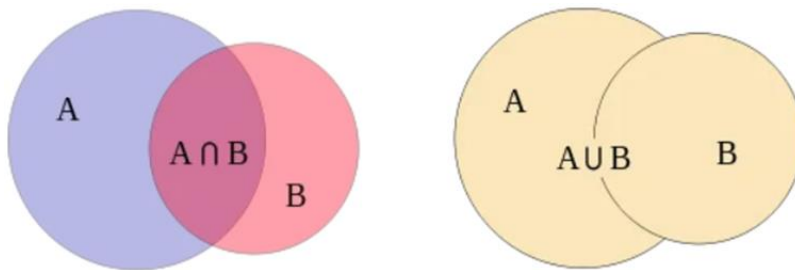
۱- مجموعه تهی ($\emptyset$) و فضای نمونه (Ω) باید در F باشند. به این معنی که $\emptyset \in F$ و $\Omega \in F$

۲- اگر A در F باشد باید مکمل آن یعنی A' نیز در F باشد.

۳- اگر $A_1, A_2, \dots$ دنباله‌ای از پیشامدهای جدا از هم باشند، اجتماع آن‌ها نیز یک پیشامد است، یعنی $\bigcup_{i=1}^{\infty} A_i \in F$

مجموعه‌ای که براساس زیر مجموعه‌های یک فضای نمونه نامتناهی ایجاد شود و در همه شرط‌های اول تا سوم صدق کند، «سیگما-میدان (σ -Field)» یا «سیگما-جبر (σ -Algebra)» (گفته می‌شود. پس می‌توان فضای پیشامد را سیگما-میدان حاصل از فضای نمونه در نظر گرفت.

نکته $A \cap B$: را پیشامد رخداد هر دو پیشامد A و B می‌نامند. همچنین پیشامد $A \cup B$ نیز به معنی رخداد پیشامد A یا B (یا هر دو) خواهد بود. پیشامد مکمل A یعنی A' را پیشامد عدم رخداد A در نظر می‌گیرند.



پیشامدهای ناسازگار

اگر A و B دو پیشامد باشند، آن‌ها را ناسازگار گویند، اگر اشتراکشان برابر مجموعه $\emptyset$ باشد. یعنی:

$$A \cap B = \emptyset$$

به این ترتیب فقط یا پیشامد A رخ داده یا B و نه هر دو.

دنباله پیشامدهای $A_1, A_2, \dots$ را دو به دو ناسازگار گویند، اگر همزمان فقط یکی از آن‌ها رخ دهد. یعنی:

$$A_i \cap A_j = \emptyset, i, j = 1, 2, \dots$$

دنباله پیشامدهای یکنوا

اگر دنباله پیشامدهای $A_1, A_2, \dots$ به صورت صعودی باشند، آن را یکنوا می‌نامیم. یعنی رابطه بین A_i ها به صورت زیر باشد.

$$A_1 \subset A_2 \subset A_3 \subset \dots$$

البته ممکن است این رابطه به صورت نزولی نیز گفته شود، ولی در هر دو حالت دنباله پیشامدها یکنوا هستند، زیرا می‌توان با ساخت پیشامدهای جدید از روی پیشامدهای نزولی، آن‌ها را به صورت صعودی درآورد.

اصول احتمال و تابع احتمال

با توجه به تعریف فضای نمونه و فضای پیشامد، امکان تعریف تابع احتمال بوجود می‌آید.

تعریف تابع احتمال: هر تابعی یا نگاشتی از فضای پیشامد به مجموعه اعداد حقیقی اگر در سه اصل زیر صدق کند یک تابع احتمال است.

- **اصل اول:** مقدار احتمال برای هر پیشامد، نامنفی است. بنابراین اگر A یک پیشامد باشد (یعنی متعلق به فضای پیشامد F باشد)

$$P(A) \geq 0$$

- **اصل دوم:** مقدار احتمال برای فضای نمونه برابر است با ۱.

$$P(\Omega)=1$$

- **اصل سوم:** احتمال اجتماع هر دنباله نامتناهی از پیشامدهای دو به دو ناسازگار، برابر با مجموع احتمال پیشامدهای دنباله است.

$$P(\bigcup_{i=1}^{\infty} A_i) = \sum_{i=1}^{\infty} P(A_i)$$

این اصول توسط «آندری کولموگوروف» (Andrey Kolmogorov) «ریاضی و آماردان روسی در سال ۱۹۳۳، به عنوان اصول تابع احتمال مطرح شد. این کار باعث شد که نظریه احتمال از پشتوانه آنالیز ریاضی و نظریه اندازه برخوردار شود. براساس این سه اصل قضیه‌های زیادی برای تابع احتمال اثبات شده است. در ادامه به معرفی چند قضیه مهم برای تابع احتمال می‌پردازیم:

۱- مقدار احتمال برای پیشامد $\emptyset$ برابر با صفر است.

$$P(\emptyset)=0$$

۲- احتمال اجتماع دو پیشامد ناسازگار برابر با مجموع احتمال آنها است.

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

۳- مقدار احتمال برای مکمل یک پیشامد A برابر با تفاضل مقدار احتمال پیشامد A از ۱ است.

$$P(A') = 1 - P(A)$$

۴- اگر پیشامد A زیر مجموعه پیشامد B باشد، مقدار احتمال نیز برای پیشامد A از B کمتر است.

$$A \subset B \rightarrow P(A) \leq P(B)$$

۵- اگر C پیشامدی باشد که از اجتماع دو پیشامد A و B ایجاد شده باشد، احتمال پیشامد C برابر است با مجموع پیشامدهای A و B منهای پیشامد $A \cap B$.

$$P(C) = P(A \cup B) = P(A) + P(B) - P(A \cap B)$$

پیوستگی تابع احتمال

برطبق اصل سوم می‌توان به بررسی پیوستگی تابع احتمال پرداخت. **قضیه پیوستگی تابع احتمال:** دنباله‌ای نامتناهی و یکنوا از پیشامدهای جدا از هم مثل A_n را در نظر بگیرید، آنگاه حد احتمال این پیشامدها برابر با احتمال حد پیشامدها است. یعنی:

$$\lim_{n \rightarrow \infty} P(A_n) = P(\lim_{n \rightarrow \infty} A_n)$$

مثال: فرض کنید فردی مسیر محل کار تا منزلش را به دو شیوه می‌تواند طی کند. فضای نمونه و فضای پیشامد برای او به چه صورت است؟

پاسخ: اگر شیوه اول را با A و شیوه دوم را با B نشان دهیم، فضای نمونه به صورت $\Omega = \{A, B\}$ نوشته خواهد شد. همچنین فضای پیشامد نیز برابر است با $F = \{A, B, \emptyset, \Omega\}$ ، زیرا واضح است که $A \cup B = \Omega$ و $A \cap B = \emptyset$ همینطور مشخص است که $B' = A$ و $A' = B$.

احتمال اینکه روش اول را انتخاب کند برابر با 0.4 است. احتمال اینکه روش دوم را برای رفتن به منزل انتخاب کند چقدر است؟

$$P(B) = P(A') = 1 - P(A) = 1 - 0.4 = 0.6$$

احتمال اینکه او از این دو روش استفاده نکند چقدر خواهد بود؟ طبق قضیه شماره ۵ می توان نوشت

$$P(A \cap B) = P(A \cup B) - P(A) - P(B)$$

حال با به کار بردن این تساوی برای پیشامدهای مکمل داریم:

$$P(A' \cap B') = P(A' \cup B') - P(A') - P(B') = p(\Omega) - (1 - 0.4) - (1 - 0.6) = 1 - 0.6 - 0.4 = 0$$

مثال: فرض کنید A ، B و C سه پیشامد باشند و فضای نمونه از این سه پیشامد تشکیل شده. حال به دنبال پاسخ برای پرسش های زیر هستیم.

- پیشامد اینکه هر سه با هم اتفاق بیافتند، چیست؟

با توجه به نکته ای که در قسمت پیشامد گفته شد، $A \cap B \cap C$ پیشامد مورد نظر است.

- احتمال اینکه هیچکدام اتفاق نیافتند کدام است؟ با توجه به اینکه سه پیشامد، فضای نمونه را می سازد، $\Omega = A \cup B \cup C$ که به معنی رخداد حداقل یکی از پیشامدها است. در نتیجه احتمال رخداد هیچکدام از پیشامدها برابر است با

$$P(A' \cap B' \cap C') = P(\Omega') = 1 - P(\Omega) = 1 - 1 = 0$$

- احتمال اینکه فقط دو پیشامد رخ دهد چگونه محاسبه می شود؟

$$P(A \cap B \cap C') + P(A \cap C \cap B') + P(B \cap C \cap A')$$

مثال: دو پیشامد A و B در فضای پیشامد F هستند و $P(A) = a$ و $P(B) = b$ حداکثر مقدار برای $P(A \cup B)$ چقدر خواهد بود؟

طبقه قضیه ۵ می دانیم $P(C) = P(A \cup B) = P(A) + P(B) - P(A \cap B)$ ولی از مقدار $P(A \cap B)$ اطلاعی نداریم. مشخص است که مقدار احتمال نامنفی است یعنی برای هر پیشامد مانند A و B داریم $P(A) \geq 0$ و $P(B) \geq 0$ از طرفی چون $A \cap B \subset A$ و $A \cap B \subset B$ ، نتیجه می گیریم که:

$$P(A \cap B) \leq P(A), P(A \cap B) \leq P(B) \rightarrow P(A \cap B) \leq \min(P(A), P(B)) = \min(a, b)$$

روش‌های آماری

برای کشف یک واقعیت، فرضیه ای تعریف می شود (hypothesis) و آزمایشی انجام گیرد (experiment)، اگر آزمایش صحت فرضیه را ثابت نمود، فرضیه مورد قبول واقع می شود و در غیر اینصورت رد می شود. اگر صحت فرضیه چندین بار به اثبات رسد فرض به صورت تئوری یا نظریه (theory) در می آید.

آزمایش: آزمایش به کلیه عملیاتی گفته می شود که جهت رد یا قبول فرضیه ای در روی تعداد محدودی از افراد یک جامعه آماری انجام می شود. آزمایش در حالت کلی دو نوع است:

۱- آزمایش مطلق

۲- آزمایش مقایسه ای

در آزمایش مطلق فقط در روی یک ماده خاص بررسی می شود مثلاً محقق در مورد یک واکنش شیمیایی تحقیق می کند، در صورتیکه در آزمایش مقایسه ای بیش از یک مورد تحت بررسی و مقایسه قرار می گیرد. در طرح های آزمایشی غالباً از آزمایشهای مقایسه ای استفاده می شود. به عنوان مثال برای مقایسه عملکرد ۴ واریته گندم (A, B, C, D) و هر یک در سه تکرار می توان آنها را در مزرعه بصورت زیر کاشت و عملکردها را مورد مقایسه قرار داد.

تکرار اول	A	C	B	D
تکرار دوم	D	A	C	B
تکرار سوم	A	B	D	C

طرح های آزمایشی Experimental Designs :

طرحهای آزمایشی الگوهای ابداع شده‌ای هستند که برای انجام آزمایشات مقایسه ای مورد استفاده قرار می گیرند. آشنایی درست با این طرح ها و آشنایی با قوانین آماری دو شرط اساسی برای موفقیت در اجرای یک آزمایش است.

تیمار Treatment:

تیمار یا رفتار که گاهی واریانت نیز نامیده می شود به هر یک از عواملی که در یک آزمایش مورد مقایسه قرار می گیرند، اطلاق می شود. مثلاً در مقایسه چند واریته گندم از لحاظ عملکرد هر یک از واریته های گندم یک تیمار است یا به فرض اثر سه نوع کود ازته، فسفره و پتاسه را در روی عملکرد یک واریته گندم در نظر بگیرند، هر نوع کود در اینجا یک تیمار است یا همینطور زمان های مختلف آبیاری، فواصل کشت، رژیم غذایی دام و

تیمار را معمولاً با حروف بزرگ نظیر A, B, C و ... نشان می دهند.

ماده آزمایشی Experimental material:

مقایسه تیمارها به کمک وسیله یا موجودی انجام می گیرد. موجود یا وسیله موردنظر را ماده آزمایشی می گویند، مثلاً برای مقایسه چند وارینه گندم بایستی آنها را در مزرعه ای کاشت، خاک یا زمین در اینجا ماده آزمایشی است، یا برای مقایسه اثر چند نوع رژیم غذایی در یک نوع نژاد دام، ماده آزمایشی دام خواهد بود.

تکرار اول	A	C	B	D
تکرار دوم	D	A	C	B
تکرار سوم	A	B	D	C

واحد آزمایشی (کرت یا پلات) Experimental unit (plot):

قسمتی از ماده آزمایشی است که تیمار در یک تکرار تحت بررسی و آزمایش قرار دارد. مثلاً در آزمایش های مربوط به کود، کرتی که در آن یک کود و در یک تکرار آزمایش می شود، واحد آزمایشی است. یا در آزمایشهای دامپروری یک راس دام که تحت تاثیر یک رژیم غذایی خاصی قرار گرفته باشد یک واحد آزمایشی محسوب می گردد. به واحد آزمایشی در آزمایشهای مزرعه ای معمولاً کرت گفته می شود.

بلوک Block:

مجموعه ای از واحدهای آزمایشی هستند که در تحت شرایط مشابهی تشکیل شده اند. لازم است که ماده آزمایشی که یک بلوک در آن قرار دارد از حداکثر یکنواختی برخوردار باشد و تفاوت ماده آزمایشی در بین واحدهای آزمایشی یک بلوک تا حد امکان به کمترین مقدار خود تقلیل یابد.

اگر در یک گروه مربوط به بلوک، کلیه تیمارهای مورد آزمایش وجود داشته باشند آنرا بلوک کامل و اگر در تشکیل بلوک فقط عده ای از تیمارها شرکت داشته باشند، آنرا بلوک ناقص می نامند. در حالت کلی بلوک می تواند منطقه کوچکی در مزرعه یا گروهی از جانوران و یا زمان های متفاوت به کار بردن تیمارها در واحدهای آزمایشی و باشد.

نمونه Sample:

در جامعه نامحدود، اندازه گیری صفات تمام افراد امکانپذیر نیست، بنابراین محقق مجبور است به شیوه ای تصادفی گروهی از افراد را گزینش نموده و صفت یا صفات موردنظر آنها اندازه گیری شده و نتایج بر اساس اصول و قوانین ویژه ای به کل جامعه تعمیم داده شود، چنین گروه یا گروههایی نمونه نامیده می شوند.

بدیهی است که با افزایش واحدهای نمونه، تخمین صفات جمعیت دقیق تر خواهد بود ولی باید توجه داشت که امکانات مورد نیاز برای انجام این کار نیز محدود است.

خطای آزمایشی Experimental Error:

به طور کلی خطای آزمایش به آن دسته از تغییراتی گفته می شود که به وسیله عواملی خارج از کنترل پژوهشگر به وجود می آید. به طور کلی خطای آزمایش به دلیل عوامل زیر به وجود می آید:

۱- اختلاف ژنتیکی هر گیاه نسبت به گیاه دیگر در کرت هایی که تیمار یا رفتار یا به طور کلی تمام عوامل اثر کننده بر آنها یکسان بوده است.

۲- اختلاف هایی که در اثر غیر یکنواختی های جزئی در بکار بستن یا در انجام آزمایش به وجود می آید. برای مثال توزیع نابرابر کود در کرت ها، نبود دقت در اندازه گیری ها، نابرابری عمق کاشت در خاک و ... به ایجاد این دسته از اختلاف ها منجر می شود. البته باید کوشید تا چنین اختلاف هایی تا حد ممکن کاهش یابد. برای کاهش خطای آزمایشی می توان از راه حل های زیر استفاده نمود:

- مواد آزمایشی همگن یا مشابه انتخاب شود
- تکرارهای آزمایش متناسب اختیار شوند
- طرح مناسب به کار برده شود

در مقایسه آزمایشهای مشابه، آزمایشی که واریانس اشتباه آزمایشی کمتری داشته باشند، دارای دقت بیشتری است و نیز از چند آزمایش انجام یافته، آزمایشی که ضریب تغییرات (C.V) کمتری دارد از دقت زیادتری برخوردار است زیرا خطای آزمایشی کمتری دارد.

اصول کلی در طرح های آزمایشی:

به کارگیری اصول زیر در طرحهای آزمایشی نتایج قابل قبول و معتبری به ارمغان می آورد. این اصول باید در کلیه آزمایشهای مقایسه ای در مزرعه و آزمایشگاه رعایت شوند:

۱- تکرار تیمارها Replication

تکرار بدین معنی است که یک تیمار چند بار تکرار شود. تکرار تیمار در برآورد خطای آزمایشی و مقایسه هر چه دقیقتر اثر تیمارها نقش دارد. اصولاً دقت یک آزمایش بستگی به تعداد تکرار تیمار مورد نظر دارد. هر چه قدر تعداد تکرار بیشتر باشد، دقت آزمایش بیشتر می شود ولی در هر آزمایش حد متوسطی برای تعداد تکرار وجود دارد که اگر از این حد تجاوز نماید دقت آزمایش نه تنها به همان اندازه افزایش نمی یابد بلکه اضافه نمودن تکرار سبب افزایش هزینه و غیر یکنواختی ماده آزمایش و به طور غیر مستقیم زیاد شدن خطای آزمایش می گردد. به تجربه ثابت شده است که در اکثر طرحهای آزمایشی تعداد تکرار بسته به حساسیت آزمایشی بین ۴ تا ۸ انتخاب می شود (با افزایش حساسیت آزمایش تعداد تکرار زیاد می شود).

۲- توزیع تصادفی (Randomization) تیمارها در واحدهای آزمایشی:

در یک آزمون آماری باید شرایطی موجود باشد تا آزمون معتبر باشد. یکی از این شرایط مستقل بودن مشاهدات و خطای آزمایشی از هم است. برای تحقق این شرط لازم است که تیمارها در واحدهای آزمایشی به صورت تصادفی قرار گیرند. به طور کلی توزیع تصادفی موجب می گردد که تیمارها و تکرار آنها شانس مساوی برای قرار گرفتن در واحدهای آزمایشی متفاوت (مثلاً حاصلخیز و یا فقیر) داشته باشند و لذا اثر عوامل غیر قابل کنترل در آزمایش روی تیمارهای موجود تعدیل می گردد.

۳- کنترل موضعی Local Control:

کاهش خطای آزمایشی با ایجاد محدودیت هایی برای انتخاب تصادفی کامل انجام می گیرد که به آن کنترل موضعی گفته می شود (مثل طرح های بلوک).

برای کاهش خطای آزمایشی چند راه وجود دارد:

الف- انتخاب طرح های مناسب که بتوان منابع متغیر را در آنها کنترل کرد.

ب- از طریق کاهش یا حذف منابع تغییرات با استفاده از اطاق های رشد و آزمایشگاه ها که نتایج حاصل از این آزمایشات اغلب قابل تعمیم به شرایط طبیعی نیست.

پاره ای از تکنیک های اجرایی در طرح های آزمایشات:

۱- اندازه واحدهای آزمایشی plot size :

تجربه نشان داده که هر قدر اندازه کرت ها کوچکتر باشد، خطای آزمایش کمتر خواهد بود. البته اندازه کرت ها به کیفیت زمین و نوع گیاه نیز بستگی دارد و نباید از حد معینی کوچکتر باشد. اگر زمین یکنواخت باشد می توان کرت ها را کوچکتر انتخاب نمود ولی اگر زمین یکنواخت نباشد، اندازه کرت ها را می توان بزرگتر انتخاب کرد.

۲- شکل واحدهای آزمایشی plot form :

اگر تغییرات حاصلخیزی خاک نامشخص باشد، بایستی بلوک ها تا حد امکان به شکل مربع یا نزدیک به آن باشد، چرا که کرت های واقع شده در یک بلوک مربع یک دست تر می باشند تا کرت هایی که در بلوک طولی قرار گرفته اند. حال اگر حاصلخیزی خاک در یک جهت مثلاً شمال به جنوب متغیر باشد و محقق از وجود آن آگاهی داشته باشد، بایستی بلوکها عمود بر جهت تغییر حاصلخیزی و کرت های داخل بلوک را به موازات تغییر حاصلخیزی قرار دهد تا حداکثر تشابه بین کرت های یک بلوک بوجود آید.

۳- حاشیه border:

اغلب اتفاق می افتد که تیمارها همدیگر را تحت تاثیر قرار می دهند در این گونه موارد برای جلوگیری از تاثیر کرت های مختلف بر روی یکدیگر حاشیه در نظر گرفته می شود، بدین معنی که معمولاً کرت های آزمایشی را در مجاورت هم قرار می دهند ولی در دو طرف هر کرت بسته به نوع گیاه چند ردیف می کارند و یا در اطراف هر کرت قسمتی را به عنوان حاشیه خالی می گذارند و یا به صورت پشته در می آورند تا آبیاری هر واحد آزمایشی به صورت مجزا انجام گیرد. برای اطمینان بیشتر، ردیف های حاشیه و قسمت های کناری در آمارگیری نمونه برداری و محاسبات منظور نمی شوند.

طرح های آزمایشی پایه:

به طور کلی طرح ها را می توان به دو دسته تقسیم کرد:

۱- طرح هایی که در آنها فقط اثر یک منبع پراکندگی بررسی می شود. در این دسته تنها یک طرح وجود دارد و آن طرح کاملاً تصادفی Completely randomized design می باشد.

۲- طرحهایی که در آنها اثر بیش از یک منبع پراکندگی بررسی می شود، از جمله این طرح ها می توان به بلوک های کامل تصادفی Randomized complete block design و طرح مربع لاتین Latin square design اشاره نمود. سه طرح ذکر شده در فوق، طرحهای پایه و اصلی را در آمار تشکیل می دهند که بقیه طرح ها به طور مستقیم یا غیر مستقیم از این سه طرح منشا می گیرند. انتخاب یکی از این سه طرح پایه برای انجام یک آزمایش با توجه به نکات زیر انجام می گیرد:

الف- تعداد و نوع تیمار

ب- تعداد تکرار (میزان دقت آزمایش)

ج- یکنواختی یا غیریکنواختی ماده آزمایشی

طرح کاملاً تصادفی:

ساده ترین آزمایش در کشاورزی، بررسی تغییرات یک عامل است. در این طرح، تیمارها به طور کاملاً تصادفی در کرت ها یا واحدهای آزمایشی قرار می گیرند. به طوری که هر یک از کرت ها شانس مساوی برای دریافت هر یک از تیمارها دارند. این طرح موقعی مورد استفاده قرار می گیرد که واحدهای آزمایشی یکنواخت باشند و یا ماده آزمایشی دارای تغییرات یکسانی در جمیع جهات باشد. بدین جهت برای آزمایش های گلخانه ای و آزمایشگاهی خیلی مناسب است چون در آنها اثرات محیط یکنواخت و قابل کنترل است ولی در آزمایشات صحرایی به علت عدم یکنواختی قطعات مختلف زمین، قابلیت استفاده از آن محدود است.

مزایای طرح کاملاً تصادفی:

۱- انعطاف پذیری: یکی از خصوصیات این طرح، قابل انعطاف بودن آن است یعنی پژوهشگر می تواند هر تعداد تیمار و برای هر تیمار، هر تعداد تکرار را انتخاب نماید. در صورتی که دو طرح اصلی دیگر، تکرارهای مساوی برای تیمارهای مختلف لازم است. اگر تعداد تکرار برای تمام تیمارها یکسان باشد، طرح را متعادل و در غیر اینصورت طرح را نامتعادل می نامند.

۲- تجزیه واریانس ساده: هنگامی که تکرار تیمارها برابر نیست. تجزیه واریانس پیچیده نیست. این موضوع در مورد طرح های پیچیده تر صادق نیست.

۳- کرتهای از دست رفته: در صورت از بین رفتن یافته هایی از کرت ها، تغییری در عملیات تجزیه واریانس پیدا نمی شود و به تخمین کرتهای از دست رفته نیازی نیست. در صورتی که در طرح های دیگر کرت یا کرتهای از دست رفته را باید حتماً اول تخمین زد و سپس به تجزیه واریانس پرداخت.

۴- حداکثر درجه آزادی خطا: تعداد درجه آزادی خطای آزمایش در این طرح از تمام طرح های دیگر بیشتر است زیرا به جز درجه آزادی تیمار، همه درجه آزادی به خطای آزمایش تعلق می گیرد.

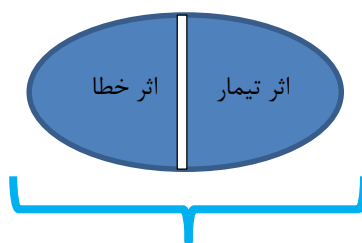
معایب طرح کاملاً تصادفی:

یکی از عیوب اصلی این طرح، کم بودن کارایی آن است و در صورت ناهمگنی واحدهای آزمایشی دقت پایین است.

مدل ریاضی طرح:

در این طرح چنانچه بیان شد فقط یک عامل پراکندگی یعنی اثر تیمارها بررسی می‌شود، لذا تغییرات کل به دو قسمت مربوط به اثر تیمار و خطای آزمایشی خواهد شد. این وضعیت با یک فرمول ساده ریاضی قابل تعریف است.

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



مقدار کل تغییرات

در این فرمول X_{ij} : نشان دهنده هر مشاهده (داده) در آزمایش، μ : میانگین کل جمعیتی که از طریق نمونه‌ها با فرض صفر مورد بررسی قرار می‌گیرد. t_j : اثر هر تیمار و ε_{ij} : اثر خطای آزمایشی است. به عبارت دیگر اجزای تشکیل دهنده هر داده در این طرح، میانگین جمعیت، اثر تیمار و اثر عوامل ناشناخته می‌باشد.

از رابطه بالا چنین بر می‌آید که اگر بین تیمارها اختلافی وجود نداشته باشد و خطای آزمایشی به صفر تقلیل یابد، مقدار هر مشاهده برابر میانگین کل جمعیت خواهد بود یعنی تمام داده‌ها برابر μ خواهد بود.

در طرح کاملاً تصادفی جدول تجزیه واریانس (ANOVA) به صورت زیر خواهد بود:

منبع تغییر	درجه آزادی df	مجموع مربعات SS	میانگین مربعات MS	F
کل	$r \times t - 1$	TSS		
تیمار	$t - 1$	VSS	VMS	VMS/EMS
خطا	$t(r - 1)$	ESS	EMS	

t: تیمار و r: تکرار

$$TSS = \sum_{i=1}^r \sum_{j=1}^t x_{ij}^2 - \frac{(\sum_{i=1}^r \sum_{j=1}^t x_{ij})^2}{r \times t} = \sum_{i=1}^r \sum_{j=1}^t x_{ij}^2 - CF \quad df = r \times t - 1$$

$$VSS = \frac{\sum_{j=1}^t x_{.j}^2}{r} - CF \quad df = t - 1$$

$$ESS = \sum_{i=1}^r \sum_{j=1}^t x_{ij}^2 - \frac{\sum_{i=1}^r x_{i.}^2}{r} = TSS - VSS \quad df = t(r - 1)$$

مثالی از طرح کاملاً تصادفی متعادل:

یک کارشناس تصمیم می گیرد که اثرات پنج نوع عملیات آماده کردن زمین را روی رشد طولی نهال های کاج کاشته شده در یک جنگلکاری با هم مقایسه نماید. وی ۲۵ پلات یا کرت آزمایشی را انتخاب می نماید و هر نوع عملیات را بطور تصادفی در ۵ پلات یا کرت انجام داده است. سپس پلات ها را با نهالهای موردنظر جنگلکاری نموده و در پایان سال پنجم روش طولی درختان را اندازه گیری نمود که میانگین رویش طولی درختان در هر پلات به شرح جدول زیر می باشد، مطلوبست بررسی این موضوع که آیا نوع عملیات آماده کردن زمین در روی رویش طولی درختان موثر است یا نه؟

انواع عملیات آماده کردن (تیمارها)

	A	B	C	D	E	
	۱۵	۱۶	۱۳	۱۱	۱۴	
	۱۴	۱۴	۱۲	۱۳	۱۲	
	۱۲	۱۳	۱۱	۱۰	۱۲	
	۱۳	۱۵	۱۲	۱۲	۱۰	
	۱۳	۱۴	۱۰	۱۱	۱۱	
مجموع	۶۷	۷۲	۵۸	۵۷	۵۹	۳۱۳
میانگین	۱۳/۴	۱۴/۴	۱۱/۶	۱۱/۴	۱۱/۸	۱۲/۵۲

فرض های آماری : $H_0 = 5$ تیمار مثل هم هستند و اختلاف معنی داری با هم ندارند

$H_A =$ حداقل یک تیمار وجود دارد که با بقیه اختلاف معنی داری دارد.

$$TSS = 15^2 + 14^2 + \dots + 11^2 - (313^2/5 \times 5) = 64.24$$

$$df = r \times t - 1 = 5 \times 5 - 1 = 24$$

$$CF = (313)^2 / 25 = 3918.76$$

$$VSS = (67^2 + \dots + 59^2) / 5 - 3918.76 = 34.64$$

$$df = 5 - 1 = 4$$

$$ESS = TSS - VSS = 64.24 - 34.64 = 29.60$$

$$df = 5(5-1) = 20$$

جدول تجزیه واریانس:

F	میانگین مربعات MS	مجموع مربعات SS	درجه آزادی df	منبع تغییر
۵/۵۸**		۶۴/۲۴	۲۴	کل
	۸/۶۶	۳۴/۶۴	۴	تیمار
	۱/۴۸	۲۹/۶۰	۲۰	خطا

$$F_{5\%, 4, 20} = 2.866$$

$$F_{1\%, 4, 20} = 4.431$$

پیش از مقایسه میانگین ها معمولا C.V یا ضریب تغییرات آزمایش را محاسبه می کنند:

$$C.V. = \left(\frac{\sqrt{MSE}}{\bar{X}_{..}} \right) \times 100 = \left(\frac{\sqrt{1.48}}{12.52} \right) \times 100 = 9.71\%$$

معمولا هر چه C.V در آزمایش بزرگتر باشد، دقت انجام آزمایش کم یا اشتباه آزمایشی بزرگ است. با وجود این، مقدار C.V به ماهیت آزمایش و نوع صفت مورد اندازه گیری بستگی دارد. برخی از صفات از C.V بزرگتری نسبت به سایر صفات در شرایط یکسان آزمایش برخوردار هستند. گاهی اوقات گفته می شود که اگر C.V آزمایش بالاتر از ۳۰ باشد آزمایش باید تکرار شود. شاید این موضوع در مواردی که تفاوت بین تیمارها معنی دار نیست صادق باشد، زیرا در این حالت، عدم معنی دار شدن می تواند ناشی از عدم تفاوت واقعی بین تیمارها یا بزرگ بودن اشتباه آزمایش باشد. اما اگر با وجود C.V بزرگ آزمایش، تفاوت بین تیمارها معنی دار شود (بخصوص در سطح ۱٪) نیازی به تکرار آزمایش نیست.

روش های مقایسه میانگین:

وقتی که مقدار F جدول تجزیه واریانس معنی دار نباشد کار تمام شده است، اما اگر مقدار F در سطوح آماری موردنظر (معمولا ۹۵ درصد و ۹۹ درصد) معنی دار باشد، این سوال مطرح می گردد که کدام تیمار و یا تیمارها با یکدیگر اختلاف معنی دار دارند. برای پاسخ به این سوال باید از آزمونهای مربوط به مقایسه میانگین ها استفاده شود. محقق می تواند بر حسب شرایط آزمایش یکی از آنها را مورد استفاده قرار دهد که در زیر به شرح آن دسته از آزمونهایی که بیشتر مورد استفاده قرار می گیرند پرداخته می شود:

آزمون حداقل اختلاف معنی دار LSD (Least Significant Different) :

روش های مقایسه میانگین:

وقتی که مقدار F جدول تجزیه واریانس معنی دار نباشد کار تمام شده است، اما اگر مقدار F در سطوح آماری موردنظر (معمولا ۹۵ درصد و ۹۹ درصد) معنی دار باشد، این سوال مطرح می گردد که کدام تیمار و یا تیمارها با یکدیگر اختلاف معنی دار دارند. برای پاسخ به این سوال باید از آزمونهای مربوط به مقایسه میانگین ها استفاده شود. محقق می تواند بر حسب شرایط آزمایش یکی از آنها را مورد استفاده قرار دهد که در زیر به شرح آن دسته از آزمونهایی که بیشتر مورد استفاده قرار می گیرند پرداخته می شود:

آزمون حداقل اختلاف معنی دار LSD (Least Significant Different) :

آزمون LSD برای اولین بار حدود ۷۵ سال پیش به کار برده شد و هنوز هم مورد استفاده است. یکی از مزایای آن سادگی و سرعت عمل آن است. زیرا در این آزمون یک مقدار ثابت به نام LSD محاسبه و اختلاف بین میانگین ها را با آن مقایسه می کنند.

شیوه انجام این آزمون آن است که پیش از انجام مقایسه ها باید مقدار LSD یا کمترین اختلاف جدا کننده بین دو تیمار را محاسبه کرد. این مقدار از فرمول زیر بدست می آید:

$$LSD_{\alpha} = S_d \times t_{\alpha/2, d_{fE}}$$

که در رابطه فوق X : سطح احتمال آماری، $X/2$: نشان دهنده دو طرفه بودن آزمون، d_{fE} : درجه آزادی خطای آزمایش، t : مقدار معین و ثابتی که از جدول t بدست می آید و S_d : اشتباه معیار یا خطای استاندارد اختلاف بین دو میانگین است که از رابطه زیر بدست می آید:

$$S_d = \sqrt{\frac{2 \text{ MSE}}{r}} = \sqrt{\frac{SE^2}{r}}$$

MS_E : واریانس خطای آزمایش و r : تعداد تکرار تیمار است.

$$S_d = \sqrt{\frac{2 \times 1.48}{5}} = 0.769$$

$$LSD \ 5\% = 0.769 \times t \ 5\%, \ 20 = 0.769 \times 2.086 = 1.61$$

$$LSD \ 1\% = 0.769 \times t \ 1\%, \ 20 = 0.769 \times 2.845 = 2.19$$

تیمار	B	A	E	C	D
میانگین	۱۴,۴	۱۳,۴	۱۱,۸	۱۱,۶	۱۱,۴
					۹۵٪
					۹۹٪

آزمون دانکن Duncan test:

مراحل انجام آزمون دانکن به شرح زیر است:

۱- محاسبه اشتباه معیار یا انحراف استاندارد میانگین تیمارها یعنی S_x که مقدار آن برای کلیه میانگین ها مساوی است.

$$S_x = \sqrt{\frac{MSE}{r}} = \sqrt{\frac{1.48}{5}} = 0.544$$

۲- محاسبه "حداقل دامنه های معنی دار LSR" که از حاصلضرب S_x در اعداد جدولی به نام "دامنه های معنی دار استیودنت SSR" با توجه به درجه آزادی خطا بدست می آیند.

$$LSR_{\alpha} = S_x \times SSR_{\alpha, d_{fE}}$$

جداول مذکور در سطوح احتمال ۵٪ و ۱٪ توسط دانکن تهیه و به افتخار آماردان مشهور گاست GOSSET که خود را با نام مستعار استیودنت نامید، نامگذاری شد. در این جداول چون برای مقایسه دست کم دو میانگین می بایست وجود داشته باشد، تعداد دامنه مورد مقایسه از ۲ شروع می شود در مثال فوق $df_E = 20$ و $t = 5$ لذا خواهیم داشت:

دامنه، درجه آزادی ۲۰	۲	۳	۴	۵
SSR 5%	۲,۹۵	۳,۱۰	۳,۱۸	۳,۲۵
SSR 1%	۴,۰۲	۴,۲۲	۴,۳۳	۴,۴۰
LSR 5%	۱,۶۰	۱,۶۹	۱,۷۳	۱,۷۷
LSR 1%	۲,۱۹	۲,۳۰	۲,۳۵	۲,۳۹

اولین نکته ای که باید توجه داشت این است که مقدار LSR در زیر دامنه ۲، برابر LSD در احتمالهای مربوطه است و این همواره صادق است بنابراین آزمون به طور غیر مستقیم جواب آزمون LSD را هم در بر دارد.

۳- ردیف کردن میانگین تیمارها از بزرگترین به کوچکترین (و یا بر عکس):

تیمار	B	A	E	C	D
میانگین	۱۴,۴	۱۳,۴	۱۱,۸	۱۱,۶	۱۱,۴

۴- مقایسه دو به دو تیمارها:

هرگاه در دو آزمایش t تیمار وجود داشته باشد $t(t-1)/2$ مقایسات لازم است. در مثال فوق $5(5-1)/2 = 10$ مقایسه لازم است. در سطح ۵٪ داریم:

$$14.4 - 13.4 = 1 < \text{LSR}(\text{damane } 2) 1.6$$

$$14.4 - 11.8 = 2.6 > \text{LSR}(\text{damane } 3) 1.69$$

$$14.4 - 11.6 = 2.8 > \text{LSR}(\text{damane } 4) 1.73$$

$$14.4 - 11.4 = 3 > \text{LSR}(\text{damane } 5) 1.77$$

و به همین ترتیب:

$$13.4 - 11.8 = 1.6 = \text{LSR}(\text{damane } 2) 1.6$$

$$13.4 - 11.6 = 1.8 > \text{LSR}(\text{damane } 3) 1.69$$

$$13.4 - 11.4 = 2 > \text{LSR}(\text{damane } 4) 1.73$$

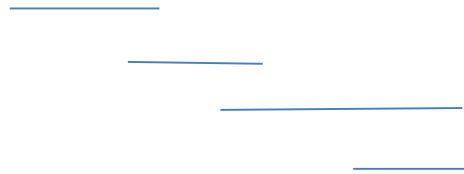
$$11.8 - 11.6 = 0.2 < \text{LSR}(\text{damane } 2) 1.6$$

$$11.8 - 11.4 = 0.4 < \text{LSR}(\text{damane } 3) 1.69$$

$$11.8 - 11.4 = 0.2 < \text{LSR}(\text{damane } 2) 1.6$$

در نتیجه می توان نتایج فوق را بصورت خطی در زیر نشان داد:

14.4 13.4 11.8 11.6 11.4



یا با توجه به نتایج مشابه در دو خط پایانی به صورت زیر می آید:

14.4 13.4 11.8 11.6 11.4



همین روش در سطح ۱٪ نیز می توان مشاهده نمود که نتایج به صورت زیر است:

14.4 13.4 11.8 11.6 11.4



می توان علاوه بر مقایسات به صورت خطی از روش های دیگری نیز بهره جست:

- مقایسه میانگین های تمام تیمارها به وسیله تشکیل جدول تفاضل ها: در این جدول قدر مطلق تفاضل هر دو میانگین در جدول وارد می شود، حال می توان هر مقدار را با LSR مربوطه مقایسه کرده و معنی دار بودن آنرا تعیین نمود: در این روش جدول مربع تشکیل داده و هر ضلع آنرا به $t-1$ جز تقسیم می کنند.

تیمار	B	A	E	C
A	1NS			
E	2.6**	1.6 ^{NS}		
C	2.8**	1.8*	0.2 ^{NS}	
D	3**	2*	0.4 ^{NS}	0.2 ^{NS}

مقایسه میانگین های تمام تیمارها به روش الفبا:

این روش، اختلاف ها را حتی به طور ساده تری نسبت به روش خط کشی نشان می دهد و بیشتر در مقاله های علمی بکار می رود. برای به کار بردن این روش ابتدا باید روش خط کشی را دنبال کرد و پس از انجام این کار، هر خطی را با یکی از حروف الفبا نشان می دهند (جایگزین می نمایند) و معمولاً کار را از بزرگترین میانگین شروع می نمایند همانند جدول زیر:

تیمار	میانگین	۵٪	۱٪
-------	---------	----	----

B	11.4	a	a
A	13.4	ab	ab
E	11.8	bc	b
C	11.6	c	b
D	11.4	c	b

آزمون توکی Tukey test:

این روش که آزمون اختلاف معنی دار قابل اعتماد HSD نامیده می شود مانند آزمون دانکن برای مقایسه همه میانگین ها به صورت دو به دو بکار می رود ولی مانند روش LSD یک مقدار ثابت برای هر سطح احتمال محاسبه می شود. روش انجام این آزمون به شرح زیر است:

۱- محاسبه اشتباه معیار میانگین S_x (بر حسب آزمون دانکن)

$$S_x = \sqrt{\frac{MSE}{r}} = \sqrt{\frac{1.48}{5}} = 0.544$$

۲- محاسبه مقدار HSD یا W :

$$HSD_{\alpha} = W_{\alpha} = S_x \times Q_{\alpha, (t, dfe)}$$

$$W_{5\%} = 0.544 \times 4.23 = 2.301$$

$$W_{1\%} = 0.544 \times 5.29 = 2.878$$

۳- مرتب کردن میانگین ها به صورت نزولی یا صعودی و مقایسه دو به دو آنها

تیمار	B	A	E	C	D
میانگین	۱۴,۴	۱۳,۴	۱۱,۸	۱۱,۶	۱۱,۴

۹۵٪

۹۹٪

در هر سطح احتمال تفاضل دو میانگین با مقدار W مقایسه می شود، اگر این تفاضل مساوی یا بزرگتر از مقدار W باشد، بین دو میانگین اختلاف معنی دار وجود خواهد داشت.

در سطح احتمال ۵٪ نتیجه دو آزمون دانکن و توکی متفاوت است و تعداد اختلافات معنی دار در آزمون دانکن بیشتر از توکی است.

آزمون استیودنت - نیومن - کلز (SNK) Student- Newman- Keuls:

روش SNK مانند LSR یک آزمون چند دامنه‌ای است، در حقیقت این آزمون از سویی همانند روش دانکن و از سویی همانند روش توکی است زیرا اولاً بیشتر از یک دامنه برای مقایسه اختلافات محاسبه می‌شود (همانند آزمون دانکن) و دوم اینکه همانند روش توکی جدول q_x به کار می‌رود، روش انجام این آزمون به شرح زیر است:

۱- محاسبه اشتباه معیار میانگین S_x همانند روش های قبلی

$$S_x = \sqrt{\frac{MSE}{r}} = \sqrt{\frac{1.48}{5}} = 0.544$$

۲- محاسبه مقدار SNK:

$$SNK_{\alpha} = S_x \times q_{\alpha, (R, dfE)}$$

بسته به سطح احتمال دلخواه (۵٪ یا ۱٪) مقادیر عددی q با توجه به درجه آزادی خطا و دامنه موردنظر R از جدول توکی استخراج می‌شود:

R	۲	۳	۴	۵
q 5% (R,20)	۲,۹۵	۳,۵۸	۳,۹۶	۴,۲۳
q 1% (R,20)	۴,۰۲	۴,۶۴	۵,۰۲	۵,۲۹
SNK 5%	۱,۶۰	۱,۹۵	۲,۱۵	۲,۳۰
SNK 1%	۲,۱۹	۲,۵۲	۲,۷۳	۲,۸۸

مقدار SNK_{α} در حالت $R=2$ برابر LSD_{α} و در حالت $R=7$ برابر HSD در آزمون توکی است.

۳- مرتب کردن میانگین و مقایسه آنها مانند روش دانکن

۴-

تیمار	B	A	E	C	D
میانگین	۱۴,۴	۱۳,۴	۱۱,۸	۱۱,۶	۱۱,۴

۹۵٪

۹۹٪

