



طراحی آزمایش‌های کشاورزی

(رشته علوم کشاورزی)

مؤلفان: بهروز فوقی، علیرضا طالعی، غلامعلی اکبری

فهرست مطالب

پیشگفتار مؤلفین	یازده
فصل اول. مختصری از آمار جهت یادآوری	۱
۱-۱ تعریف‌ها	۱
علم آمار (Statistics)	۱
جامعه (Population)	۱
نمونه (Sample)	۲
معیار نمونه‌ای (Statistic) و معیار جامعه‌ای (Parameter)	۲
مشاهده‌ها یا داده‌ها (data)	۲
منحنی نرمال (Normal curve)	۳
نشان دادن یا خلاصه کردن اندازه‌های آماری	۴
۲-۱ معیارهای تمایل به مرکز	۴
۳-۱ معیارهای پراکندگی	۵
۴-۱ ضریب تغییرات (C.V. = Coefficient of Variation)	۷
۵-۱ آزمون فرض (Test of Hypothesis)	۷
۶-۱ مقایسه دو میانگین	۸
۱-۶-۱ مقایسه دو میانگین (نمونه بزرگ با جامعه)	۹
۲-۶-۱ مقایسه دو میانگین با نمونه‌های بزرگ	۱۰
۳-۶-۱ مقایسه دو میانگین با نمونه‌های کوچک	۱۲
۷-۱ مقایسه دو واریانس	۱۹
۸-۱ مقایسه بیش از دو واقعه	۲۰
سؤال‌های فصل مختصری از آمار جهت یادآوری	۲۳

۲۷	فصل دوم. طرح‌های آزمایش‌های کشاورزی
۲۷	۱-۲ اصطلاح‌ها و تعاریف‌ها
۲۷	۱-۱-۲ علم (Science)
۲۷	۲-۱-۲ آزمایش (Experiment)
۲۸	۳-۱-۲ طرح‌های آزمایشی (Experimental Designs)
۲۸	۴-۱-۲ تیمار یا رفتار (Treatment)
۲۸	۵-۱-۲ مواد آزمایشی (Experimental Material)
۲۹	۶-۱-۲ واحد آزمایشی، کرت یا پلات (Experimental Unit = Plot)
۲۹	۷-۱-۲ بلوک (Block)
۲۹	۸-۱-۲ داده‌ها و مشاهدات (Data)
۲۹	۲-۲ مفاهیم کلی در سازمان‌دهی و قضاوت روی نتایج طرح‌های آزمایشی
۲۹	۱-۲-۲ انواع تغییرها
۳۳	۲-۲-۲ اشتباه‌های آزمایشی (Experimental Errors)
۳۳	۳-۲ اصول کلی در طرح‌های آزمایشی
۳۴	۱-۳-۲ تکرار تیمارها
۳۶	۲-۳-۲ پخش تصادفی تیمارها در واحدها (Rondomization)
۳۷	۳-۳-۲ کنترل منابع تغییرهای شناخته شده
۳۷	۴-۳-۲ استفاده از طرح مناسب با هدف
۳۸	۵-۳-۲ ارائه طرح و انتشار نتایج
۳۹	۴-۲ تکنیک‌های اجرایی در آزمایش‌های مزرعه‌ای
۳۹	۱-۴-۲ اندازه واحدهای آزمایشی
۴۱	۲-۴-۲ شکل واحدهای آزمایشی
۴۲	۳-۴-۲ حاشیه
۴۳	۴-۴-۲ روش‌های تصادفی کردن تیمارها
۴۴	۵-۴-۲ عملیات زراعی و آمارگیری در خلال آزمایش
۴۴	۶-۴-۲ طرز نمونه‌گیری
۴۵	۷-۴-۲ تجزیه آماری و تفسیر نتایج
۴۶	۸-۴-۲ انواع طرح آزمایشی
۴۷	فصل سوم. طرح کاملاً تصادفی
۴۷	۱-۳ طرح کاملاً تصادفی (Completely Rondomized Design = C.R.D)
۴۸	۲-۳ مزایای طرح کاملاً تصادفی
۴۸	۳-۳ معایب طرح کاملاً تصادفی

۴۸	۴-۳ طرز پیاده کردن طرح (قرعه‌کشی تیمارها)
۵۰	۵-۳ مدل ریاضی طرح
۵۰	۶-۳ تجزیه آماری طرح کاملاً تصادفی با تکرارهای مساوی (طرح متعادل)
۵۵	۷-۳ طرح کاملاً تصادفی با تکرارهای نامساوی (طرح نامتعادل)
۵۷	۸-۳ طرح کاملاً تصادفی با بیش از یک مشاهده برای هر تیمار در هر تکرار
۶۳	۹-۳ مقایسه‌های متعامد یا مستقل (Orthogonal comparisons)
۷۰	سؤال‌های فصل طرح کاملاً تصادفی
۷۵	فصل چهارم. مقایسه میانگین تیمارها
۷۵	۱-۴ مقایسه میانگین تیمارها
۷۵	۲-۴ روش حداقل تفاوت معنی‌دار
۷۹	۳-۴ توکی
۸۱	۴-۴ روش چند دامنه‌ای دانکن
۸۵	سؤال‌های فصل مقایسه میانگین تیمارها
۸۷	فصل پنجم. طرح بلوک‌های کامل تصادفی
۸۷	۱-۵ طرح بلوک‌های کامل تصادفی (Randomized Complete Block Design) یا RCBD
۸۸	۲-۵ بلوک‌بندی
۹۰	۳-۵ تصادفی کردن و پیاده کردن طرح RCBD
۹۱	۴-۵ مزایا و معایب این طرح
۹۱	۵-۵ مدل ریاضی طرح بلوک‌های کامل تصادفی
۹۲	۶-۵ تجزیه آماری طرح بلوک‌های کامل تصادفی با یک مشاهده در هر واحد آزمایشی
۹۳	۱-۶-۵ تجزیه واریانس آزمایش ارقام یونجه
۹۴	۲-۶-۵ مقایسه میانگین تیمارها و تعیین نتایج آزمایش
۹۷	۷-۵ برآورد مقدار مشاهده گم شده
۹۸	۱-۷-۵ برآورد یک مشاهده از بین رفته
۱۰۰	۲-۷-۵ برآورد بیش از یک مشاهده از بین رفته
۱۰۳	۸-۵ تجزیه آماری طرح بلوک‌های کامل تصادفی با چند مشاهده در هر واحد آزمایشی
۱۰۵	۹-۵ سودمندی نسبی (Relative Efficiency = R.E.) یا کارایی نسبی
۱۰۶	سؤال‌های فصل طرح کاملاً تصادفی
۱۱۱	فصل ششم. طرح مربع لاتین
۱۱۱	۱-۶ طرح مربع لاتین (Latin Square Design)
۱۱۲	۲-۶ تصادفی قرار دادن تیمارها در واحدهای آزمایشی

۱۱۴	۳-۶ مزایا و معایب طرح مربع لاتین
۱۱۴	الف) مزایای طرح مربع لاتین
۱۱۵	ب) معایب طرح مربع لاتین
۱۱۵	۴-۶ مدل ریاضی طرح مربع لاتین
۱۱۷	۵-۶ تجزیه آماری طرح مربع لاتین با یک مشاهده در هر واحد آزمایشی
۱۱۹	۶-۶ مقایسه میانگین تیمارها
۱۱۹	۶-۶-۱ آزمون حداقل اختلاف معنی دار یا LSD
۱۲۰	۶-۶-۲ آزمون توکی
۱۲۰	۶-۶-۳ آزمون دانکن
۱۲۲	۷-۶ تخمین (برآورد) کورت از بین رفته
۱۲۳	۸-۶ سودمندی نسبی یا کارایی نسبی
۱۲۵	۹-۶ طرح مربع لاتین با بیش از یک مشاهده در هر واحد
۱۲۷	سؤال‌های فصل طرح مربع لاتین
۱۳۱	فصل هفتم. آزمایش‌های فاکتوریل
۱۳۱	۱-۷ آزمایش‌های فاکتوریل (Factorial Experiments)
۱۳۷	۲-۷ تصادفی کردن و پیاده کردن طرح F.E
۱۳۸	۳-۷ مزایا و معایب طرح آزمایش‌های فاکتوریل
۱۳۸	الف) مزایای آزمایش‌های فاکتوریل
۱۳۹	ب) معایب آزمایش‌های فاکتوریل
۱۳۹	۴-۷ مدل ریاضی طرح آزمایش‌های فاکتوریل
۱۴۰	۵-۷ تجزیه آماری آزمایش‌های چند عاملی
۱۴۷	۱- تجزیه تیمار به اثرهای تشکیل دهنده با روش کلی
۱۴۹	۲- تجزیه تیمار با روش فاکتوریل
۱۵۲	آزمایش فاکتوریل 2^3
۱۵۶	۶-۷ مقایسه میانگین‌ها در آزمایش‌های دو فاکتوره و سه فاکتوره
۱۵۸	مقایسه سطوح B (عمق‌های مختلف شخم پائیزه)
۱۵۹	الف) تجزیه واریانس
۱۶۴	ب) مقایسه میانگین‌ها
۱۶۵	سؤال‌های فصل آزمایش‌های فاکتوریل
۱۷۱	فصل هشتم. منحنی‌های پاسخ
۱۷۱	۱-۸ بررسی روند عکس‌العمل صفت مورد بررسی (منحنی‌های پاسخ)

۱۷۳	۲-۸ انواع روندها
۱۸۸	۳-۸ روند عکس العمل در آزمایش های دو فاکتوره
۱۹۴	۴-۸ روند عکس العمل در آزمایش های دو فاکتوری کمی
۲۰۷	۵-۸ تجزیه اثر متقابل در تجزیه واریانس به کمک معادله های چند جمله ای متعامد
۲۱۱	سؤال های فصل منحنی های پاسخ

۲۱۷	فصل نهم. اختلاط در آزمایش های فاکتوریل (Confounding)
۲۱۷	۱-۹ کلیات
۲۱۸	۲-۹ اصول و روش اختلاط
۲۲۲	۳-۹ طرز پیاده کردن طرح های اختلاط یافته
۲۲۲	۴-۹ انواع اختلاط
۲۲۵	۵-۹ تجزیه آماری آزمایش فاکتوریل 2^n با اختلاط کامل
۲۲۷	۶-۹ تجزیه آماری آزمایش فاکتوریل 2^3 با اختلاط ناقص
۲۳۳	۷-۹ اختلاط در آزمایش 3^2
۲۳۵	۸-۹ اختلاط در آزمایش 3^3
۲۴۳	۹-۹ اختلاط در سایر آزمایش ها به جز آزمایش های K^n
۲۴۵	سؤال های فصل اختلاط در آزمایش های فاکتوریل (Confounding)

۲۴۹	فصل دهم. کرت های خرد شده (Split plot)
۲۴۹	۱-۱۰ کلیات
۲۵۴	۲-۱۰ مزایا و معایب آزمایش هایی کرت های خرد شده
۲۵۶	۳-۱۰ مدل آماری در آزمایش کرت های خرد شده
۲۵۷	۴-۱۰ تجزیه آماری در آزمایش های کرت های خرد شده
۲۶۰	۵-۱۰ طرز پیاده کردن و روش تصادفی کردن، جدول تجزیه واریانس و درجه های آزادی در آزمایش های کرت های خرد شده
۲۶۹	۶-۱۰ مقایسه میانگین ها
۲۷۷	۷-۱۰ برآورد مقدار مشاهده از بین رفته
۲۹۲	آزمون دانکن (LSR)
۲۹۲	الف) مقایسه سطوح فاکتور اصلی (A)
۲۹۲	ب) مقایسه سطوح فاکتور فرعی (B)
۲۹۳	ج) مقایسه سطوح B در یک سطح A
۲۹۴	د) مقایسه سطح یا سطوح B در سطوح مختلف
۲۹۵	آزمون توکی
۲۹۸	سؤال های فصل کرت های خرد شده

۳۰۵

۳۷۱

۳۹۱

پیوست ۱

پیوست ۲

منابع

بنام خدا

پیشگفتار مؤلفین

خدای بزرگ و یکتا را بسیار شکر می‌نماییم که توان انجام این مهم را عنایت بخشید تا بتوانیم سهم کوچکی در پیشبرد و اعتلای فرهنگ این کشور داشته و در جهت خود کفایی و استقلال کشور گام کوچکی برداریم. در این راستا آموزش و تربیت افراد متخصص و آگاه یکی از برنامه‌های اولیه می‌باشد، چرا که یکی از عوامل مهم موفقیت در امر آموزش، وجود منابع و مراجع مفید در رشته‌های مختلف علوم است. از طرف دیگر، پژوهشگران و محققان علاقه‌مند به آگاهی از اصول و مبانی ریاضی آمار بوده تا بتوانند برداشت کلی و درستی از منطق ریاضی و دلایل علمی در طرح و اجرای آزمایش داشته باشند و آزمایش‌های کشاورزی را به طور مناسب طراحی کرده و در پایان به نتایج معتبری نیز دست پیدا کنند.

در این کتاب سعی شده است که از بحث تخصصی ریاضی تا آنجا که امکان داشته باشد، خودداری شده و نیز از اثبات فرمول‌های ریاضی صرف نظر گردد و فقط به ذکر نتایج اکتفا شده است. لذا این کتاب به عنوان یک مرجع اولیه و پایه‌ای برای دانشجویان و متخصصین علوم و رشته‌های مختلف فنی، کشاورزی و نیز علوم اجتماعی و روانشناسی کاربرد دارد. همچنین سعی شده است که مباحث مختلف آماری با مثال‌ها و مسایل مختلف و متنوع به صورت ساده‌تر و سلیس بیان شده و باعث فهم

بهرتر و عمیق‌تر مطالب گردد. به همین دلیل بر روی جزئیات و مراحل قدم به قدم محاسبه مسایل تأکید شده و به تفصیل مورد بررسی قرار گرفته است.

در فصل اول مفاهیم آماری به صورت مختصر و مفید جهت یادآوری بیان شده است. در فصل دوم، اصول و مبانی طرح‌های مختلف کشاورزی از نقطه نظر اهداف، نحوه و نوع پیاده کردن، هدف از تصادفی پیاده کردن و مسایل دیگر مورد بررسی قرار گرفته است. در فصول سوم و پنجم و ششم طرح‌های پایه‌ای به دلیل اهمیت آنها در فصول جداگانه ارائه شده است. آزمایش‌های فاکتوریل، اختلاط در آزمایش‌های فاکتوریل و آزمایش‌های کرت‌های خرد شده در فصول هفتم و نهم و دهم آمده است. جهت بررسی روند عکس‌العمل صفات مورد بررسی به انواع خطی، درجه دوم و غیره، فصل هشتم در نظر گرفته شده است. مقایسه میانگین‌ها را که این نیز دارای اهمیت ویژه‌ای است در فصل چهارم گنجانده شده است. از آنجایی که هر کار چاپی دارای نواقصی می‌باشد، لذا از کلیهٔ سرووران عزیز استدعا دارد جهت تصحیح و رفع نقص، موارد اصلاحی خود را به مولفین منتقل نمایند. تا انشاء... در چاپ‌های بعدی مد نظر قرار بگیرد.

مؤلفین هم‌چنین بر خود لازم می‌دانند که از کمک‌های بی‌دریغ آقای دکتر کمال سادات اسمعیلان در چاپ و نشر، از خانم‌ها نوری، رفیعی و شاکرمی جهت تایپ و از کمیتهٔ فنی و انتشارات دانشگاه پیام نور و کلیهٔ کسانی که ما را در تهیه و نشر این کتاب یاری کرده‌اند، صمیمانه تشکر و قدردانی نماید. این کتاب تشکر نماید.

ومن... التوفیق

غلامعلی اکبری
استادیار پردیس ابوریحان
دانشگاه تهران

علیرضا طالعی
استاد پردیس کرج
دانشگاه تهران

بهروز فوقی
مربی پردیس ابوریحان
دانشگاه تهران

فصل اول

مختصری از آمار جهت یادآوری

۱-۱ تعریفها

لغت استاتستیک یا آمار، از دولت ریشه گرفته است و مفهوم اولیه آن بررسی در ظواهر اداره یک دولت است که شامل جمع‌آوری و تنظیم ارقام مربوط به قوانین و مقررات مالیات‌ها، گمرک‌ها، تاریخ تولد و غیره بوده است. امروزه روش‌های آماری تقریباً در کلیه علوم و به‌ویژه در پژوهش‌های علمی کاربرد دارد و در تحقیق‌های کشاورزی و زیستی جای به‌سزایی برای خود باز کرده است.

علم آمار (Statistics)

علم آمار عبارت از ابداع و کاربرد روش‌های ریاضی برای تخمین بوده و احتمال نتیجه‌گیری صحیح از مشاهدات در یک یا چند جامعه از طریق بررسی یک یا چند نمونه می‌باشد. در این کتاب بیشتر بر روی کاربرد روش‌ها و نه ابداع روش‌های جدید تکیه خواهد شد. چون که این کار معمولاً به عهده ریاضی‌دانان گزرده شده است.

جامعه (Population)

مجموعه نسبتاً بزرگی از افراد یا اشیاء که حداقل دارای یک صفت مشترک قابل

اندازه‌گیری باشند، جامعه نامیده می‌شود. مثلاً افراد متعلق به یک گونه گیاهی یا جانوری را جامعه آن گونه گیاهی یا جانوری می‌گویند. تعداد افراد متعلق به یک جامعه الزاماً نامحدود نیست چرا که ممکن است تعداد نسبتاً کمی از افراد تشکیل جامعه بدهند (جامعه محدود (Finite Population)). اما اغلب یک زیست‌شناس یا محقق کشاورزی با جامعه‌های بزرگ (جامعه نامحدود (Infinite Population) سر و کار دارد.

نمونه (Sample)

در جامعه‌های بزرگ، اندازه‌گیری کلیه افراد امکان‌پذیر نیست. لذا محقق باید به‌طور تصادفی (به نحوی که تمام افراد شانس مساوی برای انتخاب شدن داشته باشند) گروهی از افراد را انتخاب و صفت یا صفتهای مورد نظر را در آنها اندازه‌گیری نماید و نتایج را بر اساس روش‌های آماری مخصوص به جامعه تعمیم دهد. چنین گروه‌هایی نمونه نامیده می‌شوند. هر چه تعداد افراد نمونه بیشتر باشد برآورد مشخصه‌های جامعه دقیق‌تر خواهد بود. در آمار معمولاً اگر تعداد افراد نمونه ۳۰ و بیشتر از ۳۰ باشد نمونه را نمونه بزرگ و اگر کمتر از ۳۰ باشد، نمونه را نمونه کوچک می‌گویند. برای هر کدام از این دو حالت روش‌های ویژه‌ای به کار می‌روند که بعداً خواهیم دید.

معیار نمونه‌ای (Statistic) و معیار جامعه‌ای (Parameter)

معیار نمونه‌ای، شاخص صفت نمونه و معیار جامعه‌ای، شاخص صفت جامعه است. معیار نمونه‌ای معمولاً با حروف کوچک لاتین و معیار جامعه‌ای با حروف کوچک یونانی نشان داده می‌شوند مثلاً برای میانگین نمونه $\bar{x}$ و برای میانگین جامعه μ به کار می‌رود. مقدار عددی معیار نمونه‌ای قابل نوسان است در صورتی که مقدار عددی معیار جامعه‌ای همیشه ثابت است و تغییر نمی‌کند. در معیار نمونه‌ای، تعداد افراد نمونه، از نمونه‌ای به نمونه دیگر تفاوت می‌کند لذا شاخص‌های نمونه‌ای دائماً متفاوت می‌شود. ولی افراد جامعه چون ثابت و بدون تغییر است، شاخص‌های جامعه‌ای نیز ثابت می‌باشند.

مشاهده‌ها یا داده‌ها (data)

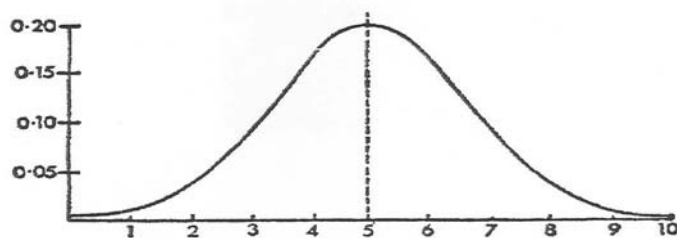
اعداد و ارقامی هستند که از اندازه‌گیری صفت مورد مطالعه در افراد یا شمارش افراد

برای داشتن یک صفت حاصل می‌شوند. یک صفت می‌تواند به صورت یک متغیر در نظر گرفته شود که قابل تغییر بوده و معمولاً آن را با x نشان می‌دهیم البته گاهی با y نیز نمایش داده می‌شود. یک متغیر می‌تواند به صورت پیوسته (Continuous) و یا ناپیوسته (Discontinuous) باشد. متغیر ناپیوسته یک صفت کیفی بوده که قابل اندازه‌گیری نیست و فقط می‌توان آن را با تعداد یا درصد نشان داد مثل مدل ماشین، نوع سیگار، جنسیت، تعداد دانه گندم در هر سنبله و غیره. ولی متغیر پیوسته یک صفت کلی بوده که قابلیت اندازه‌گیری را داراست مانند عملکرد دانه، ارتفاع گیاه، مقدار قند چغندر و غیره.

منحنی نرمال (Normal curve)

در رشته‌های علوم زیستی از جمله کشاورزی، توزیع اغلب داده‌ها یا متغیرها دارای شکلی شبیه به زنگوله است. این نوع توزیع اولین بار توسط گوس (Gauss) مطالعه و گزارش شد و به این جهت این توزیع را توزیع نرمال یا توزیع گوس گویند. اندازه‌گیری یک صفت در افراد نمونه با تعداد زیاد منجر به یک منحنی نرمال می‌شود که این منحنی با معیارهای نمونه قابل توصیف است. معمولاً منحنی نرمال با میانگین و انحراف استاندارد (انحراف معیار) مشخص می‌شود که بعداً در مورد آن بحث می‌شود. همه منحنی‌های نرمال دارای شکلی متقارن می‌باشند. یعنی افراد جمعیت به صورت متقارن در دو طرف میانگین پراکنده شده‌اند. منحنی تابع توزیع نرمال برای متغیر

پیوسته‌ای با $m=5$ و $\sigma=2$ به این صورت زیر می‌باشد:

$$p(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-m}{\sigma}\right)^2}$$


شکل ۱-۱ توزیع نرمال

نشان دادن یا خلاصه کردن اندازه‌های آماری

تعدادی اندازه (Data) را می‌توان به روش‌های مختلف مثل شکل، نمودار و یا جدول نشان داد. برای نشان دادن آمار به طریقه نمودار، از دو خط عمودی و افقی که محور مختصات گفته می‌شود، استفاده می‌شود که روی محور عمودی تعداد یا فراوانی افراد یا عناصر و روی محور افقی حدود طبقه‌ها (دسته‌ها) و یا عدد میانی طبقه‌ها مشخص می‌شود.

آخرین روش نشان دادن و خلاصه کردن آمار به وسیله جدول‌های وفور (Frequency Table) یا جدول‌های توزیع فراوانی بیان می‌گردند. معمولاً داده‌ها و یا اندازه‌ها، توسط دو خصوصیت معیارهای تمایل به مرکز و معیارهای پراکندگی مورد بحث قرار می‌گیرند.

۲-۱ معیارهای تمایل به مرکز

مرکز داده‌ها، معمولاً توسط معیارهای تمایل به مرکز مشخص می‌شود. بهترین، ساده‌ترین و مهم‌ترین معیار مرکزی رابه نام میانگین حسابی (Arithmetic Mean) می‌شناسیم که مفهوم آن مجموع مقادیر اندازه‌گیری‌ها، تقسیم بر تعداد این اندازه‌ها می‌باشد.

مثال ۱-۱: میانگین حسابی اعداد ۲، ۴، ۶، ۸ و ۱۰ چند می‌باشد؟

$$\text{میانگین حسابی} = \frac{۲ + ۴ + ۶ + ۸ + ۱۰}{۵} = ۶$$

معمولاً میانگین را در جامعه با μ و در نمونه‌ها با $\bar{x}$ نشان می‌دهند که فرمول

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_n}{n} \quad \text{آن: می‌باشد و به طور خلاصه چنین می‌باشد:}$$

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

در این فرمول صفت یا متغیر را با x و اندازه اول این متغیر را با x_1 و اندازه

آخر آن را با x_n (که می‌خوانیم x اندیس n) نمایش می‌دهیم تعداد داده‌ها را نیز با n

مشخص می‌کنیم. در این فرمول علامت $\sum$ یا سیگما به معنی مجموع ارقام است و مجموع $\sum_{i=1}^n$ از ۱ تا n خوانده می‌شود. باید خاطر نشان کرد که میانگین فقط قسمتی از اطلاعات موجود در داده‌ها را در بر دارد و لذا نمی‌تواند بیان‌کننده همه خصوصیت‌ها و مشخصه‌های اطلاعات آماری و یا داده‌ها باشد.

۳-۱ معیارهای پراکندگی

اندازه پراکندگی یا تغییرات موجود در داده‌ها، جهت توصیف بیشتر و بهتر اطلاعات آماری و یا داده‌ها لازم است. میانگین هیچ اطلاعاتی در این زمینه، در اختیار ما نمی‌گذارد. بهترین و متداول‌ترین روش، جهت اندازه‌گیری پراکندگی و نحوه توزیع داده‌ها، استفاده از انحراف معیار و مجذور آن، واریانس می‌باشد. که علامت آن در جامعه σ^2 و در نمونه s^2 می‌باشد (البته علامت واریانس در جامعه σ^2 و در نمونه‌ها s^2 می‌باشد). به منظور محاسبه واریانس یا مقدار متوسط پراکندگی داده‌ها نسبت به میانگین از فرمول زیر استفاده می‌کنیم:

$$\sigma^2 = \frac{(x_1 - \mu)^2 + (x_2 - \mu)^2 + \dots + (x_n - \mu)^2}{N}$$

و یا به طور خلاصه:

$$\sigma^2 = \frac{\sum_{i=1}^n (x_i - \mu)^2}{N}$$

و فرمول واریانس در نمونه چنین است:

$$s^2 = \frac{(x_1 - \bar{x})^2 + \dots + (x_n - \bar{x})^2}{N-1} = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{N-1}$$

چنانچه مشاهده می‌شود در فرمول واریانس نمونه $N-1$ وجود دارد که این را در حقیقت درجه آزادی یا (Degree of Freedom) گویند که معمولاً یکی کمتر از تعداد افراد نمونه است و به صورت ساده‌تر با df نمایش داده می‌شود. ضمناً به $\sum_{i=1}^n (x_i - \bar{x})^2$ معمولاً SS می‌گویند که در اینجا چون متغیر x است به آن SS_x می‌گویند.

مثال ۱-۲: میانگین و انحراف معیار جدول توزیع فراوانی زیر را به دست آورید.

طبقه‌ها	۰-۴	۴-۸	۸-۱۲	۱۲-۱۶	۱۶-۲۰	۲۰-۲۴	۲۴-۲۸
فراوانی‌ها	۲	۴	۸	۱۲	۸	۴	۲

حل: جدول را به ترتیب زیر تشکیل داده و ستون‌های مختلف را ایجاد و به محاسبه آنها می‌پردازیم.

طبقه‌ها	F_i	x_i	$(x - \bar{x})$	$(x - \bar{x})^2$	$F_i (x - \bar{x})^2$	$F_i x_i$
۰-۴	۲	۲	-۱۲	۱۴۴	۲۸۸	۴
۴-۸	۴	۶	-۸	۶۴	۲۵۶	۲۴
۸-۱۲	۸	۱۰	-۴	۱۶	۱۲۸	۸۰
۱۲-۱۶	۱۲	۱۴	۰	۰	۰	۱۶۸
۱۶-۲۰	۸	۱۸	+۴	۱۶	۱۲۸	۱۴۴
۲۰-۲۴	۴	۲۲	+۸	۶۴	۲۵۶	۸۸
۲۴-۲۸	۲	۲۶	+۱۲	۱۴۴	۲۸۸	۵۲
	$N = ۴۰$				۱۳۴۴	۵۶۰

در این جدول x_i متوسط عددی هر طبقه است که نماینده هر طبقه شمرده می‌شود. برای محاسبه میانگین در جدول توزیع فراوانی از فرمول زیر استفاده می‌کنیم.

$$\mu = \frac{\sum F_i x_i}{N}$$

و می‌دانیم $\sum F_i = N$ پس میانگین برابر است با:

$$\mu = \frac{۵۶۰}{۴۰} = ۱۴/۰$$

و انحراف معیار در جدول توزیع فراوانی برابر است با:

$$\sigma = \sqrt{\frac{\sum (x_i - \bar{x})^2 F_i}{N}}$$

و نتیجه چنین خواهد شد:

$$\sigma = \sqrt{\frac{۱۳۴۴}{۴۰}} = ۵/۸$$

قبلاً گفتیم که واریانس، مجذور انحراف معیار می‌باشد. ضمناً انحراف معیار را،

گاه انحراف استاندارد نیز می‌گویند.

۱-۴ ضریب تغییرات (C.V. = Coefficient of Variation)

ضریب تغییرات یکی از شاخص‌های پراکندگی است و فرمول آن چنین است:

$$C.V. = \frac{S}{\bar{x}} \times 100$$

این ضریب به صورت درصد بیان می‌شود و یک عدد مطلق است و دارای واحد نمی‌باشد. زیرا انحراف استاندارد و میانگین داده‌ها وقتی به یکدیگر تقسیم شوند، نتیجه عددی بدون واحد خواهد شد. کاربرد مهم این شاخص موقعی است که دو سری داده با یکدیگر مقایسه شوند و به خصوص در زمانی که واحد داده‌ها با یکدیگر از یک نوع نباشد، بهتر است.

مثال ۱-۳ اگر میانگین قد دانش آموزان ۱۵ ساله $\bar{x}_1 = 155/81$ سانتی‌متر و انحراف استاندارد قد آنها $s_1 = 8/80$ سانتی‌متر و میانگین وزن آنها $\bar{x}_2 = 44/11$ کیلوگرم و انحراف استاندارد وزن $s_2 = 7/57$ کیلوگرم باشد. ضریب تغییرات این صفت‌ها چنین خواهد بود:

$$\text{قد } C.V. = \frac{8/80}{155/81} \times 100 = 5/6\% \quad \text{وزن } C.V. = \frac{7/57}{44/11} \times 100 = 17\%$$

تفسیر این موضوع چنین است که ضریب تغییرات یا ضریب پراکندگی صفت وزن، بیشتر از ضریب تغییرات قد می‌باشد. یعنی باید انتظار داشت که نوجوانان از نظر وزن، اختلاف بیشتری با هم داشته باشند در صورتی که اختلاف قد آنها نسبت به صفت وزن، کمتر است. ضمناً به یاد داشته باشیم که ضریب تغییرات یک آزمایش از مقدار عددی ۳۰ بیشتر نباشد. در غیر این صورت، آزمایش غیر قابل قبول می‌باشد.

۱-۵ آزمون فرض (Test of Hypothesis)

یک فرض آماری، فرض مربوط به یک پارامتر یعنی اندازه آماری یک جامعه است. در حقیقت حکم یا ادعایی راجع به پارامترهای مجهول جامعه است. برای اینکه با اطمینان کامل صحت و سقم یک فرض آماری را تعیین کنیم، ممکن است لازم باشد که کلیه

افراد جامعه را بررسی نمائیم چون اغلب این امر غیرممکن یا غیر عملی است (گاه جامعه خیلی بزرگ است که بررسی آن هزینه و وقت زیادی لازم دارد که شاید بررسی آن معقول نباشد)، لذا مجبوریم نمونه‌ای از جامعه را انتخاب کرده و نتایج آن را، برای تصمیم در مورد صحت یا عدم صحت فرض، استفاده کنیم. روش آماری با نحوه تصمیم‌گیری که ما را به درک درستی یا نادرستی یک فرض راهنمایی می‌کند، یک آزمون آماری نامیده می‌شود. فرض مورد آزمون را معمولاً با H_0 نشان می‌دهیم و فرض نول گوئیم و در مقابل H_1 موضوعی را قرار می‌دهیم که در صورت رد H_0 جایگزین آن شود که آن را حکم مسئله گفته و با H_1 نشان می‌دهیم. پارامترهای جامعه که از آنها صحبت می‌شود ممکن است که میانگین باشد و یا که واریانس‌ها باشد که فرض صفر تساوی آنها را بیان می‌نماید و با روش خاص، مسئله را حل نموده و فرض صفر را قبول کرده و یا که آن را رد می‌نمائیم. در آزمون یک فرض آماری، موقعی فرض قابل قبول است که احتمال حساب شده از مقدار معین α (بخوانیم آلفا) یعنی سطح معنی‌دار بودن زیادتر باشد. اگر احتمال حساب شده کمتر از α باشد نشان دهنده غلط بودن فرض صفر است و در این صورت نتیجه را معنی‌دار گوئیم. اگر احتمال حساب شده کمتر از $\alpha = 0.05$ باشد نتیجه را معنی‌دار گفته و فرض صفر را رد می‌نماییم. اگر احتمال حساب شده کمتر از $\alpha = 0.01$ باشد، در این صورت آن را بسیار معنی‌دار گفته و مجدداً فرض صفر رد می‌شود. ضمناً فراموش نشود که اگر بخواهیم اطلاعات به‌دست آمده از یک گروه کوچک‌تر را به یک گروه بزرگ‌تر (مثلاً از یک نمونه به یک جامعه)، تعمیم دهیم و به عبارت دیگر از روی اندازه‌های نمونه، اندازه‌های جامعه را تخمین بزنیم روش استقراء یا روش جزء به کل را به کار گرفته‌ایم.

۶-۱ مقایسه دو میانگین

در مطالعه‌های بیولوژیکی همیشه یک سؤال پیش می‌آید که در چه صورت دو نمونه یا دو سری از نتایج به‌دست آمده حقیقتاً متفاوتند؟ به عبارت دیگر اختلاف بین دو میانگین در اثر نوسان‌های تصادفی در نمونه‌گیری بوجود آمده‌اند یا اینکه واقعاً دو میانگین با یک‌دیگر اختلاف معنی‌دار دارند؟ به عبارت دیگر این اختلاف مهم‌تر از آن است که بتوانیم ادعا بکنیم که دو نمونه میانگین واحد دارند و یا از یک جامعه واحد

گرفته شده‌اند؟

۱-۶-۱ مقایسه دو میانگین (نمونه بزرگ با جامعه)

اگر نمونه‌ای در اختیار داشته باشیم تعداد افراد این نمونه n و اعضای آن شامل x_1 و x_2 و ... x_n باشد میانگین و انحراف معیار این نمونه چنین به دست می‌آید:

$$\bar{X} = \frac{\sum X_i}{n} = \text{میانگین نمونه}$$

$$S_x = \sqrt{\frac{\sum (x - \bar{x})^2}{n-1}} \quad \text{به شرطی که } n < 30$$

اگر n بزرگ‌تر از ۳۰ باشد بجای $n-1$ ، عدد n را قرار می‌دهیم. ولی بهترین فرمولی که می‌تواند انحراف معیار نمونه را نشان دهد، عبارت است:

$$S_{\bar{X}} = \frac{S_x}{\sqrt{n}}$$

حال اگر به‌خواهیم بهترین انحراف معیار نمونه را با توجه به پارامتر جامعه نشان دهیم، خواهیم داشت:

$$\sigma_{\bar{x}} = \frac{\sigma_x}{\sqrt{n}}$$

مثال ۱-۴: واحدی از یک مجتمع صنعتی گزارش می‌دهد که کارگران آن واحد به طور متوسط در مدت ۱۳ دقیقه، یک دستگاه را تکمیل نموده و بدین جهت تقاضای افزایش حقوق برای کارگران خود می‌نماید. اگر متوسط مدت لازم برای تکمیل ۴۰۰ دستگاه در یک نمونه ۱۴ دقیقه با انحراف معیار ۱۰ دقیقه باشد، نسبت به ادعای این واحد چه نتیجه‌گیری می‌شود ($\alpha = 0.05$)؟

$$H_0: m = 13 \quad \text{دقیقه}$$

$$H_1: m > 13 \quad \text{دقیقه}$$

$$n = 400$$

$$\bar{x} = 14$$

$$m_{\bar{x}} = 13$$

$$\sigma_x = 10$$

$$\sigma_{\bar{x}} = \frac{\sigma_x}{\sqrt{n}}$$

$$Z = \frac{\bar{x} - m_{\bar{x}}}{\sigma_{\bar{x}}} = \frac{14 - 13}{\frac{10}{\sqrt{400}}} = 2/0 \quad Z = 2/0 \rightarrow p = 0/477$$

$$P(\bar{x} \geq 14) = 0/5000 - 0/4772 = 0/0228$$

چون مقدار p باقیمانده از $\alpha = 5\%$ کوچک‌تر است، پس فرض صفر رد می‌شود. یعنی مدت زمان تکمیل یک دستگاه بیش از ۱۳ دقیقه می‌باشد.

۱-۶-۲ مقایسه دو میانگین با نمونه‌های بزرگ

اگر n_1 افراد نمونه جامعه اول (x) و n_2 افراد نمونه جامعه دوم (y) در نظر گرفته شده و هر دو بزرگ‌تر از عدد ۳۰ باشند، در این صورت می‌توان این دو نمونه را بزرگ در نظر گرفت. جامعه x دارای میانگین m_x و انحراف معیار σ_y و جامعه دوم یا y دارای میانگین m_y و انحراف معیار σ_y می‌باشد، از هر کدام از این دو جامعه، یک نمونه اخذ شده و میانگین هر نمونه تعیین شد و نحوه مقایسه این دو میانگین چنین خواهد بود:

$$Z = \frac{(\bar{x} - \bar{y}) - m_{\bar{x} - \bar{y}}}{S_{\bar{x} - \bar{y}}}$$

که در اینجا $\bar{x}$ میانگین نمونه اخذ شده از جامعه x و $\bar{y}$ میانگین نمونه اخذ شده از جامعه y می‌باشد و $S_{\bar{x} - \bar{y}}$ انحراف معیار تفاوت دو میانگین است که چنین به دست می‌آید:

$$S_{\bar{x} - \bar{y}} = \sqrt{S_{\bar{x}}^2 + S_{\bar{y}}^2} \quad \text{و} \quad S_{\bar{x}} = \frac{\sigma_x}{\sqrt{n_1}}, \quad S_{\bar{y}} = \frac{\sigma_y}{\sqrt{n_2}}$$

ضمناً در این گونه مقایسه‌ها فرض صفر و حکم مسئله چنین خواهد بود:

$$\text{فرض } H_0: m_x - m_y = m_{\bar{x} - \bar{y}} = 0 \quad m_{\bar{x}} = m_{\bar{y}}$$

$$\text{حکم } H_1: m_x - m_y = m_{\bar{x} - \bar{y}} \neq 0 \quad m_{\bar{x}} \neq m_{\bar{y}}$$

و در پایان میزان احتمال به دست آمده یا p را با α مقایسه می‌کنیم اگر $P > \alpha$ باشد نتیجه بی‌معنی و قبول فرض صفر و در غیر این صورت نتیجه معنی‌دار بوده و فرض صفر را رد می‌نمائیم.

ضمناً $S_{\bar{x}}$ و $S_{\bar{y}}$ از روابط زیر نیز به دست می‌آیند:

$$S_{\bar{x}} = \sqrt{\frac{\sum (x - \bar{x})^2}{n_1(n_1 - 1)}} = \sqrt{\frac{S_x^2}{n}} = \frac{S_x}{\sqrt{n}} \quad S_{\bar{y}} = \sqrt{\frac{\sum (y - \bar{y})^2}{n_2(n_2 - 1)}} = \sqrt{\frac{S_y^2}{n}} = \frac{S_y}{\sqrt{n}}$$

مثال ۱-۵: دو نوع لامپ الکتریکی از نظر دوام مورد بررسی قرار گرفته و نتایج زیر به دست آمده است:

نوع ۲	نوع ۱
تعداد نمونه ۶۴	تعداد نمونه ۴۶
میانگین نمونه ۱۰۴۱ ساعت	میانگین نمونه ۱۰۷۰ ساعت
$\Sigma (y - \bar{y})^2 = 23200$	$\Sigma (x - \bar{x})^2 = 21000$

آیا می توانیم تفاوت متوسط دوام دو نوع لامپ را در سطح ۰/۰۱ متفاوت بدانیم؟
 حل: نوع ۱ را x و نوع دوم را y در نظر گرفته و چنین ادامه می دهیم: (البته چون نمونه ها بزرگ است در مخرج بجای $n-1$ ، می توانیم مقدار n را قرار دهیم:

$$S_{\bar{x}} = \sqrt{\frac{\Sigma (x - \bar{x})^2}{n_1 (n_1)}} = \sqrt{\frac{21000}{46 (46)}} = 3/15$$

$$S_{\bar{y}} = \sqrt{\frac{\Sigma (y - \bar{y})^2}{n_2 (n_2)}} = \sqrt{\frac{23200}{64 (64)}} = 2/38$$

$$S_{\bar{x} - \bar{y}} = \sqrt{S_{\bar{x}}^2 + S_{\bar{y}}^2} = \sqrt{9/92 + 5/66} = 3/94$$

$$Z = \frac{(\bar{x} - \bar{y}) - m_{\bar{x} - \bar{y}}}{S_{\bar{x} - \bar{y}}} = \frac{(1070 - 1041) - (0)}{3/94} = 7/36$$

با استفاده از جدول Z (جدول شماره ۱ پیوست) مقدار P را به دست می آوریم و چون مسئله دو دانه است (حکم مسئله چنین نشان می دهد)، مقدار P را در دو ضرب می نمائیم تا P کل به دست آید.

$$H_0 : m_{\bar{x} - \bar{y}} = 0$$

$$H_1 : m_{\bar{x} - \bar{y}} \neq 0$$

$$P \text{ کل} = (0/5 - 0/5) \times 2 = 0$$

و P کل در مقایسه با $\alpha = 1\%$ کوچک تر بوده که یعنی نتیجه معنی دار بوده و فرض صفر رد می شود یعنی دو نمونه لامپ با یک دیگر تفاوت داشته و لامپ نوع ۱ بهتر از لامپ نوع ۲ می باشد. این موضوع را با توجه به میانگین دو نمونه نتیجه گیری می نمائیم.

البته نتیجه را می‌توان چنین نیز بیان نمود که با اطلاعات به‌دست آمده از مسئله نمی‌توان دو میانگین را یکسان تلقی نمود.

۳-۶-۱ مقایسه دو میانگین با نمونه‌های کوچک

وقتی که اندازه یا تعداد داده‌ها در یکی یا هر دو نمونه کمتر از ۳۰ باشد، روش توزیع نرمال دارای دقت خوبی برای مقایسه دو نمونه نخواهد بود. برای این‌که برآورد واریانس از طریق دو نمونه جداگانه، دقت کمتری همراه دارد که البته این در حالتی است که واریانس دو نمونه مشابه باشد. در این حالت با استفاده از توزیع t استیودنت می‌بایست که مسئله را پاسخ داد و البته در نمونه‌های کوچک انحراف معیار با $\hat{S}$ (بخوانیم S هت) نشان داده می‌شود و نمونه‌های بزرگ را با S و انحراف معیار جامعه را با σ نشان می‌دهیم. برای محاسبه S می‌توان از فرمول زیر استفاده کرد:

$$S = \sqrt{\frac{n}{n-1}} \times \hat{S}$$

حال اگر $\hat{S}$ را نداشته باشیم و بخواهیم مقدار S یا انحراف معیار واقعی جامعه را

$$\text{محاسبه کنیم از فرمول } S = \sqrt{\frac{\sum (x - \bar{x})^2}{n-1}} \text{ استفاده می‌کنیم که در این فرمول داریم:}$$

$$\sum (x - \bar{x})^2 = SS_x$$

و اگر بخواهیم بهترین برآورد معیار انحراف میانگین نمونه‌هایی با n متغیر را به‌دست آوریم آن را با $S_{\bar{x}}$ نشان می‌دهیم که در اینجا نیز $\sigma = S$ منظور می‌شود.

$$S_{\bar{x}} = \frac{S}{\sqrt{n}}$$

برای حل مسایل t استیودنت، پس از محاسبه t از طریق فرمول، آن را با t جدول که بر اساس درجه آزادی $df = n - 1$ و سطح α درصد از جدول t استیودنت مخصوص استخراج می‌شود، مقایسه کرده و در صورتی که t محاسبه‌ای بزرگ‌تر از t جدول باشد نتیجه را معنی‌دار اعلام کرده و فرض صفر رد می‌شود. محاسبه و مقایسه میانگین‌ها، حالت‌های مختلفی دارد که عبارتند از:

۱-۳-۶-۱ مقایسه دو میانگین (نمونه کوچک با جامعه)

اگر متغیر x دارای توزیع نرمال با میانگین m باشد و اگر تمام نمونه‌های n تایی ممکن را از این جامعه انتخاب و میانگین آنها را با $\bar{x}$ نشان دهیم در این صورت کمیت $t = \frac{\bar{x} - m}{S_{\bar{x}}}$ دارای توزیع t استیودنت با درجه آزادی $df = n - 1$ و $S_{\bar{x}} = \frac{S}{\sqrt{n}}$ می‌باشد.

مثال ۱-۶: در یک نمونه ۱۶ تایی مقدار میانگین ۲۸/۰ و مقدار انحراف معیار ۳/۰ است آیا می‌توانیم در سطح ۰/۰۵ این فرض را که میانگین جامعه ۳۰ می‌باشد را رد نمائیم؟

$$H_0: m = 30/0$$

$$H_1: m \neq 30/0$$

$$n = 16 \text{ و } \hat{S} = 3/0, \bar{x} = 28/0, m = 30$$

$$S = \sqrt{\frac{n}{n-1}} \times \hat{S} = \sqrt{\frac{16}{15}} \times 3/0 = \frac{12}{\sqrt{15}}$$

$$S_{\bar{x}} = \frac{s}{\sqrt{n}} = \frac{\frac{12}{\sqrt{15}}}{\sqrt{16}} = 0/78 \quad t = \frac{\bar{x} - m}{S_{\bar{x}}} = \frac{28/0 - 30/0}{0/78} = -2/56, t(0/05, 15) = 2/131$$

چون t محاسبه‌ای از t جدول بزرگ‌تر می‌باشد پس نتیجه معنی‌دار بوده و فرض صفر را رد می‌نمائیم. به عبارت دیگر جامعه ما دارای میانگین ۳۰ نمی‌باشد.

۱-۳-۶-۲ حالتی که واریانس‌های دو جامعه مساوی باشند (همگن بودن واریانس‌ها)

در این حالت فرض می‌نمائیم که دو جامعه با واریانس‌های مساوی داریم و از هر کدام از این دو جامعه نمونه‌ای گرفته و تفاوت دو نمونه را مورد ارزیابی قرار می‌دهیم در اینجا ابتدا یک واریانس مشترک به دست می‌آوریم که فرمول محاسبه آن چنین است:

$$S_p^2 = S^2_{\text{مشترک}} = \frac{\sum (x - \bar{x})^2 + \sum (y - \bar{y})^2}{n_1 + n_2 - 2}$$

در این فرمول n_1 تعداد نمونه اخذ شده از جامعه x و n_2 تعداد نمونه اخذ شده از جامعه y می‌باشد. سپس با استفاده از کمیت t با فرمول زیر مقدار t را به دست می‌آوریم:

$$t = \frac{(\bar{x} - \bar{y}) - m_{\bar{x} - \bar{y}}}{S_{\bar{x} - \bar{y}}}$$

که در آن:

$$S_{\bar{x}} = \frac{S}{\sqrt{n_1}}, \quad S_{\bar{y}} = \frac{S}{\sqrt{n_2}}, \quad S_{\bar{x} - \bar{y}} = \sqrt{S_{\bar{x}}^2 + S_{\bar{y}}^2}$$

و البته فرض صفر، $H_0: m_{\bar{x}} - m_{\bar{y}} = 0$ و حکم مسئله، $H_1: m_{\bar{x}} \neq m_{\bar{y}}$ خواهد بود.

مثال ۷-۱: نتیجه کنترل کیفی در دو روش تهیه یک محصول به صورت زیر بوده است آیا این داده‌ها می‌تواند بر تفاوت کیفیت محصول در دو روش دلالت نماید؟ (α) را برابر ۱٪ در نظر بگیرید.

$$\begin{array}{l} \text{روش اول} \\ \text{روش دوم} \end{array} \quad \left| \begin{array}{l} H_0: m_{\bar{x}} = m_{\bar{y}} \\ H_1: m_{\bar{x}} \neq m_{\bar{y}} \end{array} \right.$$

$$\bar{x} = \frac{\sum x}{n_1} = \frac{10}{4} = 2.5 \quad \bar{y} = \frac{\sum y}{n_2} = \frac{18}{3} = 6$$

سپس جدول زیر را تشکیل می‌دهیم:

روش اول (x)	روش دوم (y)	$(x - \bar{x})$	$(x - \bar{x})^2$	$(y - \bar{y})$	$(y - \bar{y})^2$
۱/۵	۲/۵	-۱	۱	-۰/۵	۰/۲۵
۲/۵	۳	۰	۰	۰	۰
۳/۵	۳	۱	۱	۰	۰
۲/۵	۴	۰	۰	۱	۱
$\sum x = 10$	۳/۵		$SS_x = 2$	۰/۵	۰/۲۵
	۲			-۱	۱
	$\sum y = 18$				$SS_y = 2.5$

$$S = \sqrt{\frac{\sum (x - \bar{x})^2 + \sum (y - \bar{y})^2}{n_1 + n_2 - 2}} = \sqrt{\frac{2 + 2.5}{4 + 3 - 2}} = 0.75$$

$$S_{\bar{x}} = \frac{S}{\sqrt{n}} = \frac{0.75}{\sqrt{4}} \quad S_{\bar{y}} = \frac{0.75}{\sqrt{6}}$$

$$S_{\bar{x} - \bar{y}} = \sqrt{S_{\bar{x}}^2 + S_{\bar{y}}^2} = \sqrt{\left(\frac{0.75}{\sqrt{4}}\right)^2 + \left(\frac{0.75}{\sqrt{6}}\right)^2} = 0.48$$

$$t = \frac{(\bar{x} - \bar{y}) - m_{\bar{x}} - m_{\bar{y}}}{S_{\bar{x} - \bar{y}}} = \frac{(2/5 - 3) - 0}{0.48} = -1/0.4$$

اکنون t جدول را بر اساس 0.01 و درجه آزادی $df = n_1 + n_2 - 2 = 8$ از جدول t استیودنت (جدول شماره ۲ پیوست) استخراج می‌کنیم که برابر $t = 3/36$ می‌باشد. پس چون t محاسبه‌ای از t جدول کوچک‌تر می‌باشد، لذا نتیجه بی‌معنی است و فرض صفر قبول می‌گردد یعنی کیفیت در دو روش تهیه محصول بایکدیگر متفاوت نیست.

۱-۳-۳ حالتی که واریانس‌های دو جامعه مشابه نمی‌باشد (همگن نبودن واریانس‌ها)

در این حالت جامعه‌هایی که از آنها نمونه گرفته می‌شود دارای واریانس‌های یکسان نیستند در این مورد برای هر یک از نمونه‌ها یک واریانس محاسبه نموده و سپس واریانس تفاوت دو میانگین را به دست آورده و آنرا در فرمول t به صورت زیر می‌گذاریم:

$$t = \frac{(\bar{x} - \bar{y}) - m_{\bar{x}} - m_{\bar{y}}}{S_{\bar{x} - \bar{y}}}$$

که در این فرمول داریم:

$$S_{\bar{x} - \bar{y}} = \sqrt{S_{\bar{x}}^2 + S_{\bar{y}}^2}$$

$$S_{\bar{x}} = \frac{S_x}{\sqrt{n_1}} \quad , \quad S_x = \sqrt{\frac{SS_x}{n_1 - 1}} \quad , \quad S_{\bar{y}} = \frac{S_y}{\sqrt{n_2}} \quad , \quad S_y = \sqrt{\frac{SS_y}{n_2 - 1}}$$

پس از آن برای هر کدام درجه‌های آزادی $df_1 = n_1 - 1$ و $df_2 = n_2 - 1$ دو t به نام‌های t_1 و t_2 از جدول t استیودنت استخراج نموده، سپس t محاسبه شده را با دو t جدول (t_1 و t_2) مقایسه می‌کنیم. اگر t محاسبه شده از هر دو t بزرگ‌تر باشد آزمون ما معنی‌دار و فرض صفر را قبول می‌کنیم، ولی اگر t محاسبه شده از هر دو t جدول کوچک‌تر باشد آزمون ما معنی‌دار نبوده و فرض صفر را قبول می‌کنیم، ولی اگر t محاسبه شده مابین دو