



کاربرد کامپیوتر در علوم تربیتی (Spss ۱۳)

(رشته علوم تربیتی)

محمدحسن صیف دکتر محمدرضا سرمدی

سرشناسه	صیف، محمدحسن
عنوان و نام پدیدآور	کاربرد کامپیوتر در علوم تربیتی (۱۳ Spss) (رشته علوم تربیتی) / محمدحسن صیف، محمدرضا سرمدی.
مشخصات نشر	تهران: دانشگاه پیام نور، ۱۳۸۷.
مشخصات ظاهری	نه، ۱۶۵ص: مصور، جدول، نمودار.
فروست	دانشگاه پیام نور؛ ۱۳۷۶. گروه علوم تربیتی؛ ۵۴/ک.
شابک	978-964-387-485-8
وضعیت فهرست نویسی	فیپا
یادداشت	واژه نامه .
یادداشت	کتابنامه : ص. ۱۶۵.
موضوع	اس. پی. اس. اس تحت ویندوز
موضوع	علوم اجتماعی -- روش های آماری -- برنامه های کامپیوتری
موضوع	علوم اجتماعی -- روش های آماری -- آموزش برنامه ای
شناسه افزوده	سرمدی، محمدرضا، ۱۳۴۴ -
شناسه افزوده	دانشگاه پیام نور
شناسه افزوده	Payam Noor University:
رده بندی کنگره	۱۳۸۷ ک۹ص/HA۳۲
رده بندی دیویی	۳۰۰/۲۸۵۵۳۶:
شماره کتابشناسی ملی	۱۲۵۲۴۳۵:



دانشگاه پیام نور

کاربرد کامپیوتر در علوم تربیتی (۱۳ Spss)

دکتر محمدحسن صیف دکتر محمدرضا سرمدی

ویراستار علمی: مهندس جعفر تنها

حروفچینی و نمونه خوانی: مدیریت تولید مواد و تجهیزات آموزشی

طراح جلد: آتنا عسگری فر

لیتوگرافی، چاپ و صحافی: انتشارات دانشگاه پیام نور

تعداد: ***

چاپ.....۱۳۸۷

قیمت: ***

کلیه حقوق برای دانشگاه پیام نور محفوظ است.

بسم الله الرحمن الرحيم

پیشگفتار ناشر

کتاب‌های دانشگاه پیام نور حسب مورد و با توجه به شرایط مختلف یک درس در یک یا چند رشته دانشگاهی، به صورت کتاب درسی، متن آزمایشگاهی، فرادرسی، و کمک‌درسی چاپ می‌شوند.

کتاب درسی ثمرهٔ کوشش‌های علمی صاحب اثر است که براساس نیازهای درسی دانشجویان و سرفصل‌های مصوب تهیه و پس از داوری علمی، طراحی آموزشی، و ویرایش علمی در گروه‌های علمی و آموزشی، به چاپ می‌رسد. پس از چاپ ویرایش اول اثر، با نظرخواهی‌ها و داوری علمی مجدد و با دریافت نظرهای اصلاحی و متناسب با پیشرفت علوم و فناوری، صاحب اثر در کتاب تجدیدنظر می‌کند و ویرایش جدید کتاب با اعمال ویرایش زبانی و صوری جدید چاپ می‌شود.

متن آزمایشگاهی (م) راهنمایی است که دانشجویان با استفاده از آن و کمک استاد، کارهای عملی و آزمایشگاهی را انجام می‌دهند.

کتاب‌های فرادرسی (ف) و **کمک‌درسی** (ک) به منظور غنی‌تر کردن منابع درسی دانشگاهی تهیه و بر روی لوح فشرده تکثیر می‌شوند و یا در وبگاه دانشگاه قرار می‌گیرند.

مدیریت تولید مواد و تجهیزات آموزشی

فهرست

نه	پیشگفتار
۱	فصل اول. مفاهیم اساسی آمار
۱	مقدمه ۱
۲	جایگاه Spss در تجزیه و تحلیل اطلاعات ۲
۲	تاریخچه نرم افزار Spss
۳	آمار توصیفی و استنباطی
۴	داده‌های آماری
۵	جامعه آماری
۵	نمونه
۶	متغیر
۷	مقیاس‌های اندازه‌گیری
۷	الف) مقیاس اسمی
۷	ب) مقیاس رتبه‌ای
۸	ج) مقیاس فاصله‌ای
۹	د) مقیاس نسبی (نسبتی)
۱۰	آزمون فرضیه
۱۱	خطای نوع اول و دوم
۱۲	آزمون‌های یک دامنه‌ای و دو دامنه‌ای
۱۵	تفسیر تأیید یا رد فرض صفر
۱۷	فصل دوم. نحوه تعریف و به‌کارگیری متغیرها در Spss
۲۰	نام متغیر (name)
۲۱	نوع متغیر (Type)
۲۳	تعداد ارقام صحیح متغیر (Width)
۲۳	تعداد ارقام اعشار متغیر (Decimals)
۲۴	بر چسب متغیر (Label)
۲۴	کد گذاری مقوله‌های متغیر (Values)
۲۶	مقادیر نامشخص متغیر (Missing)
۲۷	اندازه ستون نمایش دهنده متغیر (Columns)
۲۸	چگونگی نمایش هر ستون (Align)
۲۸	مقیاس اندازه‌گیری متغیر (Measure)

۲۹	تغییر در داده‌ها (compute)
۳۰	تبدیل داده‌های کمی به کیفی (Recode)
۳۳	انتخاب داده‌ها (select cases)
۴۱	فصل سوم. توزیع فراوانی و ترسیم نمودارها
۴۱	توزیع فراوانی
۴۴	فراوانی مطلق
۴۴	فراوانی تراکمی درصدی
۴۵	فراوانی نسبی درصدی
۴۷	نمودار فراوانی
۴۷	الف) نمودار هیستوگرام
۵۰	ب) نمودار ستونی
۵۱	ج) نمودار دایره‌ای
۵۵	فصل چهارم. شاخص‌های مرکزی
۵۵	شاخص‌های مرکزی
۵۶	نما
۵۸	میان
۶۵	میانگین
۶۲	مقایسه شاخص‌های مرکزی و کاربرد آن‌ها
۶۹	فصل پنجم. شاخص‌های پراکندگی
۶۹	شاخص‌های پراکندگی
۷۰	دامنه تغییرات
۷۱	چارک
۷۴	واریانس و انحراف استاندارد (معیار)
۷۶	تفسیر انحراف استاندارد
۷۷	رتبه درصدی و دهک
۸۰	ضریب تغییرپذیری
۸۲	کجی و کشیدگی
۸۳	الف) کجی
۸۶	ب) کشیدگی
۸۹	فصل ششم. پیش‌فرض‌های آزمون پارامتریک و ناپارامتریک
۸۹	پیش‌فرض‌های آزمون پارامتریک و ناپارامتریک
۹۰	نرمال بودن توزیع جامعه
۹۲	استقلال داده‌ها (تصادفی بودن داده‌ها)
۹۵	فصل هفتم. آزمون‌های آماری پارامتریک
۹۵	آزمون‌های آماری پارامتریک
۹۶	آزمون t یک نمونه
۹۶	فرض‌ها
۹۹	آزمون t برای گروه‌های وابسته
۱۰۰	پیش‌نیازهای طرح
۱۰۰	فرض‌ها
۱۰۲	آزمون t دو نمونه مستقل

۱۰۳	پیش‌نیازهای طرح
۱۰۳	فرض‌ها
۱۰۷	تحلیل واریانس
۱۰۷	پیش‌نیازهای طرح
۱۰۸	فرض‌ها
۱۱۱	مقایسه‌های بعد از تجربه
۱۱۲	آزمون شفه
۱۱۳	آزمون توکی (HSD)
۱۱۹	فصل هشتم. آزمون‌های آماری ناپارامتریک
۱۱۹	آزمون‌های آماری ناپارامتریک
۱۲۰	محاسن آزمون‌های آماری ناپارامتریک
۱۲۱	معایب آزمون‌های ناپارامتریک
۱۲۲	آزمون ناپارامتریک کای دو (Chi-Square)
۱۲۲	کاربرد
۱۲۲	معایب
۱۲۵	آزمون ناپارامتریک مربوط به دو گروه مستقل (2 Independent samples)
۱۲۵	کاربرد
۱۲۶	محاسن
۱۲۶	معایب
۱۳۰	آزمون ناپارامتریک مربوط به دو گروه وابسته (2 Related samples)
۱۳۰	کاربرد
۱۳۱	محاسن
۱۳۱	معایب
۱۳۴	آزمون ناپارامتریک مربوط به k گروه مستقل (K Independent samples)
۱۳۴	کاربرد
۱۳۴	محاسن
۱۳۴	معایب
۱۳۸	مقایسه‌های بعد از تجربه زوجی (یک‌به‌یک)
۱۳۹	آزمون ناپارامتریک مربوط به K گروه وابسته (K Related samples)
۱۳۹	کاربرد
۱۳۹	محاسن
۱۴۰	معایب
۱۴۷	فصل نهم. آزمون‌های همبستگی
۱۴۷	مفهوم همبستگی
۱۴۹	نمودار پراکنش
۱۵۳	محاسبه ضریب همبستگی
۱۵۳	ضریب همبستگی پیرسون
۱۵۶	ضریب همبستگی اسپیرمن
۱۵۸	همبستگی جزئی
۱۶۵	مراجع

پیشگفتار

با توجه به نقش کلیدی رایانه در زندگی بشر و اهمیت آن، امروزه تمامی جوانب دنیای بشری وابسته به رایانه است و بدون در نظر گرفتن آن، نمی‌توان زندگی بشرا روزی را تصور نمود. یکی از اساسی‌ترین وجوه بشری علوم مختلف هستند که موجبات آسایش و رفاه انسانی را فراهم می‌کنند. امروزه می‌توان رایانه‌ها را وجه مشترک تمامی علوم در نظر گرفت و به استفاده گسترده در این امر پی برد در تحلیل‌های آماری در علوم مختلف از نرم‌افزارهای زیادی استفاده می‌شود که بدلیل کاربرد وسیع نرم‌افزار Spss در تحلیل داده‌های علوم مختلف و بروز شدن این نرم‌افزار با توجه به پیچیده‌تر شدن علوم مختلف، پرداختن به کاربرد این نرم‌افزار امری بسیار مهم تلقی می‌شود. به دلیل کمبود منابع لازم درخصوص آشنایی با این نرم‌افزار در علوم تربیتی و روان‌شناسی، مولفان این اثر باهدف آشنایی بیشتر دانشجویان و محققان با نرم‌افزار فوق و تحلیل بهینه داده‌های جمع‌آوری شده، کتاب حاضر را جهت تدریس در دانشگاه پیام نور تدوین نمودند. به دلیل ساختار آموزش از راه دور این دانشگاه، در تألیف این کتاب سعی شده است مطالب قابل فهم‌تر و به بیان ساده مطرح شود و دانشجویان بتوانند شیوه‌های تحلیل آماری را به نحو مقتضی به کار گیرند. توصیه می‌شود دانشجویان قبل از مطالعه این کتاب با مهارت‌های ICDL آشنایی داشته باشند تا بتواند راحت‌تر از این کتاب بهره بگیرد.

مهم‌ترین ویژگی‌های این کتاب عبارتند از:

۱. تبعیت مطالب کتاب حاضر از مطالب کتاب‌های مرجع آمار و روش تحقیق
۲. ارائه مثال‌های کاربردی ساده و قابل درک برای دانشجویان و محققان

۳. ارائه مدل‌های واقعی از نرم‌افزار فوق جهت درک هر چه بیشتر مراحل انجام کار
۴. ارائه مبانی نظری هر یک از روش‌های مورد استفاده در نرم‌افزار
۵. محتوای خودخوان کتاب حاضر و فرض این کتاب بر هیچ‌گونه آشنایی قبلی دانشجو با نرم‌افزار فوق
۶. استفاده از آخرین ویرایش نرم‌افزار Spss.

آدرس پست الکترونیک مولفین:

hassanseif@gmail.com
sarmadi@pnu.ac.ir

فصل اول

مفاهیم اساسی آمار

هدف‌های یادگیری

از دانشجو انتظار می‌رود که پس از خواندن فصل اول بتواند:

۱. مفهوم آمار، موارد و کاربرد آن را شرح دهد.
۲. فرق بین جامعه و نمونه، پارامتر و آماره را تشخیص دهد.
۳. متغیر را تعریف و انواع آن را نام ببرد.
۴. انواع مقیاس‌های اندازه‌گیری را نام برده، ویژگی‌ها و موارد کاربرد آن‌ها را نام ببرد.
۵. فرضیه و انواع آن را تعریف کند.
۶. خطای نوع اول و دوم را تعریف کند.
۷. توان آزمون و عوامل مؤثر بر آن را شرح دهد.
۸. آزمون‌های یک دامنه‌ای و دو دامنه‌ای را تعریف و تفاوت‌های آن را بیان کند.

مقدمه

روش‌های آماری در همه علوم به‌عنوان وسیله‌ای برای تنظیم، تحلیل و تفسیر داده‌ها یکی از مبانی اساسی را تشکیل می‌دهد. این روش‌ها و فنون در واقع پایه دستوری برای استنباط و استقرا بوده و قابل تکرار و تأیید در پژوهش‌های علمی به‌شمار می‌روند. آمار از واژه انگلیسی statistics است که این واژه از ریشه یونانی status به معنی ارقام و اطلاعات عددی مربوط به وضع اجتماعی و اقتصادی کشور که برای اداره حکومت لازم و مفیدند، گرفته شده است. (نجفی زند و پاشا شریفی، ۱۳۷۵: ۱۱) اولین جلوه آمار در گذشته، در امر مالیات‌گیری و خراج بستن بر دارایی‌ها و امور نظامی و ارتشی و پس

از آن در طبابت ظاهر گردیده است. امروزه آمار اصول و روش جمع‌آوری اطلاعات اولیه، مرتب کردن و خلاصه کردن و نمایش دادن آن‌ها و بالاخره تجزیه و تحلیل اطلاعات اولیه و استخراج نتایج را مورد بحث قرار می‌دهد. (حسینی، ۱۳۸۲: ۳)

جایگاه Spss در تجزیه و تحلیل اطلاعات

جایگاه Spss در تحقیقات علمی روز به روز کار برد بیشتری پیدا کرده است. این وسیله کار تحقیق را ساده کرده است، صرفه‌جویی در امر نیروی انسانی، هزینه‌ها و زمان باعث شده تا استفاده از Spss در تحقیق جایگزین روش‌های دستی شود و زمان اختصاص یافته برای تجزیه و تحلیل اطلاعات را به‌طور عجیبی کوتاه نماید. در مرحله تجزیه و تحلیل، Spss وسیله بسیار ارزشمندی است این وسیله می‌تواند محاسبات فراوان و بسیار سنگین و پیچیده را در زمان بسیار کوتاهی انجام دهد. توضیح اینکه بعضی محاسبات پیچیده با حجم بالای داده را نمی‌توان با دست انجام داد ولی Spss مشکل چنین محاسباتی را به راحتی حل می‌نماید. برای تجزیه و تحلیل داده‌ها به‌وسیله Spss محقق نیاز به برنامه Spssی و نرم‌افزار مناسب دارد. (کرلینجر ترجمه پاشا شریفی و نجفی زند، ۱۳۷۶: ۵۰۳) در حال حاضر مناسب‌ترین نرم‌افزار در تحقیقات علوم اجتماعی و انسانی، نرم‌افزار Spss^۱ است. و استفاده از آن نیز چندان مشکل نیست. برنامه Spss قابلیت‌های بسیاری دارد. نظیر: تهیه جداول توزیع فراوانی، تهیه لیست داده‌ها، تهیه جداول تقاطعی دو بعدی و چند بعدی، انجام بررسی‌های آمار توصیفی (شاخص‌های مرکزی و پراکندگی)، انجام بررسی‌های آمار استنباطی شامل همبستگی‌ها، کای ۲، انواع تست‌ها و آزمون‌های پارامتریک و ناپارامتریک، محاسبات ریاضی، رگرسیون، تغییر، اصلاح، جابجایی و مرتب کردن داده‌ها.

تاریخچه نرم‌افزار Spss

در سال ۱۹۶۸ سه جوان مبتکر نیل، هال و بنت اقدام به خلق نرم‌افزاری نمودند که بر اساس آن داده‌های خام آماری به اطلاعات بنیادی برای تصمیم‌گیری تبدیل می‌شد این تحول در سیستم نرم‌افزار آماری Spss نام گرفت. اولین کار صورت گرفته در زمینه

Spss در دانشگاه استنفورد صورت گرفت که هدف اولیه آن رفع نیازهای محلی و منطقه‌ای بود تا پخش و گسترش بین‌المللی آن. نیل جامعه شناس و دانشجوی دکترای دانشگاه استنفورد تدارکات لازم را در این زمینه مهیا کرد، بنت دانشجوی دکترای تحقیق عملیاتی به تحلیل تخصصی و طراحی ساختار سیستم Spss پرداخت و هال نیز به برنامه‌ریزی نرم‌افزار مورد نظر پرداخت. در سال ۱۹۶۹ به مرکز ملی تحقیقات افکارسنجی دانشگاه شیکاگو پیوست. دانشگاه شیکاگو Spss را ابزاری سودمند تشخیص داد و او را تشویق نمود تا به توسعه نرم‌افزاری آن اقدام نماید. بعد از مدتی هال نیز به او پیوست و موفق به ایجاد نرم‌افزار کاربردی جامع جهت تجزیه و تحلیل داده‌های آماری شدند. موسسه مک‌گرو هیل^۱ اولین نسخه نرم‌افزاری Spss را برای کاربران منتشر کرد. این نرم‌افزار در کتاب‌فروشی‌های دانشگاه‌های آمریکا در معرض فروش قرار گرفت و نیاز برای استفاده از این نرم‌افزار روزبه‌روز افزایش یافت. با فروش بالای نرم‌افزار Spss مرکز ملی تحقیقات افکار سنجی دانشگاه شیکاگو تصمیم به تأسیس شرکتی در خصوص این نرم‌افزار در سال ۱۹۷۱ گرفت. در دهه ۱۹۷۰ نرم‌افزار Spss رشد چشمگیری و فراوانی داشت به‌طوری‌که سازمان فضایی ناسا برنامه‌های زمانبندی پرتاب شاتل‌های خود را با استفاده از این نرم‌افزار محاسبه می‌نمود. در دهه ۱۹۸۰ اولین بسته نرم‌افزاری کامپیوتر شخصی Spss ویرایش گردید و در سال ۱۹۹۲ برای اولین بار نرم‌افزار Spss ماکروسافت ویندوز ارائه گردید و تا امروزه گسترش چشمگیری در سراسر دنیا داشته است. (به نقل از سایت Spss، ۲۰۰۷)

آمار توصیفی و استنباطی

هنگامی که توده‌ای از اطلاعات برای تفسیر گردآوری می‌شود، ابتدا سازمان‌بندی و خلاصه کردن آن‌ها به طریقی که به‌صورت معنی‌داری قابل فهم و ارتباط باشد، ضروری است، روش‌های آمار توصیفی به همین منظور به‌کار برده می‌شوند در این تجزیه و تحلیل، پژوهشگر داده‌های جمع‌آوری شده را با استفاده از شاخص‌های آمار توصیفی خلاصه و طبقه‌بندی می‌کند به عبارت دیگر، در تجزیه و تحلیل توصیفی، پژوهشگر

ابتدا داده‌های جمع‌آوری شده را با تهیه جدول توزیع فراوانی خلاصه می‌کند و سپس به کمک نمودار آن‌ها را نمایش می‌دهد. (دلاور، ۱۳۸۰: ۶)

نقش توصیف آماری در واقع جمع‌آوری خلاصه کردن و توصیف اطلاعات کمی به‌دست آمده از نمونه‌ها یا جامعه‌ها است. ولی محقق معمولاً کار خود را با توصیف اطلاعات پایان نمی‌دهد بلکه سعی می‌کند آنچه را که از بررسی گروه نمونه به‌دست آورده است به گروه‌های مشابه بزرگ‌تر تعمیم دهد.

تئوری‌های روان‌شناسی یا جامعه‌شناسی از طریق تعمیم نتایج یک یا چند مطالعه به آنچه که ممکن است در مورد کل افراد جامعه صادق باشد به‌وجود می‌آیند. مربیان تعلیم و تربیت با استفاده از نتایج یک مطالعه به‌خصوص با یک سری مطالعات روی گروه‌های نمونه کوچک و تعمیم آن‌ها به کل جامعه دانش‌آموزان نظریات خود را در زمینه اصلاح تعلیم و تربیت بیان می‌دارند.

محقق از میان افراد جامعه مورد نظر از طریق انتخاب تصادفی یک گروه نمونه را انتخاب می‌کند. رفتار افراد گروه نمونه مورد مشاهده قرار می‌گیرد. بر اساس رفتارهای مشاهده شده آنچه که احتمالاً واقعیت جامعه مورد نظر می‌باشد استنباط می‌گردد. این شیوه اساس استنباط آماری^۱ (آمار استنباطی) را تشکیل می‌دهد. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۱)

داده‌های آماری

اطلاعاتی که با شیوه‌های گوناگون از قبیل اندازه‌گیری، آزمایش، مشاهده، مصاحبه و مانند آن‌ها به‌صورت کمی یا کیفی به‌دست می‌آید، داده‌های آماری نامیده می‌شوند. در برخی موارد داده‌ها را از راه شمارش تمامی افراد یک مجموعه که جامعه آماری تلقی می‌شود، به‌دست می‌آید که به این امر سرشماری یا آمارگیری گفته می‌شود. اما در بسیاری از موارد سرشماری تمام افراد یک جامعه (نظیر کلیه دانش‌آموزان سال اول راهنمایی ایران) مقدور نیست، در این حالت از نمونه‌گیری یا نمونه‌برداری استفاده می‌شود. (جعفر نجفی زند، پاشا شریفی، ۱۳۸۵: ۱۴)

جامعه آماری

جامعه^۱ عبارت است گروهی از افراد، اشیا یا حوادث که حداقل دارای یک صفت یا ویژگی مشترک هستند و تنها به گروهی از افراد اطلاق نمی‌شود مانند تمام دانشجویانی که در این ترم در دانشکده ثبت نام کرده‌اند.

در پژوهش، مفهوم جامعه آماری به کلیه افرادی اطلاق می‌شود که عمل تعمیم‌پذیری در مورد آنان صورت می‌گیرد به عبارت دیگر منظور از جامعه آماری همه اعضای واقعی یا فرضی که علاقمند هستیم یافته‌های پژوهش را به آن‌ها تعمیم دهیم. (دلاور، ۱۳۸۰: ۸)

جامعه آماری می‌تواند محدود باشد و آن جامعه آماری است که تعداد اعضای تشکیل دهنده آن معین و محدود است مانند مجموعه دارندگان دیپلم متوسطه در یک شهر در یک سال معین. یا اینکه جامعه آماری ممکن است نامحدود باشد و آن در صورتی است مشخص کردن همه افراد یا اجزای تشکیل دهنده مجموعه میسر نباشد در چنین مواقعی، تعیین و شناسایی یک جامعه مناسب مستلزم صرف وقت و هزینه است. به همین دلیل عده‌ای سعی می‌کنند این عمل را از طریقی کوتاه و یا به اصطلاح نمونه‌گیری انجام دهند.

نمونه^۲

نمونه عبارت است از انتخاب درصدی از یک جامعه به‌عنوان نماینده آن جامعه اولین قدم در پژوهش علمی، تعریف جامعه بر اساس ویژگی مورد علاقه و سپس انتخاب یک نمونه از جامعه با استفاده از روش‌های مناسب است. در صورتی که اطلاعات جمع‌آوری شده از نمونه به منظور برآورد پارامتر جامعه به‌کار برده شود، در این صورت نمونه باید معرف یا نماینده واقعی جامعه باشد. مراد از معرف یا نماینده واقعی بودن این است که بین ویژگی‌های نمونه و جامعه‌ای که نمونه از آن انتخاب شده است، شباهت تقریباً کاملی وجود داشته باشد بر اساس این شباهت است که برآورد پارامتر جامعه امکان‌پذیر می‌شود. اندازه‌هایی که از نمونه به‌دست می‌آید، آماره^۳ نامیده می‌شود

1. Population
2. Sample
3. Statistic

به عبارت دیگر آماره ویژگی یا ویژگی‌هایی کمی است یک نمونه را توصیف می‌کند و پارامتر، ویژگی‌های جامعه را تعیین می‌کند. پارامتر را با حروف یونانی و آماره را با حروف اول انگلیسی نشان می‌دهند. برای روشن شدن این مطلب به جدول زیر توجه کنید. (دلاور، ۱۳۸۰: ۱۰).

ویژگی‌های شاخص‌های آماری در نمونه و جامعه

ویژگی شاخص	میانگین	واریانس	انحراف استاندارد	نسبت	همبستگی	تعداد مشاهدات
آماره	$\bar{X}$	S^2	S	P	r	n
پارامتر	μ	σ^2	σ	π	ρ	N

متغیر

متغیر^۱، خانه‌ای از حافظه است که دارای اسم و نوع می‌باشد و با توجه به نوع آن مقداری اطلاعات در آن ذخیره می‌شود. متغیر یک مفهوم است که بیش از دو یا چند ارزش یا عدد به آن اختصاص داده می‌شود کرلینجر^۲ (۱۹۸۶) معتقد است، متغیر یک سمبل است که می‌توان عدد یا ارزش را جایگزین آن کرد. در برخی مطالعات تنها یک متغیر مورد بررسی قرار می‌گیرد، چنین پژوهشی را تک متغیری^۳ می‌خوانند مانند تعیین میزان هوش دانش‌آموزان مقطع متوسطه در یک شهر و در زمان معین. اگر بررسی دارای دو متغیر باشد که غالباً نیز این چنین است مطالعه را دو متغیری^۴ می‌نامند، مانند بررسی رابطه بین هوش دانش‌آموزان مقطع متوسطه با میزان پیشرفت تحصیلی آنان. در مواردی که بررسی بیش از دو متغیر را شامل شود آن را چند متغیری^۵ می‌نامند مانند بررسی رابطه عوامل گوناگون از قبیل میزان تحصیلات والدین، هوش و پایگاه اقتصادی و اجتماعی بر پیشرفت تحصیلی. (نجفی زند و پاشاشریفی، ۱۳۷۶: ۱۵)

1. Variable
2. Kerlinger
3. univariate
4. Bivariate
5. multivariate

مقیاس‌های اندازه‌گیری

معمولاً اندازه‌هایی که برای تعیین مقدار متغیر مورد سنجش اختصاص می‌یابند از معنا و دقت یکسان برخوردار نیستند. از این لحاظ، کمیت‌ها دارای چهار مفهوم یا مقیاس می‌شوند. این مقیاس‌ها از ساده به پیچیده و به صورت سلسله مراتبی تنظیم شده است یعنی هر طبقه دارای تمامی ویژگی‌های طبقات قبل به اضافه ویژگی‌های مخصوص به آن طبقه است در زیر به بررسی ویژگی‌های مقیاس‌های چهارگانه می‌پردازیم.

الف) مقیاس اسمی^۱

این مقیاس صرفاً به تعیین طبقات می‌پردازد، طبقاتی که افراد، اشیاء یا حوادث ممکن است در آن‌ها جایگزین شوند. البته این طبقات بایستی مانع‌الجمع باشند، به این معنا که چنانچه از اعداد برای طبقه‌بندی افراد استفاده می‌شود، یک نفر را در بیش از یک طبقه جا نداد. و به هر طبقه یک عدد (مثل کدگذاری) اختصاص می‌دهیم. این به مفهوم ریاضی ارزش کمی ندارند و فقط برای مشخص کردن و یا نامیدن متغیرها به کار می‌روند. مثلاً هنگامی که افراد یک نمونه یا جامعه را به دو باسواد و بیسواد، متاهل و مجرد و طبقه‌بندی می‌کنیم. (دلاور، ۱۳۸۰: ۲۱)

عملیات مجاز آماری و ریاضی درباره مقیاس اسمی: عملیات مجاز درباره این مقیاس بسیار کم و به قرار زیر هستند:

عملیات مجاز آماری: شمارش فراوانی، تعیین نما

عملیات مجاز ریاضی: انجام هیچ‌یک از چهار عمل اصلی در این مقیاس میسر نیست (سیف، ۱۳۷۶: ۴۶).

ب) مقیاس رتبه‌ای^۲

هرگاه اعداد موقعیت فرد یا شیء مورد مطالعه را بر حسب صفت معین در میان افراد یا اشیاء گروه مشخص سازند، چنین کمیت‌هایی دارای مقیاس رتبه‌ای هستند. بنابراین مقیاس رتبه‌ای (ترتیبی) به منظور مرتب کردن افراد یا اشیاء از بالاترین به پایین‌ترین

1. Nominal scale

2. Ordinal scale

(براساس ویژگی‌های مورد اندازه‌گیری) بکار برده می‌شود. مثال مقاطع تحصیلی نمونه‌ای از مقیاس ترتیبی است (ابتدایی > راهنمایی > متوسطه) به عبارت دیگر مقیاس رتبه‌ای مشخص می‌کند که هر دانش‌آموز از نظر مقطع تحصیلی برتر یا پایین‌تر از فرد دیگر است و رابطه $(A > B)$ برقرار است (نجفی زند و پاشا شریفی، ۱۳۷۵: ۱۸).

عملیات مجاز آماری و ریاضی درباره مقیاس رتبه‌ای: در این مقیاس عملیات بیشتری از مقیاس اسمی مجازند. با این وجود، عملیات مورد نیاز یک مقیاس خوب اندازه‌گیری در این مقیاس نیز مجاز نیستند.

عملیات مجاز آماری: شمارش فراوانی، تعیین نما، محاسبه میانه، محاسبه درصدها و محاسبه ضریب همبستگی رتبه‌ای اسپیرمن^۱.

عملیات مجاز ریاضی: انجام هیچ‌یک از چهار عمل اصلی در این مقیاس میسر نیست. (سیف، ۱۳۷۶: ۴۳)

ج) مقیاس فاصله‌ای^۲

این مقیاس علاوه بر طبقه‌بندی، نام‌گذاری و مرتب کردن، به ما اجازه می‌دهد که فاصله‌های موجود بین افراد، اشیاء یا حوادث را دقیقاً مشخص کنیم (دلاور، ۱۳۸۰: ۲۲). در این مقیاس کمیت‌ها برحسب یک مبنا یا صفر قراردادی نسبت به هم در روی یک پیوستار قرار می‌گیرند نه برحسب صفر مطلق. بدین ترتیب میزان تفاوت میان افراد گروه مورد سنجش را از لحاظ متغیر مورد بررسی نشان می‌دهند، مثال نمرات دانشجویان در درس آمار. این نمره‌ها تفاوت دانشجویان را به صورت مدرج نشان می‌دهد و چنانچه فرد (الف) در آزمون آمار نمره ۲۰ و فرد (ب) نمره ۳۰ گرفته باشد می‌توان گفت که فرد (ب) ۱۰ واحد بیشتر از فرد (الف) آمار می‌دانسته است، اما نمی‌توان گفت که فرد (ب) ۱/۵ برابر فرد (الف) از آمار اطلاع داشته است، زیرا مبنای اولیه مشترک یا صفر مطلق حاکم نیست (نجفی زند و پاشا شریفی، ۱۳۷۵: ۲۲).

عملیات مجاز آماری و ریاضی درباره مقیاس فاصله‌ای: از آنجا که در این مقیاس فواصل بین واحدها برابر هستند و در نتیجه این مقیاس از یکی از ویژگی‌های مهم یک

1. Spearman Rank-Order Correlation

2. Interval Scale

مقیاس خوب اندازه‌گیری برخوردار است، لذا همه محاسبات آماری و غالب محاسبات ریاضی را درباره آن می‌توان انجام داد.

عملیات مجاز آماری: شمارش فراوانی، تعیین نما، محاسبه میانه، محاسبه درصدها، ضریب همبستگی رتبه‌ای و ضریب همبستگی گشتاوری پیرسون. عملیات مجاز ریاضی: جمع و تفریق مجاز است؛ ضرب و تقسیم مجاز نیست. (سیف، ۱۳۷۶: ۴۴)

د) مقیاس نسبی (نسبی)^۱

مقیاس نسبی بالاترین سطح اندازه‌گیری است و حدود فعالیت آن مشتمل بر کلیه عملیاتی که در مقیاس‌های اسمی، ترتیبی و فاصله‌ای صورت می‌پذیرد و هم از صفر مطلق برخوردار باشد، چنین اعدادی دارای مقیاس نسبی هستند. مثال قد افراد که با وسیله‌ای مانند متر اندازه‌گیری می‌شود. (دلاور، ۱۳۸۰: ۲۳)

عملیات مجاز آماری و ریاضی درباره مقیاس نسبی: در این مقیاس همه عملیات آماری و ریاضی مجاز است. (سیف، ۱۳۷۶: ۴۶)

خلاصه ویژگی‌های مهم مقیاس‌های چهارگانه اندازه‌گیری (سیف، ۱۳۷۶: ۴۶)

ویژگی‌ها					مقیاس
وجود ترتیب در طبقات	فواصل مساوی طبقات	وجود صفر مطلق	عملیات مجاز آماری	عملیات مجاز ریاضی	
-	-	-	فراوانی، نما	هیچ‌یک از چهار عمل اصلی	اسمی
+	-	-	فراوانی، نما، میانه، ضریب همبستگی اسپیرمن	هیچ‌یک از چهار عمل اصلی	رتبه‌ای
+	+	-	فراوانی، نما، میانه، میانگین، واریانس، انحراف معیار، ضریب همبستگی پیرسون	جمع و تفریق	فاصله‌ای
+	+	+	کلیه عملیات آماری	کلیه عملیات ریاضی	نسبی

آزمون فرضیه

فرضیه به دو دسته تقسیم می‌شود:

الف) فرض صفر^۱ ب) فرض خلاف^۲

فرض صفر: فرض صفر را با H_0 نشان می‌دهند. فرض صفر اصل را بر این قرار می‌دهد که بین پارامترهای مورد مطالعه اختلاف یا ارتباط معنی‌داری وجود ندارد. این فرض با آزمون‌های آماری به کار برده می‌شود و پژوهشگر برای تفسیر نتایج از فرض صفر استفاده می‌کند، این فرض معمولاً بدین‌منظور طرح می‌شود که بعداً آن را رد کنیم. (دلاور، ۱۳۷۶:۳۷۲) مثال:

الف) بین هوش و پیشرفت تحصیلی دانش‌آموزان رابطه معنی‌داری وجود ندارد.
ب) بین سطح سواد والدین و پیشرفت تحصیلی دانش‌آموزان رابطه معنی‌داری وجود ندارد.

فرض خلاف: فرض خلاف را با H_A یا H_1 نشان می‌دهند. این فرض، مخالف فرض صفر است و عیناً مطابق با فرضیه تحقیق است. وقتی که فرضیه پوچ (صفر) رد شد، فرضیه مقابل (H_1) مورد قبول واقع می‌شود یعنی اگر در یک پژوهشی نشان داده شود که بین متغیرهای مورد پژوهش همبستگی یا ارتباط معنی‌داری وجود دارد، فرض صفر رد می‌شود و ذکر این نکته نیز لازم است که فرض خلاف بیان‌کننده انتظار پژوهشگر درباره نتایج آتی پژوهش است. (دلاور، ۱۳۷۶:۳۷۲)

مثال:

الف) بین هوش و پیشرفت تحصیلی دانش‌آموزان رابطه معنی‌داری وجود دارد.
ب) بین سطح سواد والدین و پیشرفت تحصیلی دانش‌آموزان رابطه معنی‌داری وجود دارد.

خطای نوع اول^۱ و خطای نوع دوم^۲

بر اساس اطلاعات به دست آمده از نمونه تصادفی، محقق تصمیمی می گیرد که فرضیه صفر را رد کند یا رد نکند. این تصمیم یا صحیح خواهد بود یا صحیح نخواهد بود.

	H0 صحیح	H0 غیر صحیح
عدم رد	صحیح	خطای نوع دوم
رد	$(1-\alpha)$	$(1-\beta)$
H0	خطای نوع اول	صحیح

برای جلوگیری از این خطا (نوع اول) رد فرضیه صفر وقتی که فرضیه صفر صحیح است را به صورتی کاملاً محتاطانه $0/05$ یا $0/01$ (پنج درصد یا یک درصد) انتخاب می کنیم. احتمال اخذ تصمیم صحیح- رد نکردن فرضیه صفر وقتی که فرضیه صفر صحیح است- برابر است با $(1-\alpha)$ اگر $(\alpha=0/05)$ باشد، احتمال اخذ تصمیم صحیح وقتی که فرضیه صفر صحیح است $(0/95=1-0/05)$ خواهد بود.

اما زمانی که فرضیه صفر صحیح نمی باشد مسئله تصمیم گیری شکل دیگری به خود می گیرد. وقتی که فرضیه صفر صحیح نباشد، تصمیم محقق مبنی بر عدم رد فرضیه صفر خطا خواهد بود. احتمال چنین خطایی را با β نشان داده و آن را خطای نوع دوم می نامند

وقتی که فرضیه صفر صحیح نیست و محقق هم تصمیم به رد فرضیه صفر می گیرد (یعنی آنچه را در تحقیق جستجو می کند) تصمیم او صحیح خواهد بود. با احتمال اخذ تصمیم صحیح یعنی، رد فرضیه صفر غیر صحیح^۳ توان آماري گفته می شود. توان آزمون آماری بیانگر میزان احتمالی است که اگر واقعاً تفاوتی وجود داشته باشد تحقیق آن تفاوت را نشان خواهد داد. (دلاور، ۱۳۷۶: ۳۷۵)

چهار عامل بر توان آزمون تأثیر می گذارد که عبارتند از:

1. Type I Error
2. Type II Error
3. Power Of Test

۱. سطح معنی داری: سطح معنی داری (α) در کنترل محقق می باشد. با افزایش مقدار (α)، توان آزمون افزایش می یابد بنابراین اگر تمام عوامل مؤثر بر توان آزمون را ثابت نگه داریم، با افزایش مقدار از ۰/۰۱ به ۰/۰۵، توان آزمون افزایش خواهد یافت.
۲. میزان تأثیر عمل آزمایشی: هر چه تأثیر عمل آزمایشی افزایش یابد، توان آزمون افزایش می یابد.
۳. پراکندگی در جامعه: هر چه مقدار پراکندگی جامعه کمتر باشد (انحراف معیار^۱ کوچک تر باشد) توان آزمون آماری بیشتر خواهد بود.
۴. حجم نمونه: با افزایش حجم نمونه و ثابت نگه داشتن سایر عوامل، پراکندگی توزیع نمونه گیری کاهش می یابد در نتیجه توان آزمون افزایش می یابد. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۱۴۵)

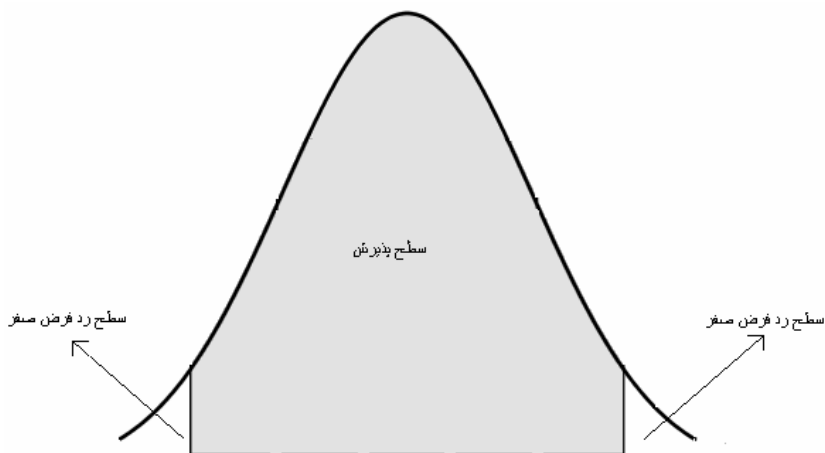
آزمون های یک دامنه ای^۲ و دو دامنه ای^۳

انتخاب نوع فرضیه جهت دار یا غیر جهت دار به میزان اطلاع محقق از موضوع تحت بررسی بستگی دارد. اگر نظریه جهت نتایج تحت بررسی را پیش بینی کرده باشد، از فرضیه جهت دار استفاده می شود. همچنین اگر در مورد جهت نتایج تحت بررسی شواهد تجربی قوی موجود باشد از فرضیه جهت دار استفاده می شود. که لازم است که اگر هر دو یعنی نظریه و شواهد تجربی جهت بازده های تحقیق را پیش بینی کنند، محقق از فرضیه جهت دار استفاده خواهد نمود. ($H_A: \mu_1 > \mu_2$) پس در آزمون های یک دامنه ای جهت تأثیر متغیر مستقل بر متغیر وابسته مشخص است. ولی اگر در مورد نتایج مورد مطالعه تردید وجود داشته باشد، مثلاً شواهد تجربی قبلی پایا نباشد و یا اینکه نظریه قوی و مناسب در اختیار نباشد، از فرضیه غیر جهت دار باید استفاده نمود. در آزمون مورد بحث هیچ نظری درباره جهت اختلاف بین میانگین ها مطرح نشده است به عبارت دیگر در این آزمون فرضیه ای دال بر اینکه کدام یک از میانگین ها بزرگ تر یا کوچک ترند، ذکر نشده است چنین آزمونی را بدون جهت (دو دامنه ای) می گویند.

1. Standard Deviation

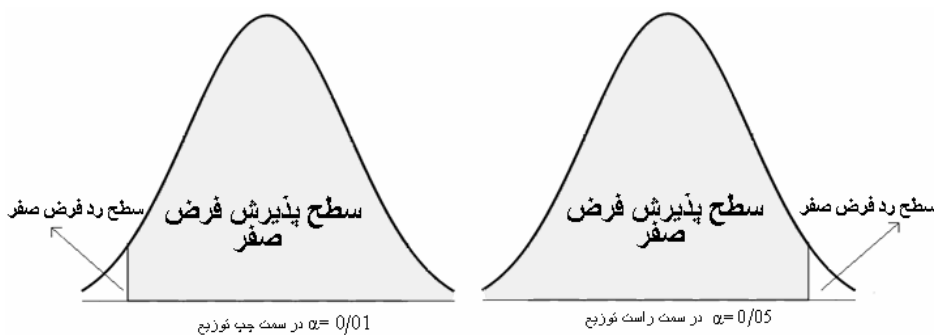
2. One taile

3. Two taile



اگر در مورد انتخاب آزمون یک‌سویه یا دوسویه تردید وجود داشته باشد از فرضیه مقابل غیر جهت‌دار استفاده می‌شود. اگر فرضیه صفر بر اساس آزمون دوسویه رد گردد، بر اساس آزمون یک‌سویه هم رد خواهد شد. هنگام اتخاذ تصمیم در مورد رد یا عدم رد فرضیه صفر، تصمیم محتاطانه‌تر ترجیح داده می‌شود. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۷۱)

بر اساس اطلاعات گروه نمونه، در مورد رد یا عدم رد فرض صفر تصمیم گرفته می‌شود. از نظر آماری، از طریق مقایسه مقدار مشاهده شده از گروه نمونه با مقدار بحرانی تصمیم گرفته می‌شود اگر مقدار مشاهده شده بزرگ‌تر یا مساوی باشد با مقدار بحرانی فرضیه صفر را رد می‌کنیم



مثلاً آیا بین میانگین ریاضی دانش‌آموزان کلاس پنجم استان و میانگین نمره ریاضی در سراسر کشور تفاوت وجود دارد؟ فرض کنید با استفاده از نورم ملی (آزمون استاندارد

شده در سطح کشور) میانگین نمره ریاضی دانش‌آموزان ۷/۸۷ و انحراف معیار ۱/۱ محاسبه شده است.

فرض کنید که آموزش و پرورش استان دلایل کافی در اختیار دارد که میانگین نمره‌های دانش‌آموزان از ۷/۸۷ بالاتر است. در این صورت فرضیه مقابل به صورت جهت‌دار مطرح می‌شود. (سطح اطمینان برای آزمون فرضیه $\alpha = 0/05$ انتخاب شده است).

برای مطالعه فرضیه صفر، یک گروه ۱۰۰ نفری به‌طور تصادفی از میان دانش‌آموزان استان انتخاب شده است. میانگین نمره‌های این گروه ۸/۲ محاسبه شده است. برای تصمیم‌گیری در مورد رد یا عدم رد فرضیه صفر، میانگین مشاهده شده از گروه نمونه را به نمره استاندارد تبدیل می‌کنیم

$$Z = \frac{\bar{X} - \mu}{\sigma_x}$$

برای محاسبه آزمون Z، ابتدا باید مقدار خطای معیار میانگین را محاسبه کرد

$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{N}}$$

$$\sigma_{\bar{x}} = \frac{1/1}{\sqrt{100}} = 0/11$$

پس از محاسبه مقدار خطای معیار میانگین با استفاده از فرمول Z مقدار نمره استاندارد را محاسبه می‌کنیم:

$$Z = \frac{8/2 - 7/87}{0/11} = 3$$

چون مقدار نمره استاندارد مشاهده شده (۳) از مقدار بحرانی (۱/۶۵) بزرگ‌تر است، فرضیه صفر رد می‌شود و فرضیه مقابل پذیرفته می‌شود. یعنی نتیجه می‌گیریم که میانگین نمره‌های دانش‌آموزان استان از دانش‌آموزان کشور بیشتر است به دلیل اینکه فرضیه مقابل جهت‌دار است، بنابراین ناحیه بحرانی^۱ فقط در یک طرف توزیع - سمت راست - قرار می‌گیرد. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۷۴)



تفسیر تأیید یا رد فرض صفر

هنگامی که فرض صفر تأیید می‌شود بدین معنی است که مدارک کافی برای اینکه نشان داده شود این فرض غلط است، وجود ندارد. تأیید فرض صفر به این معنی نیست که این فرض واقعاً درست است. در واقع تأیید فرض صفر به این معنی است که مدارک و اطلاعات کافی برای نتیجه‌گیری بر خلاف آن وجود ندارد. هنگامی که فرض صفر رد می‌شود به این معنی است که بین شاخص‌های آماری مورد مقایسه اختلاف معنی‌دار از نظر آماری وجود دارد. یعنی چنانچه فرض صفر واقعاً درست باشد احتمال اینکه اختلاف بین شاخص‌های مورد مقایسه شانسی باشد برابر α (خطای نوع اول) است. (دلاور، ۱۳۸۰: ۳۷۸)

خودآزمایی

۱. پژوهشگری علاقمند است که ویژگی‌های شخصیتی دانش‌آموزان تیزهوش یک منطقه را که بهره‌هوشی آن‌ها بالاتر از ۱۵۰ است، مطالعه کند او به ۳۰ نفر از آن‌ها دسترسی پیدا می‌کند و از بین آن‌ها ۱۰ نفر را انتخاب می‌نماید

الف) جامعه (ب) نمونه را در این پژوهش تعیین کنید

۲. آماره و پارامتر را تعریف و اطلاعات مورد نیاز در جدول زیر را تکمیل کنید.

فصل دوم

نحوه تعریف و به کارگیری متغیرها در Spss

هدف‌های یادگیری

از دانشجو انتظار می‌رود که پس از خواندن فصل دوم بتواند:

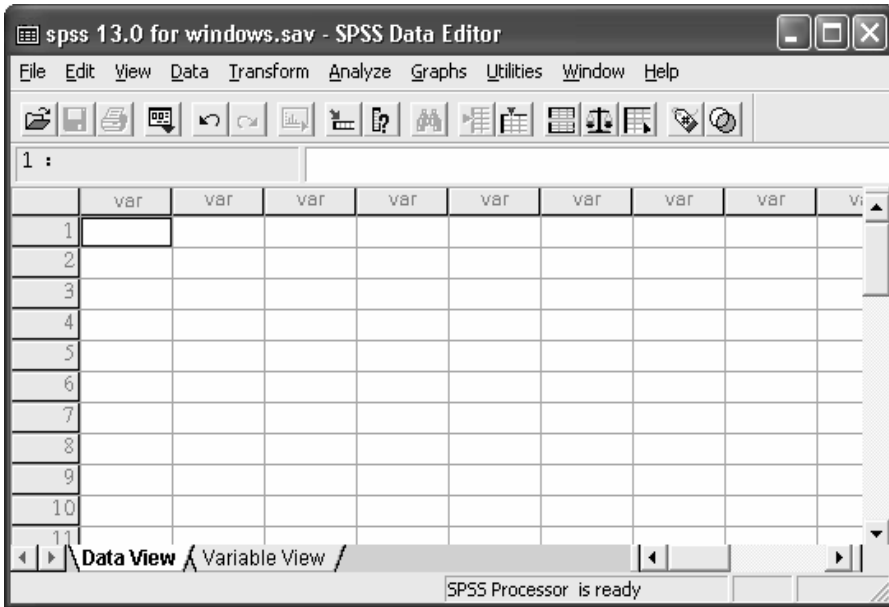
۱. متغیرها را به Spss معرفی کند.
۲. متغیرها را کد گذاری کند.
۳. فرمان Compute را انجام دهد.
۴. داده‌های کمی را به کیفی تبدیل کند.
۵. متغیرها را نام گذاری کند.

نحوه تعریف متغیر

قبل از معرفی متغیرها به Spss می‌باید نرم افزار spss را اجرا کرد. برای اجرای این نرم افزار در صفحه اصلی ویندوز، (start) را کلیک نمایید و سپس در منوی مربوط به آن، گزینه (programs) را کلیک و در فهرست برنامه‌ها (spss 13.0 for windows) را کلیک نمایید:

Start
Programs
Spss 13.0 for windows

بعد از اجرا، صفحه spss به صورت زیر ظاهر می‌شود:



بعد از اجرای این برنامه می‌باید متغیرهای تحقیق معرفی شوند. قبل از این کار محقق می‌باید متغیرها و مقادیر آن‌ها را کدگذاری کرده باشد.

مثال: مقطع تحصیلی متغیری است رتبه‌ای و با مقوله‌های (ابتدایی، راهنمایی، متوسطه) فرض کنید محقق نزد خود مقوله‌ها را به ترتیب زیر کدگذاری کند:

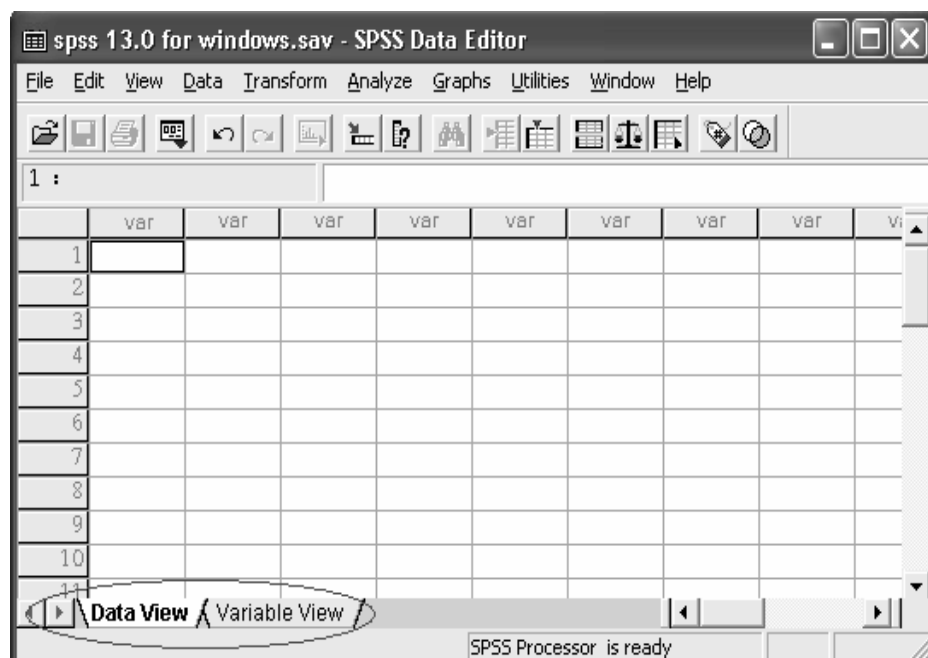
۱ = ابتدایی

۲ = راهنمایی

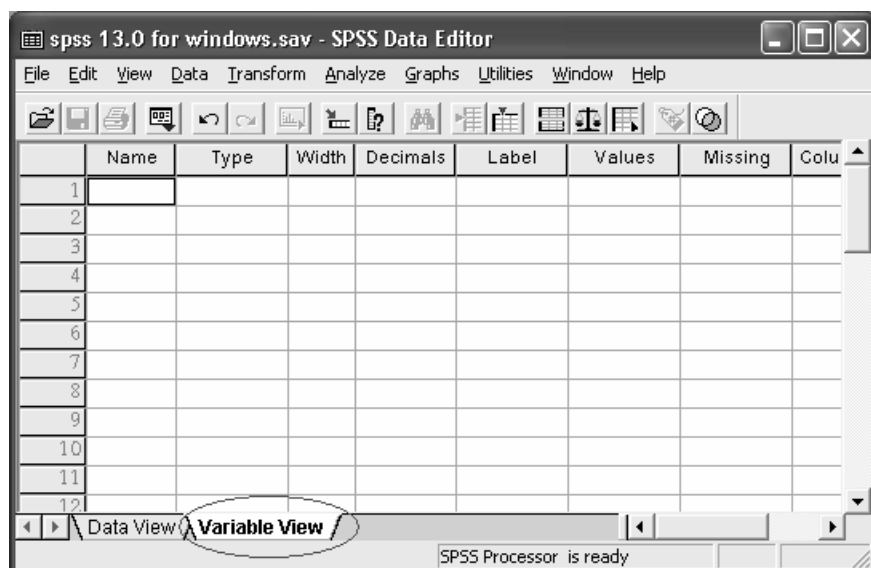
۳ = متوسطه

برای ورود اطلاعات به نکات زیر توجه نمایید:

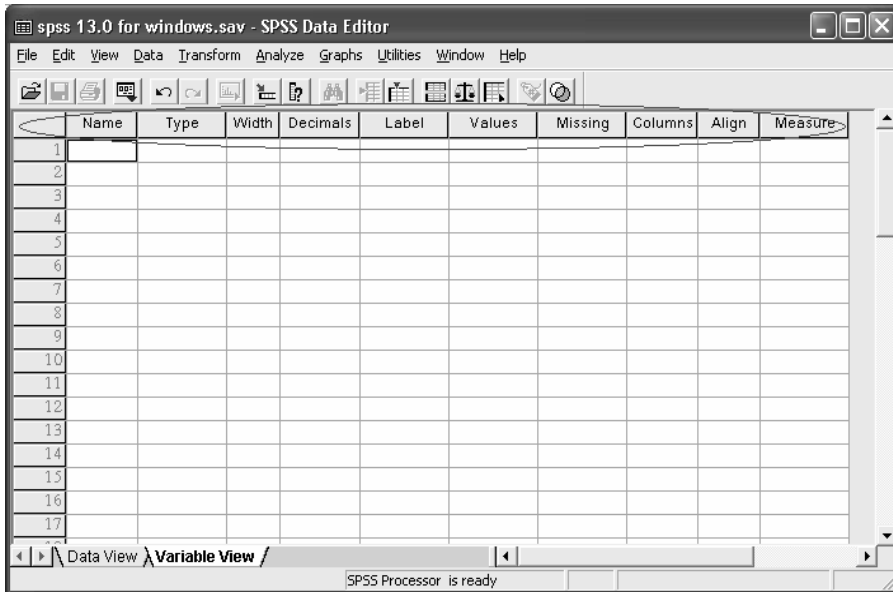
در گوشه سمت چپ و پایین صفحه اولیه spss دو لایه وجود دارد این دو لایه در صفحه اصلی spss تحت عنوان (data view) و (Variable View) نامگذاری شده‌اند. به شکل زیر توجه نمایید:



برای معرفی متغیر خود بر روی لایه (Variable View) کلیک نمایید تا منوی این صفحه ظاهر شود. شکل زیر این پنجره را نشان می‌دهد.



همچنان که دیده می‌شود این صفحه دارای ۱۰ ستون است. این ستون‌ها، به ترتیب در شکل زیر نشان داده شده است:



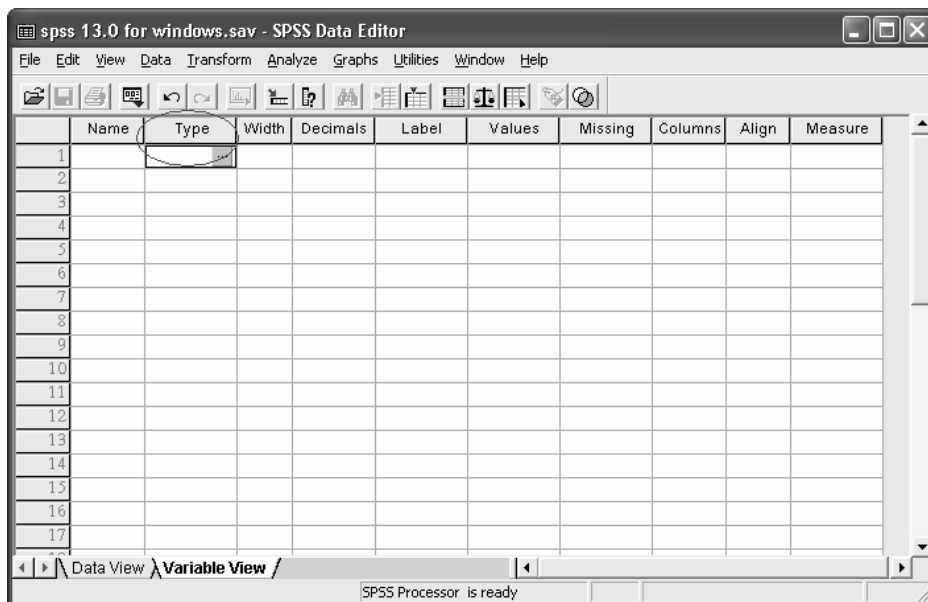
برای تعریف یک متغیر می‌باید به تعریف این ده ویژگی پرداخته شود.

نام متغیر (name)

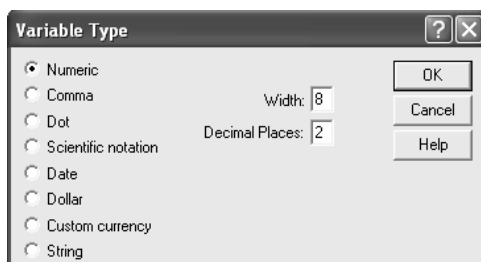
در این گزینه محقق باید برای متغیر خود نامی انتخاب کند. برای مثال متغیر درجه را در نظر آورید. در (spss13.0) قابلیت استفاده از حروف لاتین و فارسی برای نامگذاری متغیرها وجود دارد. بنابراین می‌توان نام این متغیر را به فارسی درجه بنویسید و یا کلمه لاتین آن را وارد نمایید و یا هر اسم دیگری را که مناسب است، انتخاب نمایید. توصیه می‌شود که اسامی متغیرها، را طوری انتخاب کنید که دارای مفهوم باشد و از کلمات بی‌مفهوم استفاده نکنید. همچنین از کاراکترهای مجاز برای نامگذاری خود استفاده کنید.

نوع متغیر (Type)

در این قسمت می‌بایست نوع متغیر مشخص شود، برای این منظور روی دومین خانه جدول (Type) کلیک نمایید:

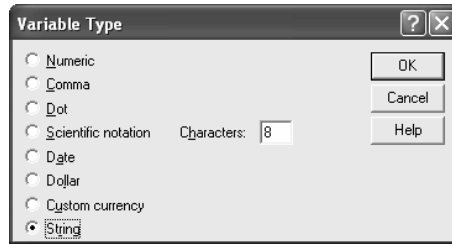


تا پنجره‌ای به شکل زیر ظاهر گردد:

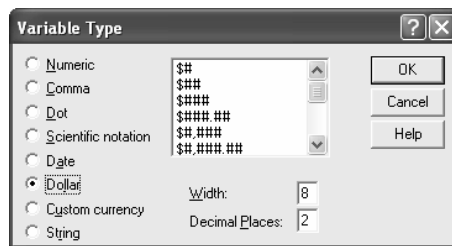


در این پنجره شما می‌توانید یکی از انواع متغیرها را انتخاب نمایید. مهم‌ترین انواع متغیرها در Spss عبارتند از:

۱. (string): این گزینه برای متغیرهای رشته‌ای و یا متغیرهایی که حروف در آن‌ها به کار برده می‌شود. انتخاب می‌گردد. (به‌عنوان مثال، نام افراد، فامیلی، آدرس و غیره). توجه کنید که برای هر رشته باید طول آن را مشخص کنید.

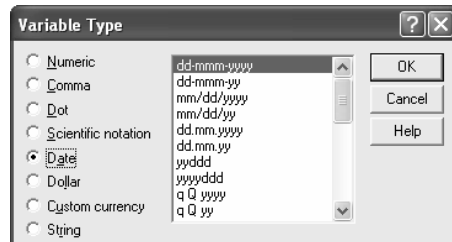


۲. (Dollar): این گزینه برای داده‌های کمی به کار می‌رود که واحد اندازه‌گیری آن‌ها، دلار می‌باشد:



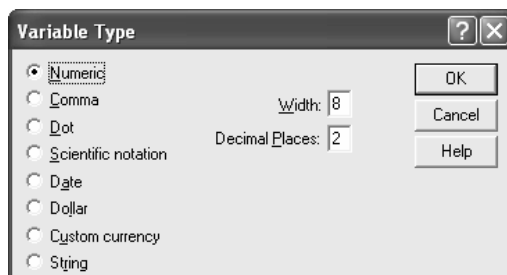
در این نوع از متغیرها باید طول متغیر و تعداد اعشار آن مشخص شود.

۳. (Date): این گزینه برای متغیرهایی به کار می‌رود که حاوی تاریخ (سال، ماه، روز) می‌باشند و بسته به نوع اطلاعات می‌تواند انواع مختلفی داشته باشد.



پنجره بالا روش‌های مختلف نمایش تاریخ را نشان می‌دهد که باید یکی از آن‌ها انتخاب شود.

۴. (Numeric): این گزینه برای داده‌های عددی به کار می‌رود در این گزینه محقق می‌تواند تعداد ارقام صحیح و اعشار متغیر خود را انتخاب کند.

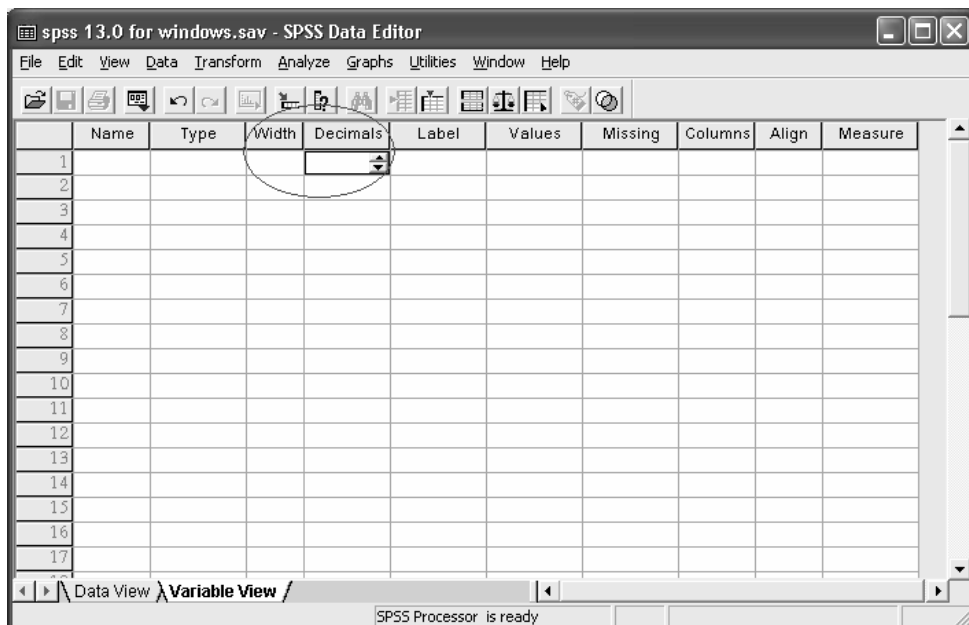


تعداد ارقام صحیح متغیر (Width)

در این قسمت محقق می‌تواند تعداد اعداد صحیح متغیر خود را با کلیک کردن روی آن افزایش و یا کاهش دهد.

تعداد ارقام اعشار متغیر (Decimals)

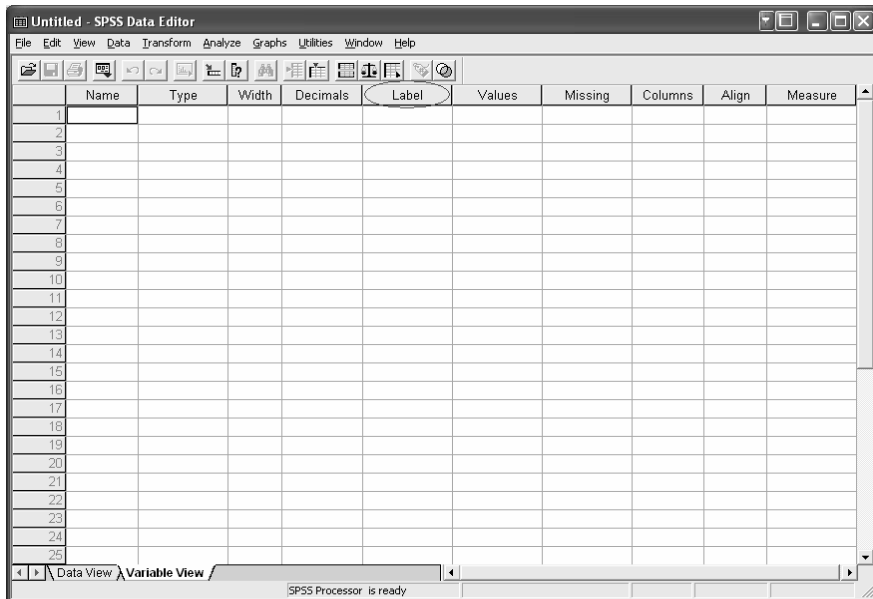
در این قسمت نیز محقق می‌تواند تعداد ارقام اعشار را برای متغیر خود افزایش و یا کاهش دهد. برای این منظور، روی این گزینه کلیک و به تغییر آن مبادرت نمایید. بقیه انواع متغیرها را می‌توانید از Help، نرم‌افزار مطالعه بفرمائید.



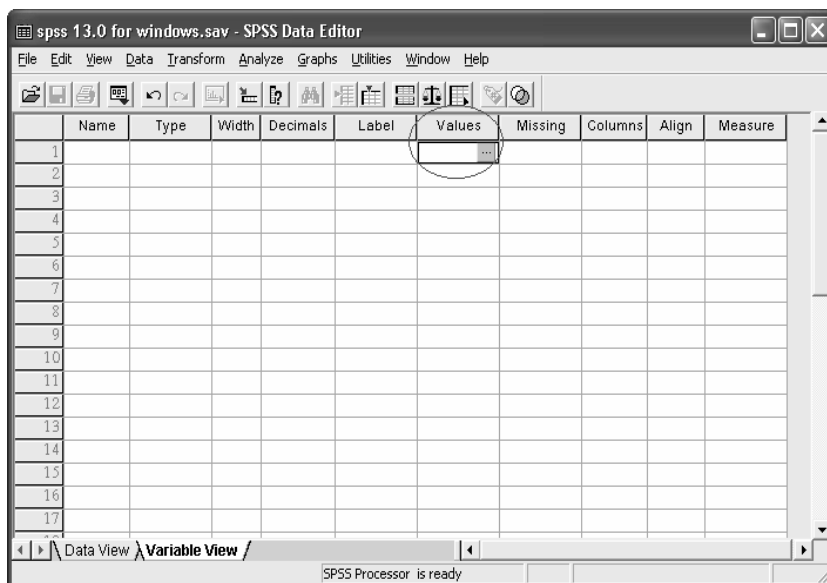
برچسب متغیر (Label)

برچسب یک کد شناسایی است که می‌تواند خالی باشد، ولی محقق چنان‌که مایل باشد می‌تواند یک عدد، رقم، و یا حرف را به متغیر خود نسبت بدهد. برای این منظور، با کلیک کردن بر روی آن، خانه آماده نوشتن می‌شود و می‌توان برچسب خود را وارد کرد.

باز توصیه می‌شود از برچسب‌های با مفهوم استفاده ننماید.

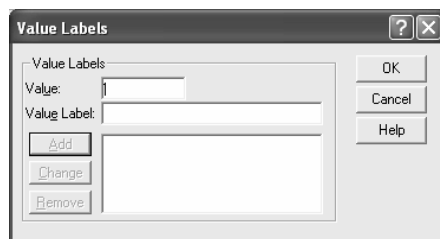
**کد گذاری مقوله‌های متغیر (Values)**

کد گذاری پس از اتمام مرحله گردآوری اطلاعات، محقق انبوهی اطلاعات را در اختیار دارد که باید از آن برای انجام اقدامات بعدی استفاده کند، یعنی اطلاعات موجود را استخراج و طبقه‌بندی نماید تا برای مرحله اساسی تجزیه و تحلیل آماده شود از این رو محقق ناگزیر است با در نظر داشتن نوع اطلاعات، ابزار سنجش و گردآوری داده‌ها، روش تجزیه و تحلیل و غیره، به کدگذاری اطلاعات بپردازد.



آنچه باید محقق در کدگذاری مورد توجه قرار دهد عبارتند از:

۱. پرسشنامه، هر پرسشنامه باید دارای شماره و کد ویژه باشد.
 ۲. منطقه و ناحیه، اگر پرسشنامه در سطح منطقه وسیعی اجرا می‌شود و جامعه مورد مطالعه از گستردگی جغرافیایی برخوردار است به هر یک از واحدهای تقسیمات جغرافیایی کد ویژه‌ای اختصاص دهید.
 ۳. سؤالات، هر یک از سؤالات نیاز به کد دارند ولی شماره ترتیب سؤال نقش کد را ایفا می‌کند.
 ۴. گزینه‌ها، گزینه‌های مربوط به پاسخ‌های احتمالی که در مقابل یا ذیل هر سؤال درج می‌شود، نیاز به کدگذاری دارند. برای سؤالات باز پاسخ باید ابتدا فهرستی از پاسخ‌های داده شده به هر سؤال را تهیه کرد و سپس باید این فهرست را کدگذاری نمود به نحوی که هر کدام از پاسخ‌ها در حکم گزینه تلقی شده و کدگذاری شوند. (حافظ نیا، ۱۳۸۰: ۱۸۷)
- در این گزینه شما می‌توانید به معرفی مقوله‌های متغیر خود اقدام کند، چنان‌که بر روی گزینه (Values) کلیک شود، پنجره گفت‌وگوی زیر ظاهر خواهد شد:



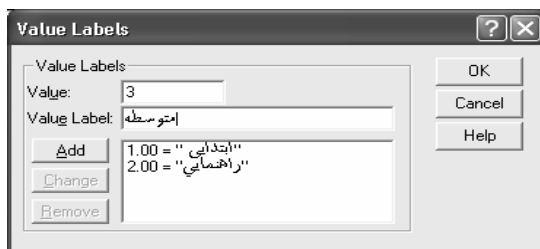
در این پنجره می‌توان به معرفی مقادیری که به مقوله‌های هر متغیر نسبت داده می‌شود. مبادرت کرد مثلاً، متغیر درجه نظامیان را به‌صورت زیر کدگذاری نمایید و کدهای مربوط را وارد نمایید:

۱ = ابتدایی

۲ = راهنمایی

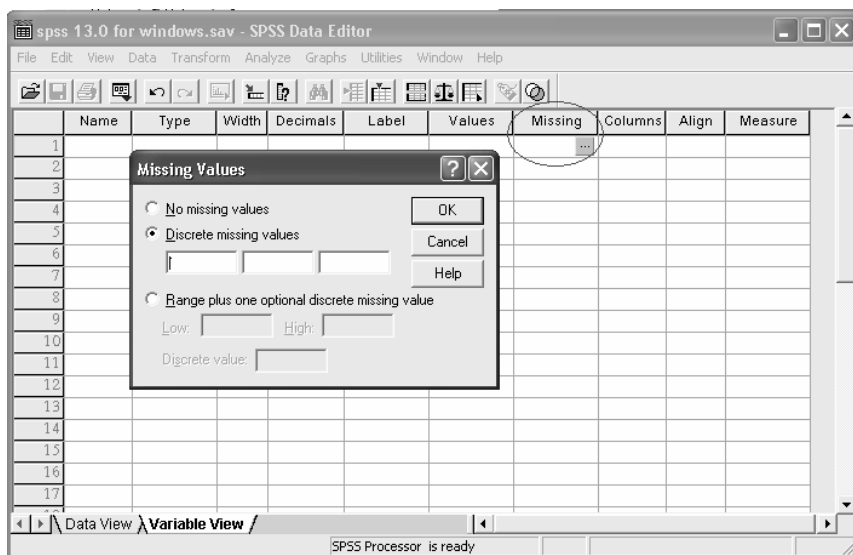
۳ = متوسطه

برای این منظور در فضای خالی (Value) عدد یک وارد نموده سپس در فضای خالی (Value Label)، کلمه سرباز وارد کنید و گزینه (Add) را کلیک نمایید. در این حالت در مربع پایین عبارت (ابتدایی = 1.00) اضافه می‌شود چنان‌که قبلاً ذکر شد، می‌توانید از حروف فارسی هم برای معرفی مقوله‌های هر متغیر استفاده نمایید. مجدداً به (Value) بروید و عدد دو را وارد کنید سپس در قسمت (Value Label) کلمه راهنمایی را وارد نمایید سپس (Add) کنید و به همین ترتیب ادامه دهید تا تمامی مقوله‌های متغیر را وارد کنید. همچنین می‌توانید از (Remove) برای پاک کردن و (Change) هم برای تغییر دادن مقوله‌های یک متغیر استفاده نمایید.

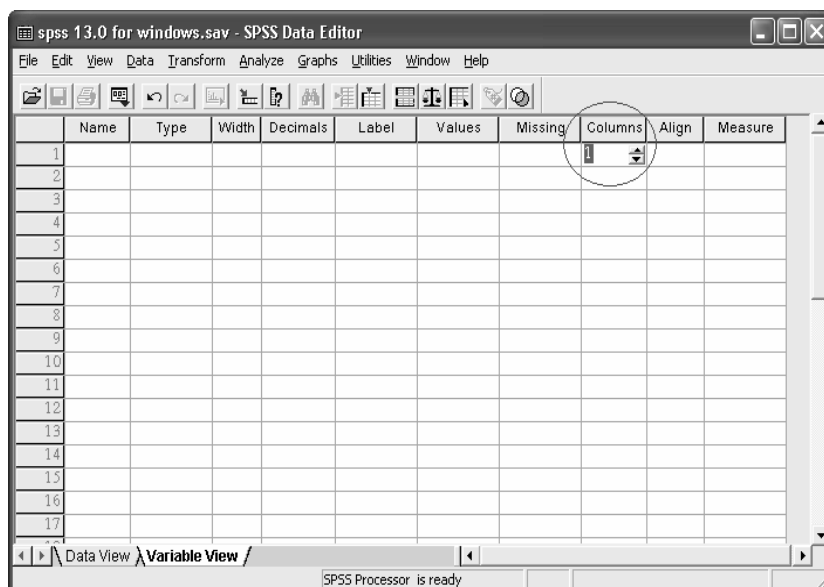


مقادیر نامشخص متغیر (Missing)

بارها اتفاق می‌افتد که پاسخگویان به متغیری خاص و یا پرسشی خاص جواب ندهند. این مقادیر در روش تحقیق به مقادیر نامشخص مشهورند. در این گزینه چگونگی برخورد Spss با این داده‌ها را مشخص نمایید.

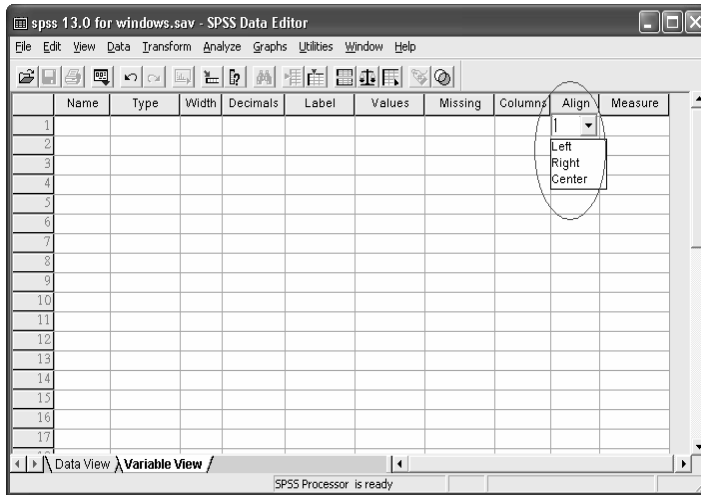


اندازه ستون نمایش‌دهنده متغیر (Columns)
این گزینه طول ستون متغیر را در صفحه اصلی مشخص می‌نماید. با کلیک کردن بر روی آن می‌توانید در منوی اصلی جای بیشتری را برای این متغیر در نظر بگیرید.



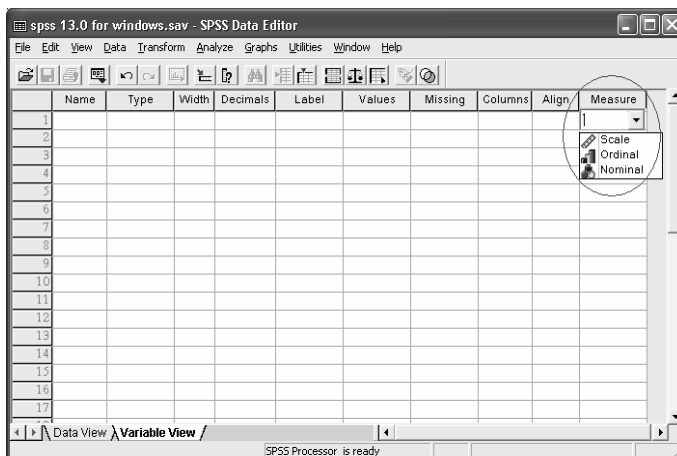
چگونگی نمایش هر ستون (Align)

این قسمت چگونگی نمایش اعداد را در یک ستون مشخص می‌کند. گزینه (Left) باعث می‌شود که اعداد در قسمت چپ سلول، گزینه (Right) در سمت راست سلول و گزینه (Center) باعث قرار گرفتن اعداد در وسط سلول می‌شود.



مقیاس اندازه‌گیری متغیر (Measure)

در این گزینه، مقیاس اندازه‌گیری متغیر مشخص می‌شود (Nominal) برای مقیاس اسمی، گزینه (Ordinal) برای متغیرهایی با مقیاس رتبه‌ای و گزینه (Scale) برای متغیرهای با مقیاس نسبی / فاصله‌ای به کار برده می‌شود.



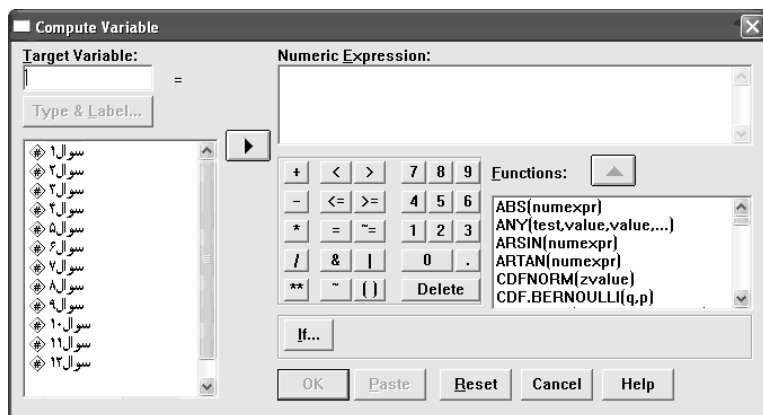
پس از تعیین این ده ویژگی برای متغیر نخست، می‌توانید در ردیف دوم به معرفی متغیر بعدی خود اقدام نمایید و پس از معرفی همه متغیرها، در قسمت پایین صفحه Spss (در سمت چپ) بر روی لایه (Data View) کلیک نمایید، چنان‌که مشاهده خواهید نمود، هر متغیر شما تحت عنوان یک ستون ساخته شده است و Spss آماده ورود اطلاعات می‌باشد و می‌توانید، اطلاعات هر متغیر را، به ترتیبی که کدگذاری نموده‌اید، در Spss وارد کنید.

تغییر در داده‌ها (Compute)

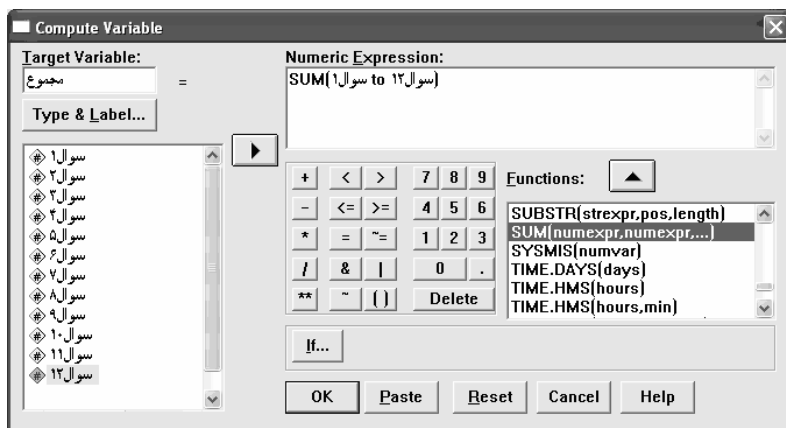
بارها اتفاق می‌افتد که می‌خواهید از روی اطلاعات متغیرهای موجود، متغیرهای جدیدی را بسازید مثلاً تصمیم دارید برای سؤالات ۱-۱۲ یک مجموع (Sum) محاسبه کنید این اعمال در زیر منوی (Compute) از منوی (Transform) امکان‌پذیر است برای تغییر در داده‌ها مراحل زیر را انجام می‌دهید:

Transform
Compute

با اجرای فرمان فوق پنجره زیر باز می‌شود:



با باز شدن پنجره فوق گزینه (SUM) را از قسمت (Function) به بالا منتقل کنید و به جای گزینه (numexpr,numexpr,...) سؤالات ۱-۱۲ را قرار داده و در قسمت (Target Variable) نام متغیر جدید خود را تایپ نمایید.



با انجام مراحل مذکور یک مجموع در صفحه اصلی spss ایجاد می‌شود. همان‌طور که ملاحظه می‌کنید در پنجره Compute تقریباً تمام دستورات مورد نیاز قرار گرفته‌اند و شما در طول تحقیق خود می‌توانید از این دستورات استفاده کنید.

spss 13.0 for windows.sav - SPSS Data Editor

	سوال ۴	سوال ۵	سوال ۶	سوال ۷	سوال ۸	سوال ۹	سوال ۱۰	سوال ۱۱	سوال ۱۲	sum	var
1	2	3	3	3	3	2	3	2	2	31.00	
2	4	4	2	3	4	4	2	3	4	41.00	
3	3	3	2	3	2	2	4	3	3	35.00	
4	2	3	3	3	3	2	2	3	2	30.00	
5	2	3	4	2	4	4	3	3	2	37.00	
6	2	3	3	3	3	2	3	3	3	33.00	
7	2	2	3	4	2	3	2	3	2	31.00	
8	3	3	3	4	3	3	3	3	3	36.00	
9	3	4	3	2	3	3	4	3	3	37.00	
10	4	4	2	4	2	4	3	4	2	38.00	
11	3	4	4	4	3	3	2	1	1	33.00	
12	4	4	4	4	2	2	2	2	2	35.00	
13	4	4	4	4	3	4	3	4	3	42.00	
14	3	4	4	4	3	3	2	2	2	33.00	
15	4	2	4	3	4	4	3	3	4	42.00	
16	4	4	3	3	3	4	3	2	2	38.00	
17	4	4	3	3	3	4	3	3	2	39.00	
18	3	2	2	3	2	2	2	2	2	30.00	
19	3	3	4	4	2	1	2	4	3	37.00	
20	4	2	4	3	4	4	3	3	4	42.00	

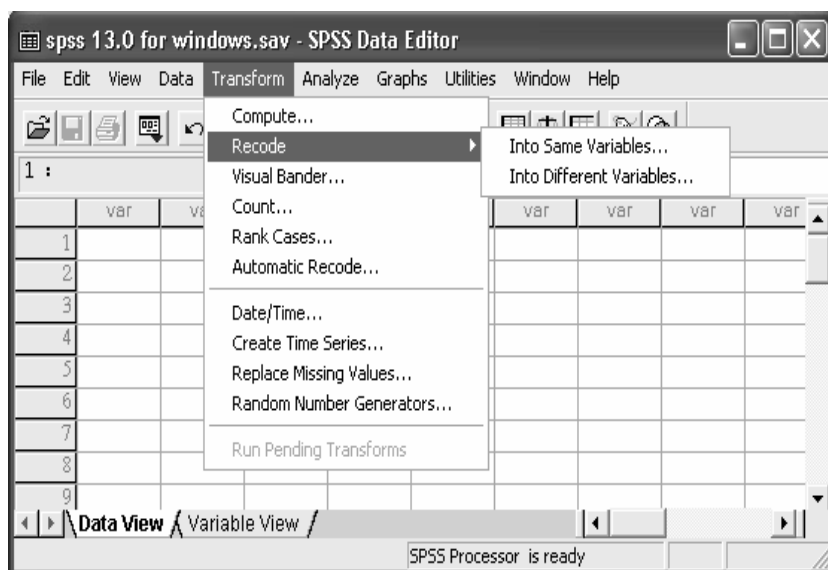
Data View Variable View | SPSS Processor is ready

تبدیل داده‌های کمی به کیفی (Recode)

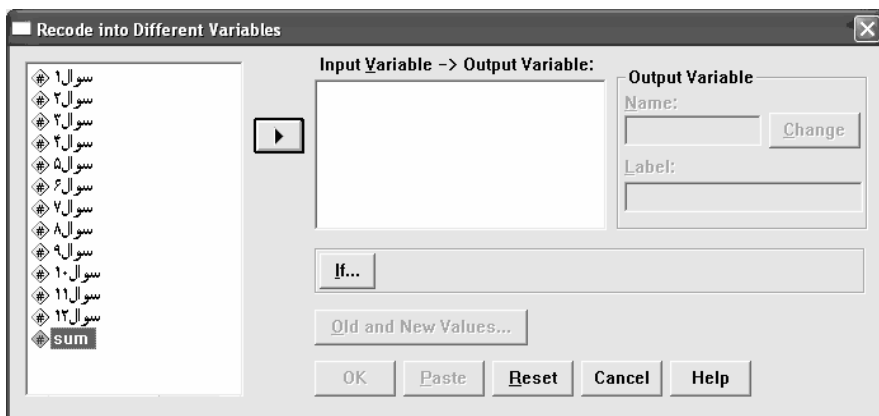
در بعضی از مواقع تصمیم دارید، داده‌های خود را از یک مقیاس کمی به مقیاس کیفی تبدیل کنید. برای این کار از فرمان (Recode) استفاده کنید برای اجرای فرمان (Recode) مراحل زیر را انجام دهید:

Transform

Recode Into Different variables....



با اجرای فرمان فوق پنجره گفت و گوی زیر باز می‌شود:



متغیر مربوطه را به قسمت Input Variable → Output Variable منتقل کنید. سپس در قسمت (Name) نام متغیر جدید خود را وارد و سپس گزینه (Old and New Values) را کلیک نمایید آنگاه پنجره زیر ظاهر خواهد شد:

Recode into Different Variables: Old and New Values

Old Value

- ☐ Value: []
- ☐ System-missing
- ☐ System- or user-missing
- ☒ Range: [] through []
- ☐ Range: Lowest through []
- ☐ Range: [] through highest
- ☐ All other values

New Value

- ☒ Value: []
- ☐ System-missing
- ☐ Copy old value(s)

Old --> New:

[Add] [Change] [Remove]

☐ Output variables are strings Width: 8

☐ Convert numeric strings to numbers ('5' -> 5)

[Continue] [Cancel] [Help]

در این پنجره، گزینه (Range) را انتخاب و حدود خود را به ترتیب وارد نمایید در این قسمت می‌توانید (Through) را وارد و سپس در کادر (New Values) در قسمت (Value) عدد یک را به این طبقه منتسب کنید و سپس (Add) را کلیک نمایید و مراحل گفته شده را مجدداً طی نمایید:

Recode into Different Variables: Old and New Values

Old Value

- ☐ Value: []
- ☐ System-missing
- ☐ System- or user-missing
- ☒ Range: 42 through 46
- ☐ Range: Lowest through []
- ☐ Range: [] through highest
- ☐ All other values

New Value

- ☒ Value: 5
- ☐ System-missing
- ☐ Copy old value(s)

Old --> New:

[Add] [Change] [Remove]

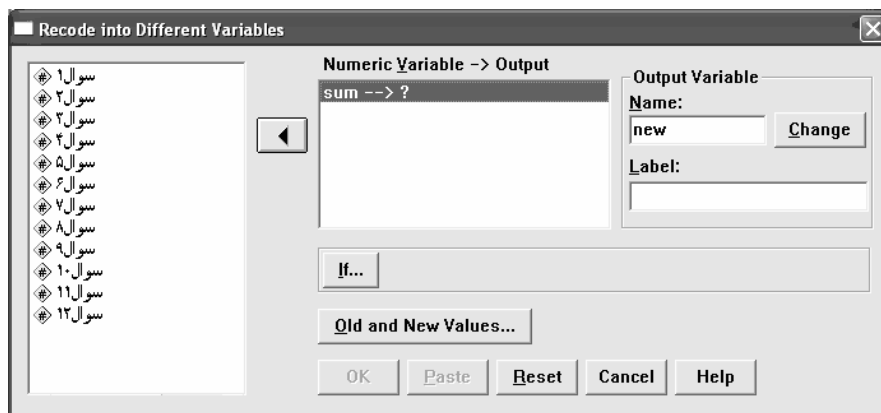
23 thru 27 --> 1
28 thru 31 --> 2
32 thru 36 --> 3
37 thru 41 --> 4

☐ Output variables are strings Width: 8

☐ Convert numeric strings to numbers ('5' -> 5)

[Continue] [Cancel] [Help]

بعد از سه مرحله انجام شده، متغیرهای شما به قسمت (Old-New) منتقل می‌شود که گزینه (Continue) را فعال نمایید تا به (Recode into Different Variables) می‌رسید در این پنجره در قسمت (Name) نام متغیر جدید خود را تایپ کنید و گزینه (Change) و سپس (OK) را کلیک نمایید:

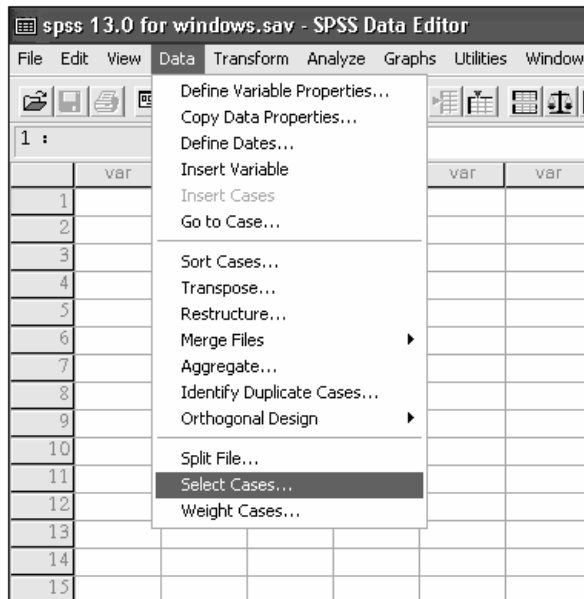


به این ترتیب متغیر جدیدی با مقیاس رتبه‌ای با مقوله‌های ۱، ۲، ۳، ۴ و ۵ در صفحه اصلی spss ساخته خواهد شد.

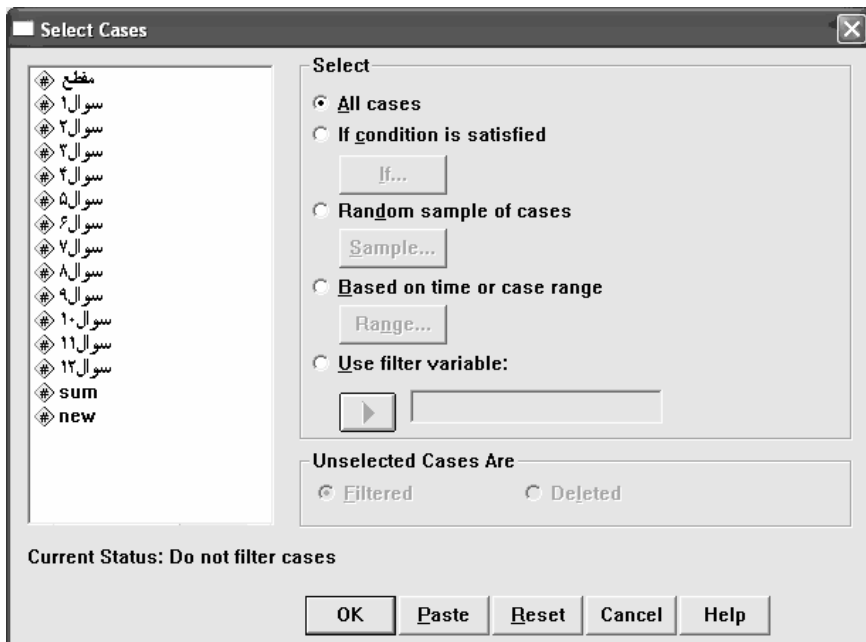
	۷ سوال	۸ سوال	۹ سوال	۱۰ سوال	۱۱ سوال	۱۲ سوال	sum	new	var	var	var
1	3	3	2	3	2	2	31.00	2.00			
2	3	4	4	2	3	4	41.00	4.00			
3	3	2	2	4	3	3	35.00	3.00			
4	3	3	2	2	3	2	30.00	2.00			
5	2	4	4	4	3	3	37.00	4.00			
6	3	3	2	3	3	3	33.00	3.00			
7	4	2	3	2	3	2	31.00	2.00			
8	4	3	3	3	3	3	36.00	3.00			
9	2	3	3	4	3	3	37.00	4.00			
10	4	2	4	3	4	2	38.00	4.00			
11	4	3	3	2	1	1	33.00	3.00			
12	4	2	2	2	2	2	35.00	3.00			
13	3	4	3	3	4	3	42.00	5.00			
14	4	3	3	2	2	2	33.00	3.00			
15	3	4	4	3	3	4	42.00	5.00			
16	3	3	4	3	2	2	38.00	4.00			
17	3	3	4	3	3	2	39.00	4.00			
18	3	2	2	2	2	2	30.00	2.00			
19	4	2	1	2	4	3	37.00	4.00			
20	3	4	4	3	3	4	42.00	5.00			

انتخاب داده‌ها (Select Casess)

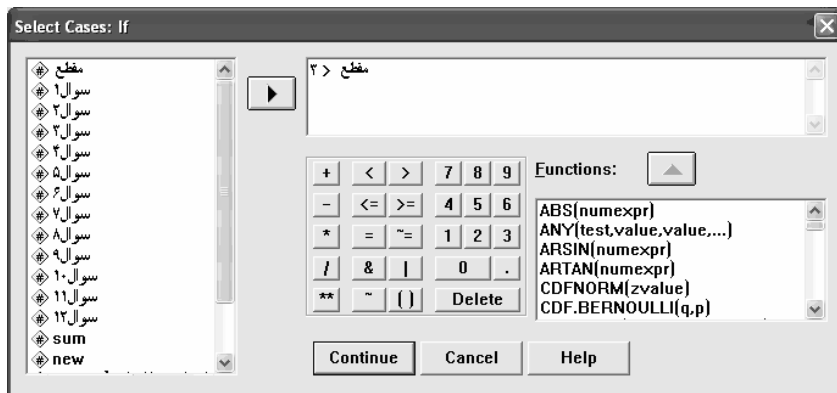
بعضی از مواقع می‌خواهید که مجموعه‌ای از اعداد در محاسبات لحاظ نشوند. برای چنین حالتی می‌توان فرمان (Select Casess...) را از منوی (Data) در پنجره اصلی نرم‌افزار (Spss) اجرا کرد:



در قسمت راست منوی (Select) انتخابهای زیادی وجود دارد که شما می‌توانید هر کدام را انتخاب نمایید. انتخاب پیش‌فرض نرم‌افزار بر (All Cases) قرار دارد این گزینه بیانگر آن است که تمامی افراد در محاسبات منظور شده‌اند.

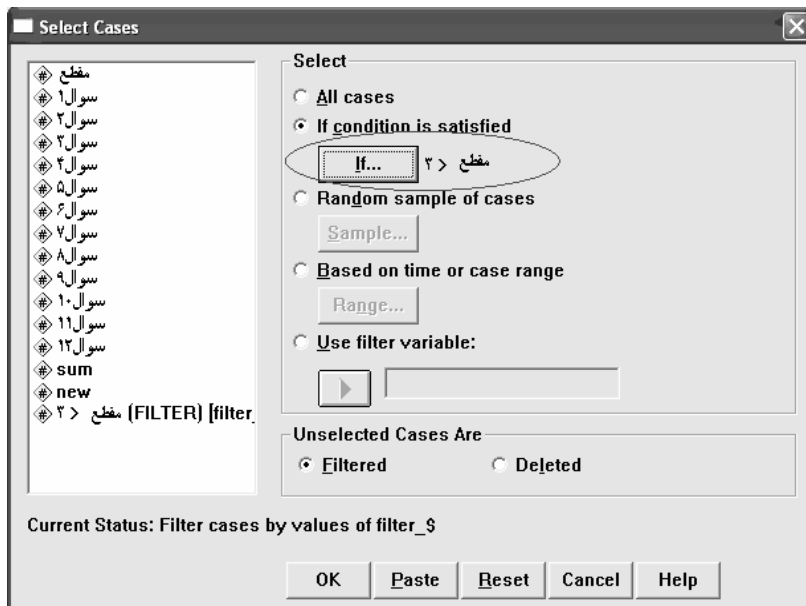


برای انتخاب متغیرهایی با ویژگی خاص، گزینه (If Condition is satisfied) را انتخاب و (If) را کلیک نمایید پنجره گفتگوی زیر ظاهر خواهد شد:



در حقیقت از پنجره بالا برای بیان شرطها استفاده می‌کند. و توسط آن هر نوع شرطی را برای مقادیر متغیرها می‌توان در نظر گرفت.

در این پنجره می‌توانید هر متغیر را که بخواهید به سمت راست منتقل نمایید و با استفاده از نمادهای ریاضی برای آن شرط تعیین کنید. سپس گزینه (continue) را کلیک کنید تا پنجره زیر ظاهر شود:



در پنجره فوق در قسمت (If Condition is Satisfied) مشاهده می‌کنید که عبارت (مقطع > ۳) آشکار می‌شود. سپس بر روی گزینه (OK) کلیک کنید و مشاهده نمایید که در صفحه اصلی (Spss) بر روی رکوردهایی که در محاسبات لحاظ نمی‌شوند خط مورب کشیده شده است.

spss 13.0 for windows.sav - SPSS Data Editor

File Edit View Data Transform Analyze Graphs Utilities Window H

سوالات: ۱۰ 3

	مقطع	سوال ۱	سوال ۲	سوال ۳	سوال ۴	filter_\$
1	1.00	1.00	4.00	1.00	1.00	1
2	2.00	2.00	1.00	2.00	4.00	1
3	3.00	1.00	4.00	5.00	1.00	0
4	1.00	2.00	2.00	1.00	3.00	1
5	3.00	1.00	3.00	4.00	1.00	0
6	2.00	2.00	5.00	3.00	2.00	1
7	3.00	5.00	1.00	3.00	5.00	0
8	3.00	1.00	5.00	3.00	1.00	0
9	2.00	4.00	2.00	3.00	4.00	1
10	1.00	3.00	1.00	5.00	3.00	1
11	3.00	3.00	4.00	1.00	3.00	0
12	1.00	3.00	1.00	5.00	3.00	1
13	3.00	2.00	4.00	2.00	2.00	0
14	2.00	1.00	1.00	1.00	3.00	1

در منوی (Select Cases) می‌توانید قسمت (Random sample of Cases) را انتخاب و (Sample) را کلیک نمایید که در صفحه گفتگوی زیر آشکار می‌شود:

Select Cases: Random Sample

Sample Size

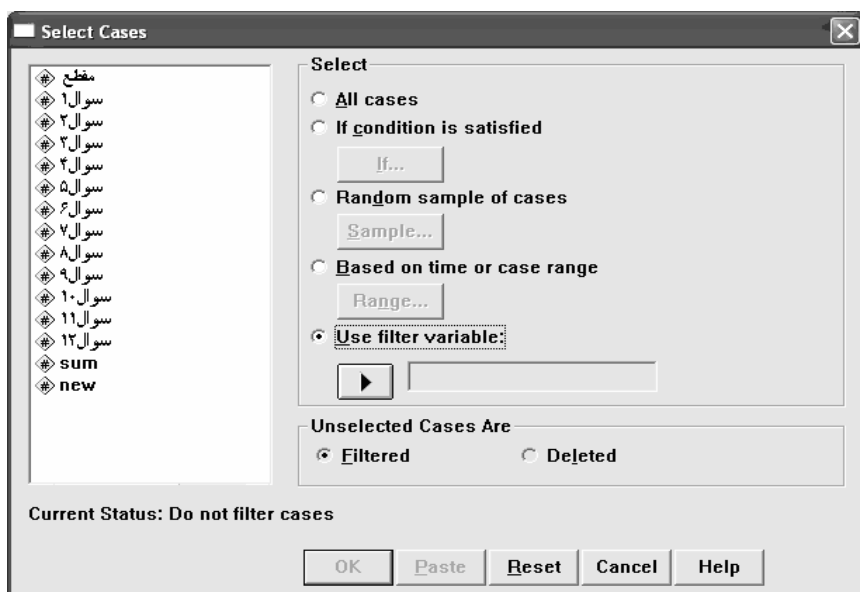
☒ Approximately 90 % of all cases

☐ Exactly cases from the first cases

Continue Cancel Help

در قسمت (Approximately) می‌توانید هر درصد دلخواه را وارد کنید تا به صورت تصادفی انتخاب شود. مثلاً چنانچه بخواهید ۹۰٪ از کل نمونه در محاسبات منظور شوند، می‌توانید در این قسمت ۹۰ را وارد نمایید یا اگر بخواهید ۴ داده از ۱۲ داده نخست انتخاب شوند، در قسمت خالی اولیه عدد ۴ و در فضای خالی بعدی عدد ۱۲ را وارد، سپس گزینه (Continue) و پس از آن (OK) را کلیک نمایید.

همچنین در منوی (Select Cases) می‌توانید قسمت (Use Filter Variable) را انتخاب کنید:



برای انتخاب متغیرها می‌توانید از گزینه‌های (Use Filter Variable) استفاده نمائید این گزینه را انتخاب و سپس هر متغیری را که خواهان آن هستید به فضای خالی منتقل کنید. سپس (Filtered) را انتخاب و (OK) را کلیک نمائید. این عمل باعث می‌شود که رکوردهایی که این متغیر در آن مقدار ندارد، از جمع نمونه‌ها حذف شود چنانچه در کادر گفتگوی به جای (Filtered) گزینه (Deleted) را انتخاب نمائید، این رکوردها در صفحه به کلی ناپدید می‌شوند.

نکته مهم: بعد از تعریف متغیرهای مختلف و مقداری آن‌ها برای ذخیره داده‌ها در حافظه جانبی رایانه که توسط نرم‌افزار Spss ضرورت گرفته، گزینه Save در منوی

File Spss را انتخاب نمایید. سپس برای این پروژه یک اسم فایل انتخاب کنید تا در مراجعات بعدی به این پروژه بتوانید داده‌ها را رویت کنید. در غیراین صورت ممکن است با هر اتفاق کوچک (قفل کردن کامپیوتر، قطع برق) مقادیر متغیرها از بین برود.

خود آزمایی

۱. نمرات زیر نتیجه امتحان درس آمار ۴۰ دانشجوی علوم تربیتی و روان‌شناسی است.

علوم تربیتی		روان شناسی	
۱۷	۱۵	۱۹	۱۰
۱۹	۱۴	۱۶/۲۵	۱۲
۱۶/۲۵	۱۲	۱۴	۱۵
۱۰	۱۵	۲۰	۱۶/۲۵
۱۲	۱۷	۲۰	۱۷
۱۲/۵	۱۸	۱۲/۵	۱۷
۱۷	۱۹	۱۴	۱۰
۱۸	۱۶	۱۲	۱۶
۱۹	۱۶/۲۵	۱۹	۱۴
۹	۲۰	۱۰	۲۰

الف) نمرات را با تعریف متغیرهایی در Spss وارد کنید.

ب) نمرات را به صورت زیر دسته‌بندی کنید:

- نمرات ۱۷-۲۰ درجه عالی
- نمرات ۱۵-۱۶ درجه خوب
- نمرات ۱۵ به پایین درجه ضعیف

ج) برای نمرات مذکور، مجموع (SUM) را محاسبه نمایید.

د) چنانچه معلم بخواهد تعدادی از نمرات (مثلاً نمره دانشجویان علوم تربیتی) در محاسبات لحاظ نشوند، از چه فرمانی استفاده می‌کند، فرمان را اجرا کنید.

۲. پژوهشگری به بررسی عواملی مؤثر بر کارایی معلمان پرداخته است متغیرهای اثرگذار در این پژوهش عبارتند از:

- مدرک تحصیلی (دیپلم، فوق‌دیپلم، لیسانس، فوق‌لیسانس)
- رشته تحصیلی (علوم انسانی، علوم تجربی، ریاضی و فنی، هنر)
- جنسیت (مرد، زن)

اکنون هر یک از متغیرهای این پژوهش را در Spss، کدگذاری کنید.

۳. کاربرد هریک از آیکن‌های زیر را بیان کنید:

- Type
- Width
- Decimals
- Lable
- Values
- Measure

۴. به منظور اندازه‌گیری رابطه بین اقتصاد خانواده و پذیرش نظام جدید آموزش و پرورش از دانش‌آموزان یک منطقه آموزشی، نمونه‌ای به‌صورت تصادفی انتخاب شده و اطلاعات زیر جمع‌آوری گردیده است، اکنون اطلاعات جدول زیر را به نرم‌افزار انتقال دهید

مخالف	موافق	پذیرش نظام جدید
		وضعیت اقتصادی
۱۰	۱۰	بالا
۱۵	۵	متوسط
۱۳	۷	پایین

فصل سوم

توزیع‌های فراوانی و ترسیم نمودارها

هدف‌های یادگیری

از دانشجو انتظار می‌رود که پس از خواندن فصل سوم بتواند:

۱. جدول توزیع فراوانی تشکیل دهد.
۲. فراوانی مطلق را محاسبه نماید.
۳. فراوانی نسبی درصدی و فراوانی تراکمی درصدی را محاسبه نماید.
۴. نمودارهای زیر را رسم کند.
 - ستونی (میله‌ای)
 - هیستوگرام
 - دایره‌ای

توزیع فراوانی

پس از جمع‌آوری اطلاعات با توده‌ای از اطلاعات روبرو هستیم که به‌صورت انبوهی از اعداد با مقیاس‌های گوناگون است و نیاز به تفسیر دارند اگر این اعداد به‌نحوی تنظیم و دسته‌بندی نشوند، استفاده از آن‌ها برای یافتن برخی پاسخ‌ها با مراجعه به این اعداد دشوار خواهد بود.

مثال: نمره‌های هوش ۵۰ دانش‌آموز در زیر آمده است با مراجعه به این اعداد نمی‌توان مشخص کرد که کمترین یا بیشترین میزان هوش در چه حدی است یا چند نفر از شرکت‌کنندگان دختر یا پسر هستند؟

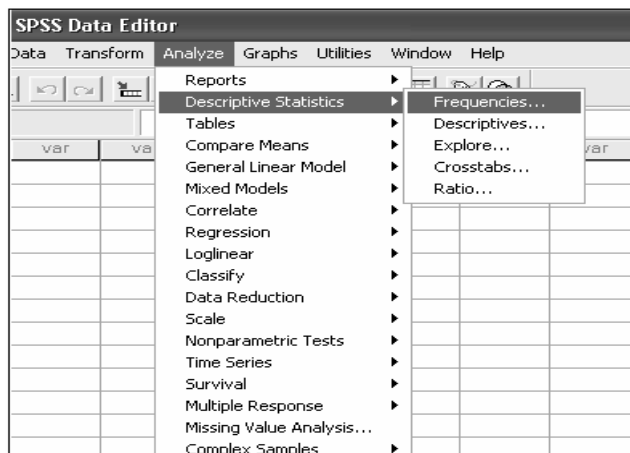
۱۲۰	۱۰۵	۱۱۳	۱۲۰	۱۲۵
۱۱۲	۱۲۵	۱۲۵	۱۱۲	۱۲۴
۱۱۰	۱۲۴	۱۲۴	۱۱۰	۱۱۳
۱۲۳	۱۱۳	۱۱۳	۱۲۳	۱۰۶
۱۰۵	۱۰۶	۱۰۶	۱۰۶	۱۲۸
۱۲۵	۱۲۸	۱۲۸	۱۲۸	۱۰۵
۱۲۴	۱۲۰	۱۰۵	۱۲۰	۱۲۵
۱۱۳	۱۱۲	۱۲۵	۱۱۲	۱۲۴
۱۰۶	۱۱۰	۱۲۴	۱۱۰	۱۱۳
۱۲۸	۱۲۳	۱۱۳	۱۲۳	۱۰۶

یکی از روش‌هایی که برای تنظیم و سازمان دادن به این قبیل داده‌ها به کار می‌رود تهیه توزیع فراوانی است.

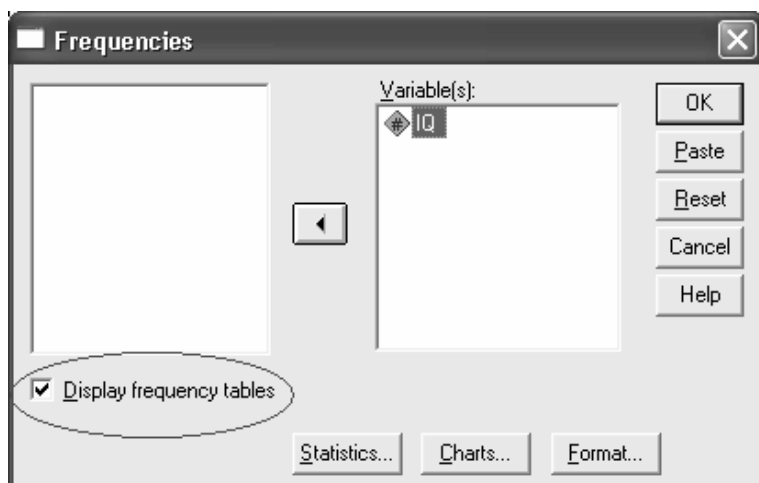
پس از جمع‌آوری داده‌های آماری لازم است آن‌ها را به گونه‌ای تنظیم و سازمان‌دهی شوند که به توان اطلاعات مورد نیاز را به روشنی به دست آورد و درک کرد. برای این منظور باید داده‌ها را به صورت توزیع فراوانی تنظیم شوند. در واقع، توزیع فراوانی عبارتند از سازمان دادن به اندازه‌ها یا مشاهده‌ها به وسیله در آوردن آن‌ها در قالب طبقه‌ها همراه با ذکر فراوانی هر طبقه. تفسیر نمره‌ها در جدول توزیع فراوانی آسان‌تر از موقعی است که نمره‌ها سازمان‌بندی نشده‌اند در این جدول به آسانی می‌توان بزرگ‌ترین و کوچک‌ترین نمره را مشخص کرد.

برای ترسیم جدول توزیع فراوانی در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies....



بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



با انجام مراحل مذکور نتایج به صورت خروجی زیر ظاهر خواهد شد:

نمرات بهره هوشی

IQ

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	105	4	8.0	8.0	8.0
	106	6	12.0	12.0	20.0
	110	4	8.0	8.0	28.0
	112	4	8.0	8.0	36.0
	113	7	14.0	14.0	50.0
	120	4	8.0	8.0	58.0
	123	4	8.0	8.0	66.0
	124	6	12.0	12.0	78.0
	125	6	12.0	12.0	90.0
	128	5	10.0	10.0	100.0
Total		50	100.0	100.0	

فراوانی مطلق^۱

بنابراین ساده‌ترین راه برای تنظیم و جدول‌بندی این است که وجوه مختلف متغیر را به‌صورت طبقات مشخص کنیم و تکرار هر وجه را که فراوانی می‌نامند در مقابل طبقه مربوطه بنویسیم:

برونداد فراوانی بهره هوشی

۱۰

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	105	4	8.0	8.0	8.0
	106	6	12.0	12.0	20.0
	110	4	8.0	8.0	28.0
	112	4	8.0	8.0	36.0
	113	7	14.0	14.0	50.0
	120	4	8.0	8.0	58.0
	123	4	8.0	8.0	66.0
	124	6	12.0	12.0	78.0
	125	6	12.0	12.0	90.0
	128	5	10.0	10.0	100.0
	Total	50	100.0	100.0	

فراوانی تراکمی درصدی^۲

گاهی لازم است جدول به گونه‌ای تدوین و تنظیم شود که نشان دهد چند درصد افراد مورد مطالعه بالاتر یا پایین‌تر از یک عدد معین قرار دارند. چنین اطلاعاتی را می‌توان از جدول توزیع فراوانی تراکمی درصدی، به‌دست آورد. برای محاسبه فراوانی تراکمی درصدی هر طبقه فراوانی تراکمی هر طبقه را به مجموع اعداد تقسیم می‌کنیم و سپس حاصل را در ۱۰۰ ضرب می‌کنیم (دلاور، ۷۱:۱۳۸۰)

$$\frac{cf}{n} \times 100$$

1. Frequency
2. Cumulative Percent

جدول زیر فراوانی تراکمی درصدی را نشان می‌دهد:

جدول فراوانی تجمعی بهره هوشی

۱۰

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	105	4	8.0	8.0	8.0
	106	6	12.0	12.0	20.0
	110	4	8.0	8.0	28.0
	112	4	8.0	8.0	36.0
	113	7	14.0	14.0	50.0
	120	4	8.0	8.0	58.0
	123	4	8.0	8.0	66.0
	124	6	12.0	12.0	78.0
	125	6	12.0	12.0	90.0
	128	5	10.0	10.0	100.0
Total		50	100.0	100.0	

با در نظر گرفتن فراوانی تجمعی می‌توان یک دید کلی از وضع گروه در آزمون به‌دست آورد چنانچه در جدول معلوم است ۵۸٪ دانش‌آموزان بهره هوشی کمتر از ۱۲۰ و ۹۰٪ دانش‌آموزان شرکت‌کننده در پژوهش بهره هوشی کمتر از ۱۲۵ دارند.

فراوانی نسبی درصدی^۱

معمولاً ذکر فراوانی مطلق در جدول به تنهایی کافی نیست و چنانچه نسبت یا درصد فراوانی محاسبه شده نوشته شود اطلاعات جدول گویاتر خواهد بود. برای محاسبه درصد کافی است که فراوانی مطلق هر طبقه را به تعداد کل تقسیم کنیم و نسبت محاسبه شده هر طبقه را در ۱۰۰ ضرب نماییم:

$$\frac{f}{n} \times 100$$

فراوانی درصدی بهره هوشی

۱۰

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	10 5	4	8.0	8.0	8.0
	10 6	6	12.0	12.0	20.0
	11 0	4	8.0	8.0	28.0
	11 2	4	8.0	8.0	36.0
	11 3	7	14.0	14.0	50.0
	12 0	4	8.0	8.0	58.0
	12 3	4	8.0	8.0	66.0
	12 4	6	12.0	12.0	78.0
	12 5	6	12.0	12.0	90.0
	12 8	5	10.0	10.0	100.0
Total		50	100.0	100.0	

تهیه جدول برحسب درصد، این حسن را دارد که اگر بخواهیم داده‌های دو یا چند جدول مختلف را با هم مقایسه کنیم این ارقام گویاتر و معتبرتر از فراوانی مطلق است. مثال: اگر نمرات درس آمار دو کلاس که یکی ۲۰ نفر دانش‌آموز و کلاس دیگر ۳۰ نفر دانشجو دارد در نظر بگیریم.

و معلوم شود که در هر دو کلاس ۱۱ نفر نمره قبولی (۲۰-۱۰) کسب کرده‌اند به ظاهر چنین به نظر می‌رسد که تعداد قبولی هر دو کلاس یکسان است اما اگر تعداد دانش‌آموزان این دو کلاس را در رابطه با جمعیت کلاس در نظر بگیریم و درصد آن‌ها را محاسبه کنیم معلوم می‌شود که نسبت یا درصد قبولی درس آمار دانش‌آموزان کلاس یک (۶۳/۳٪) بیشتر از کلاس دو (۴۵٪) است.

student1

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	10-20	11	36.7	36.7	36.7
	0-9	19	63.3	63.3	100.0
Total		30	100.0	100.0	

student2

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 10-20	11	55.0	55.0	55.0
0-9	9	45.0	45.0	100.0
Total	20	100.0	100.0	

نمودار فراوانی

استفاده از نمودارها روش مناسب جهت تفهیم بهتر مطلب به دیگران است. گاهی مفهومی را که با دیدن یک تصویر درک می‌کنیم بیش از مطلبی است که از راه خواندن به‌دست می‌آوریم. بنابراین نمودارها روش مناسبی برای توصیف و نمایش داده‌های جمع شده است و با آن می‌توان تصویر روشن‌تری (در مقایسه با جدول توزیع فراوانی) از اطلاعات گردآوری شده به‌دست آورد و ارائه توزیع فراوانی به‌صورت نمودار کار تحلیل و تفسیر داده‌ها را آسان می‌سازد.

در زیر به معرفی نمودارهای معروف مورد استفاده برای نشان دادن توزیع فراوانی داده‌ها می‌پردازیم:

(۱) هیستوگرام (۲) ستونی (۳) دایره‌ای

الف) نمودار هیستوگرام^۱

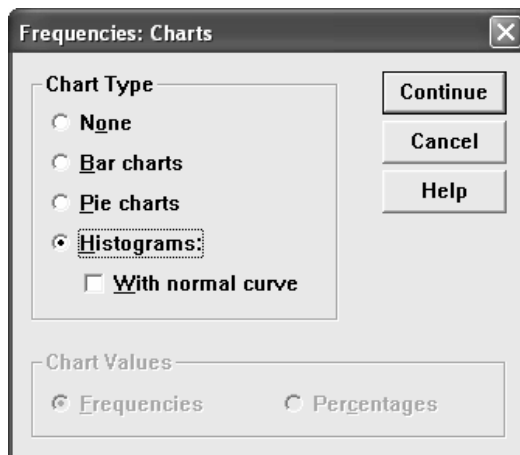
در نمودار هیستوگرام، توزیع فراوانی داده‌ها با استفاده از محورهای مختصات نشان داده می‌شود برای ترسیم نمودار هیستوگرام، فاصله طبقات را بر روی محور افقی و فراوانی‌ها را بر روی محور عمودی مشخص می‌کنیم، سپس بر روی هر طبقه ستون یا مربع مستطیلی می‌سازیم به‌گونه‌ای که عرض آن برابر با فاصله طبقه و ارتفاعش به اندازه فراوانی همان طبقه باشد. در بعضی مواقع از نمودار هیستوگرام برای نمایش متغیرهایی که با استفاده از مقیاس‌های فاصله‌ای و نسبی اندازه‌گیری شده‌اند، به‌کار برده می‌شود.

مثال: نمره‌های هوش ۵۰ دانش‌آموز در زیر آمده است ، نمودار هیستوگرام را رسم کنید؟

۱۲۰	۱۰۵	۱۱۳	۱۲۰	۱۲۵
۱۱۲	۱۲۵	۱۲۵	۱۱۲	۱۲۴
۱۱۰	۱۲۴	۱۲۴	۱۱۰	۱۱۳
۱۲۳	۱۱۳	۱۱۳	۱۲۳	۱۰۶
۱۰۵	۱۰۶	۱۰۶	۱۰۶	۱۲۸
۱۲۵	۱۲۸	۱۲۸	۱۲۸	۱۰۵
۱۲۴	۱۲۰	۱۰۵	۱۲۰	۱۲۵
۱۱۳	۱۱۲	۱۲۵	۱۱۲	۱۲۴
۱۰۶	۱۱۰	۱۲۴	۱۱۰	۱۱۳
۱۲۸	۱۲۳	۱۱۳	۱۲۳	۱۰۶

برای ترسیم نمودار هیستوگرام در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies...
Charts...
Histograms



با انجام مراحل مذکور نتایج به صورت خروجی زیر ظاهر خواهد شد:

Frequencies

Statistics

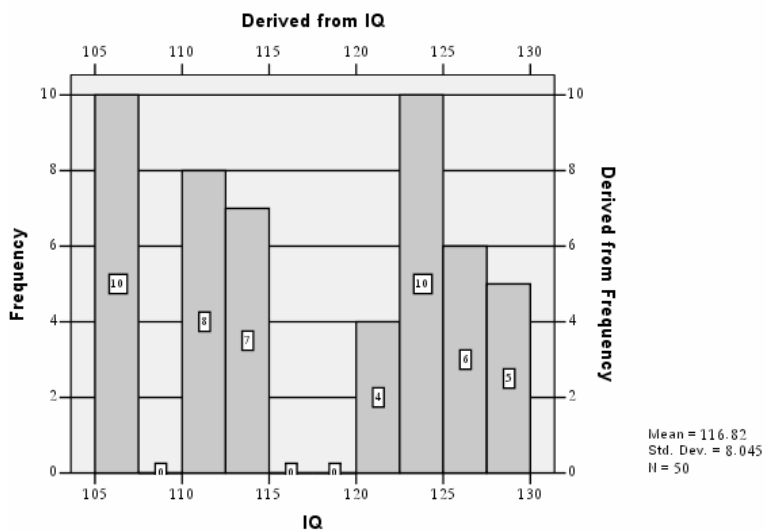
IQ

N	Valid	50
	Missing	0

IQ

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	105	4	8.0	8.0	8.0
	106	6	12.0	12.0	20.0
	110	4	8.0	8.0	28.0
	112	4	8.0	8.0	36.0
	113	7	14.0	14.0	50.0
	120	4	8.0	8.0	58.0
	123	4	8.0	8.0	66.0
	124	6	12.0	12.0	78.0
	125	6	12.0	12.0	90.0
	128	5	10.0	10.0	100.0
	Total	50	100.0	100.0	

Histogram



ب) نمودار ستونی^۱

هر زمانی که با داده‌های مربوط به مقیاس‌های اسمی، ترتیبی، فاصله‌ای و نسبی سروکار داریم می‌توان از نمودار ستونی یا میله‌ای استفاده کرد. به هنگام رسم نمودار ستونی برای داده‌های حاصل از مقیاس‌های اسمی دو نکته باید در نظر گرفت:

نکته اول: در ترسیم این نوع نمودارها هیچ گونه ترتیبی در نظر گرفته نمی‌شود یعنی طبقات مختلف را می‌توان بر روی محور افقی با هر ترتیب دلخواهی به دنبال هم قرار داد.

نکته دوم: در رسم این نوع نمودارها بر خلاف هیستوگرام، ستونها یا میله‌ها را باید از هم جدا کرد تا نشان دهد که هیچ گونه پیوستگی بین داده‌ها وجود ندارد.

مثال: در پژوهشی که به منظور مشخص نمودن بهره‌هوشی دانش‌آموزان انجام شد که در این پژوهش ۱۹ دانش‌آموز پسر و ۳۱ دانش‌آموز دختر شرکت داشته‌اند. اکنون با توجه به اطلاعات مربوط، نمودار ستونی آن را رسم کنید.

برای ترسیم نمودار ستونی در Spss مراحل زیر را انجام دهید:

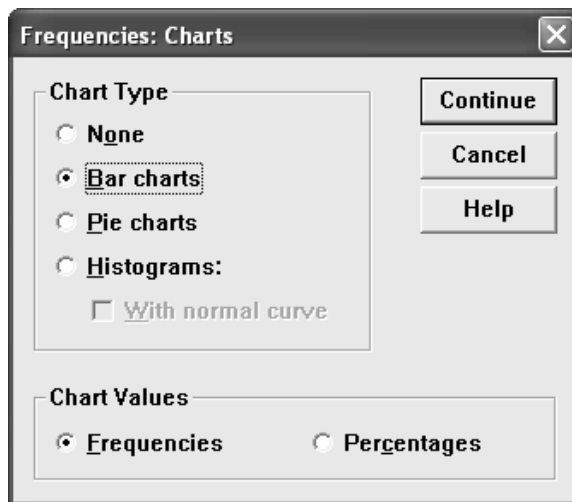
Analyze

Descriptive Statistics

Frequencies...

Charts...

Bar charts



با انجام مراحل مذکور نتایج به صورت خروجی زیر ظاهر خواهد شد:

Frequencies

Statistics

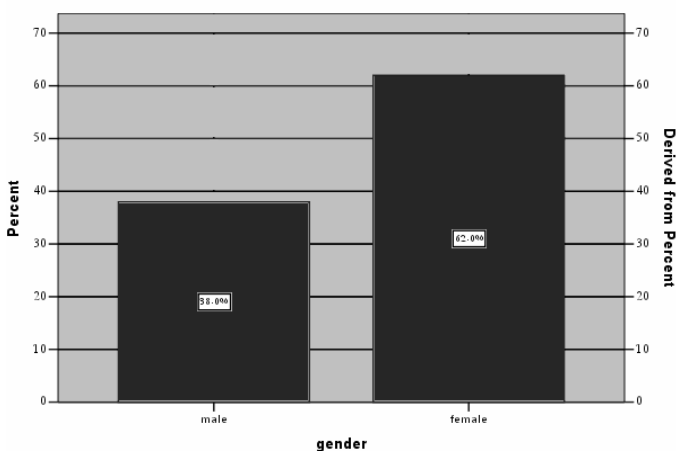
gender

N	Valid	50
	Missing	0

gender

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	male	19	38.0	38.0	38.0
	female	31	62.0	62.0	100.0
	Total	50	100.0	100.0	

gender



ج) نمودار دایره‌ای^۱

این نوع نمودار معمولاً برای نشان دادن رابطه اجزاء با کل داده‌ها به کار می‌رود. برای رسم آن ابتدا یک دایره رسم کنید (محیط و سطح آن نشانگر کل (۱۰۰٪) است) سپس محیط دایره (۳۶۰ درجه) را بر حسب فراوانی مطلق (f) یا فراوانی درصدی (P%) هر

طبقه یا بخشی از داده‌ها، تقسیم می‌کنیم با مشخص کردن قطره‌های به‌دست آمده به صورت رنگی و یا هاشور زده با یک نگاه، شمایی کامل از رابطه اجزاء با کل نمایان می‌شود. (پاشا شریفی و نجفی زند، ۱۳۷۵: ۵۱)

مثال: نمره‌های هوش ۵۰ دانش‌آموز در زیر آمده است، نمودار دایره‌ای را رسم کنید

۱۲۰	۱۰۵	۱۱۳	۱۲۰	۱۲۵
۱۱۲	۱۲۵	۱۲۵	۱۱۲	۱۲۴
۱۱۰	۱۲۴	۱۲۴	۱۱۰	۱۱۳
۱۲۳	۱۱۳	۱۱۳	۱۲۳	۱۰۶
۱۰۵	۱۰۶	۱۰۶	۱۰۶	۱۲۸
۱۲۵	۱۲۸	۱۲۸	۱۲۸	۱۰۵
۱۲۴	۱۲۰	۱۰۵	۱۲۰	۱۲۵
۱۱۳	۱۱۲	۱۲۵	۱۱۲	۱۲۴
۱۰۶	۱۱۰	۱۲۴	۱۱۰	۱۱۳
۱۲۸	۱۲۳	۱۱۳	۱۲۳	۱۰۶

برای ترسیم نمودار دایره‌ای در Spss مراحل زیر را انجام دهید:

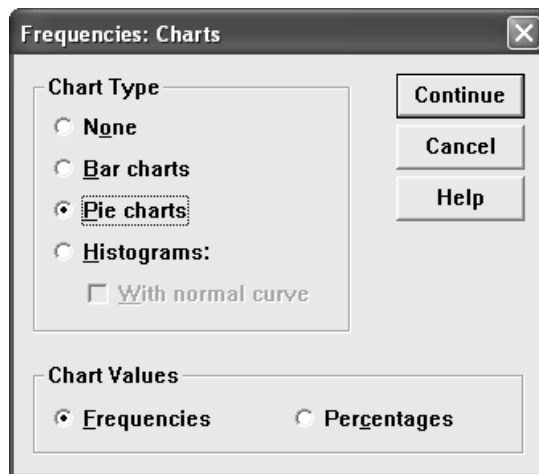
Analyze

Descriptive Statistics

Frequencies...

Charts...

Pie charts



با انجام مراحل مذکور نتایج به صورت خروجی زیر ظاهر خواهد شد:

Frequencies

Statistics

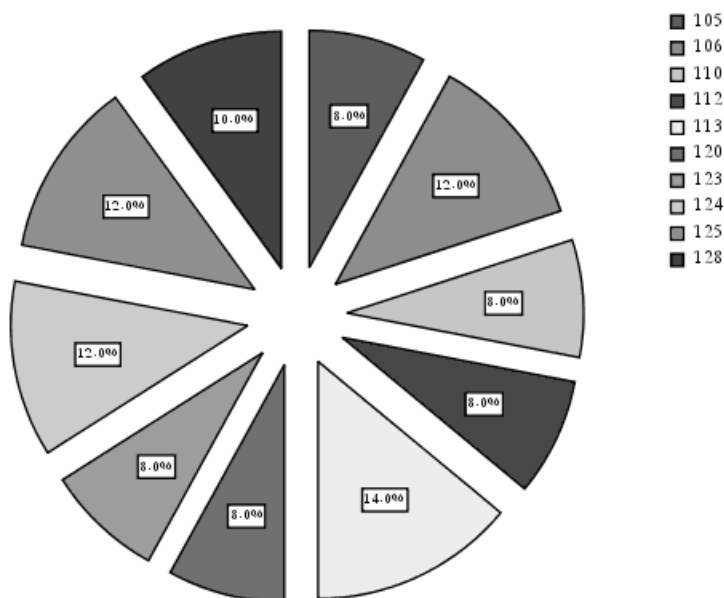
IQ

N	Valid	50
	Missing	0

IQ

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	105	4	8.0	8.0	8.0
	106	6	12.0	12.0	20.0
	110	4	8.0	8.0	28.0
	112	4	8.0	8.0	36.0
	113	7	14.0	14.0	50.0
	120	4	8.0	8.0	58.0
	123	4	8.0	8.0	66.0
	124	6	12.0	12.0	78.0
	125	6	12.0	12.0	90.0
	128	5	10.0	10.0	100.0
Total		50	100.0	100.0	

درصد ثمرات هوش دانش آموزان مورد مطالعه



همان‌طور که در انتهای فصل اشاره کردیم بعد از اتمام هر عملی در پروژه خود، تغییرات انجام شده را ذخیره نمایید برای ذخیره تغییرات در یک فایل موجود گزینه Save را از فایل یا از میله‌ابزار (Tool bar) انتخاب نمایید تا تغییرات در فایل مربوط، ذخیره شود تا بعداً بتوانید به این نتایج دسترسی داشته باشید. در صورتی که تغییرات را ذخیره نکنید اعمال انجام شده پس از خروج از فایل پروژه از بین می‌رود و قابل دسترسی نخواهد بود.

خود آزمایی

۱. نمره‌های زیر نتیجه آزمون بسندگی زبان انگلیسی دانشجویان است:

۴۲,۶۷,۸۴,۵۲,۴۵,۶۷,۷۶,۶۳,۵۰,۹۷,۷۱,۵۷
۵۶,۹۹,۴۸,۵۱,۶۹,۹۰,۶۱,۶۲,۵۶,۹۷,۶۸,۷۷,۵۹
۴۲,۶۷,۳۶,۵۲,۴۵,۶۷,۷۶,۵۰,۶۳,۹۷,۷۲,۴۹,۷۱,۵۷
۵۲,۲۹,۴۸,۵۱,۶۹,۹۰,۷۱,۶۲,۵۶,۹۷,۶۸,۷۷,۵۹
۴۲,۶۷,۸۴,۵۲,۴۵,۲۷,۷۶,۵۰,۲۳,۹۷,۷۲,۴۹,۷۱,۵۷
۷۲,۵۶,۹۹,۴۸,۵۱,۶۹,۴۰,۶۱,۶۲,۵۶,۹۷,۶۸,۷۷,۵۹

الف) جدول توزیع فراوانی تشکیل دهید.

ب) فراوانی نسبی درصدی و فراوانی تراکمی درصدی محاسبه کنید.

ج) نمودارهای زیر را رسم کنید.

- هیستوگرام
- دایره‌ای

۲. پژوهشگری نمونه‌ای از کارکنان یک منطقه آموزش و پرورش را انتخاب و آن‌ها را در سطح آمادگی، دبستان، راهنمایی و دبیرستان تقسیم‌بندی نمود، سپس نظر آنان را در مورد تشکیل انجمن معلمان سؤال و به‌صورت جدول زیر تنظیم کرد. (دلاور، ۵۰۱، ۱۳۸۰)

دبیرستان	راهنمایی	دبستان	آمادگی	
۲۰	۲۵	۲۰	۶	موافق
۱۰	۱۴	۱۵	۱۶	مخالف

الف) جدول توزیع فراوانی تشکیل دهید.

ب) فراوانی نسبی درصدی و فراوانی تراکمی درصدی محاسبه کنید.

ج) نمودار ستونی (میله‌ای) رسم کنید.

فصل چهارم

شاخص‌های مرکزی

هدف‌های یادگیری

۱. سه شاخص مرکزی که دارای بیشترین استفاده است را نام ببرد.
۲. نمای یک توزیع فراوانی را تعریف و تعیین کند.
۳. میانه را در مجموعه‌ای از اعداد محاسبه و تعیین کند.
۴. میانگین را در مجموعه‌ای از اعداد محاسبه و تعیین کند.
۵. ارتباط بین شاخص‌های مرکزی در توزیع‌های متقارن و نامتقارن بیان کند.
۶. ویژگی‌های شاخص‌های مرکزی را بیان کند.

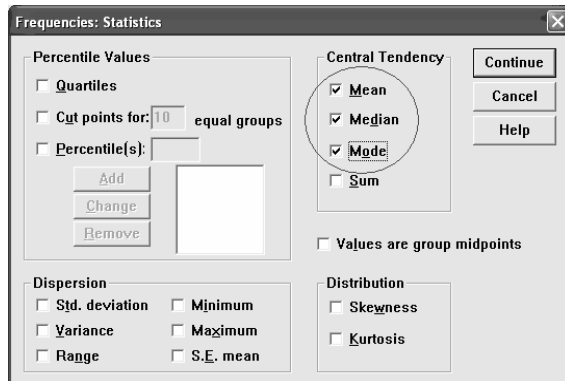
شاخص‌های مرکزی

اگر پژوهشگری بخواهد، با استناد به یک نمره واحد، چگونگی تمرکز نمرات یک گروه (نمونه) را مشخص کند از اندازه‌های گرایش مرکزی استفاده می‌کند. بنابراین شاخص‌های گرایش مرکزی، شاخص‌هایی هستند که با استفاده از آن‌ها مجموعه‌ای از اطلاعات در یک اندازه یا عدد که نماینده آن مجموعه است خلاصه می‌شود. این شاخص‌ها عبارتند از نما، میانه و میانگین.

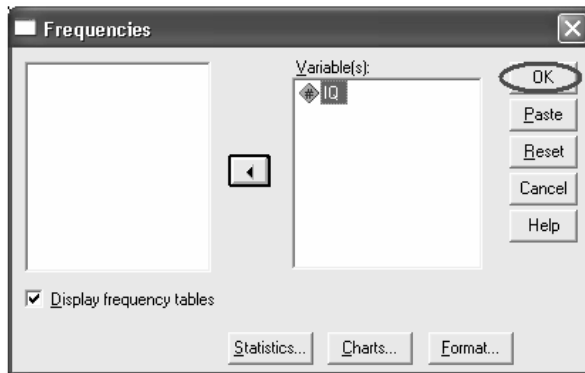
برای محاسبه شاخص‌های مرکزی در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies...
Statistics...
Mode/Mean/Median

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده گزینه (Mean/Median/Mode) را انتخاب کنید و سپس (continue) را کلیک کنید با کلیک بر روی گزینه (Continue) به پنجره (Frequencies) باز می‌گردیم:



سپس در پنجره (Frequencies) گزینه (OK) را کلیک کنید و نتایج آمار توصیفی (نما، میانه و میانگین) ظاهر خواهد شد. اکنون به منظور آشنایی بیشتر به معرفی هر یک از شاخص‌های مذکور می‌پردازیم.

نما^۱

نما در یک مجموعه داده‌ها، به عددی گفته می‌شود که بیشترین تکرار یا فراوانی را دارا باشد. برای این کار ابتدا میانه (Mn) و میانگین ($\bar{X}$) را محاسبه و سپس با قرار دادن مقادیر به‌دست آمده در فرمول مقدار (MO) مشخص می‌شود. (پاشا شریفی نجفی زند، ۱۳۷۵)

$$MO = 3Mn - 2\bar{X}$$

مثال: نمره‌های هوش ۵۰ دانش‌آموز در زیر آمده است معرف بهره هوشی دانش‌آموزان شرکت‌کننده در تحقیق کدام نمره است:

۱۲۰	۱۰۵	۱۱۳	۱۲۰	۱۲۵
۱۱۲	۱۲۵	۱۲۵	۱۱۲	۱۲۴
۱۱۰	۱۲۴	۱۲۴	۱۱۰	۱۱۳
۱۲۳	۱۱۳	۱۱۳	۱۲۳	۱۰۶
۱۰۵	۱۰۶	۱۰۶	۱۰۶	۱۲۸
۱۲۵	۱۲۸	۱۲۸	۱۲۸	۱۰۵
۱۲۴	۱۲۰	۱۰۵	۱۲۰	۱۲۵
۱۱۳	۱۱۲	۱۲۵	۱۱۲	۱۲۴
۱۰۶	۱۱۰	۱۲۴	۱۱۰	۱۱۳
۱۲۸	۱۲۳	۱۱۳	۱۲۳	۱۰۶

برای محاسبه نما در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies...
Statistics...
Mode

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:

Frequencies: Statistics

Percentile Values

☐ **Quartiles**

☐ **Cut points for:** 10 **equal groups**

☐ **Percentile(s):**

Central Tendency

☐ **Mean**

☐ **Median**

☒ **Mode**

☐ **Sum**

☐ **Values are group midpoints**

Dispersion

☐ **Std. deviation** ☐ **Minimum**

☐ **Variance** ☐ **Maximum**

☐ **Range** ☐ **S.E. mean**

Distribution

☐ **Skewness**

☐ **Kurtosis**

در پنجره ظاهر شده گزینه (Mode) را انتخاب و گزینه (Continue) را کلیک کنید و سپس در پنجره (Frequencies) گزینه (OK) را کلیک نموده تا نتایج به صورت خروجی زیر مشاهده گردد:

Frequencies

Statistics

IQ

N	Valid	50
	Missing	0
Mode		113

IQ

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	105	4	8.0	8.0	8.0
	106	6	12.0	12.0	20.0
	110	4	8.0	8.0	28.0
	112	4	8.0	8.0	36.0
	113	7	14.0	14.0	50.0
	120	4	8.0	8.0	58.0
	123	4	8.0	8.0	66.0
	124	6	12.0	12.0	78.0
	125	6	12.0	12.0	90.0
	128	5	10.0	10.0	100.0
	Total	50	100.0	100.0	

چنان‌که مشاهده می‌شود طبقه نمایی، طبقه ۱۱۳ است یعنی نمره هوش اکثریت دانش‌آموزان شرکت‌کننده در تحقیق، ۱۱۳ است.

میانه^۱

نقطه‌ای است در روی مقیاس نمرات یا توزیع نمرات که نصف بالا و نصف پایین نمرات را از یکدیگر جدا می‌کند به عبارت دیگر میانه نقطه‌ای است که ۵۰ درصد نمرات در بالای آن و ۵۰ درصد نمرات در پایین آن قرار دارد.

بنابراین میانه عددی است که جمعیت مورد مطالعه را از لحاظ توانایی یا گروهی مورد سنجش به دو نیمه تقسیم می‌کند.

$$md = L + \frac{\frac{N}{2} - cf}{f}(i)$$

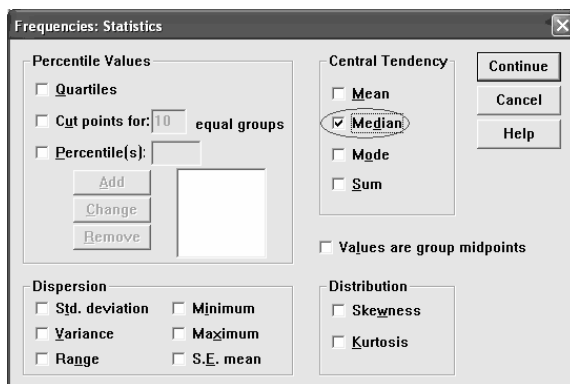
مثال: اضافه کار ۲۰ کارمند بر حسب تومان به ترتیب زیر است میانه دستمزدها چقدر است؟

۲۹۰-۲۶۰-۲۵۰-۲۰۰-۱۹۰-۱۸۰-۱۵۰-۱۵۰-۱۵۰-۱۴۰
۱۵۰-۳۰۰-۳۰۰-۳۰۰-۳۰۰-۱۹۰-۱۹۰-۲۵۰-۲۵۰-۲۵۰

برای محاسبه میانه در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies...
Statistics...
Median

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده گزینه (Median) را انتخاب و گزینه (Continue) را کلیک کنید سپس در پنجره (Frequencies) گزینه (OK) را کلیک نموده تا نتایج به صورت خروجی زیر مشاهده گردد:

Statistics

اضافه کار

N	Valid	20
	Missing	0
Median		225.0000

نقطه میانه عدد ۲۲۵ تومان است یعنی ۵۰ درصد از کارمندان اضافه کار بالای ۲۲۵ و ۵۰ درصد دیگر اضافه کار زیر ۲۲۵ تومان دارند.

میانگین^۱

کمیتی است که مقدار متوسط و یا به عبارت دیگر مرکز ثقل داده‌های به دست آمده را نشان می‌دهد.

میانگین از طریق جمع کردن تمام نمرات و تقسیم حاصل جمع بر تعداد کل نمره‌ها، به دست می‌آید. (دلاور، ۱۳۸۰: ۱۱۸)

$$\bar{X} = \frac{\sum X}{N}$$

مثال: نمرات آزمون ریاضی ۱۰ دانش‌آموز کلاس اول دبستان در زیر آمده است میانگین نسبی این دانش‌آموزان چقدر است؟

نمره ریاضی	دانش‌آموز
۱۹	۱
۱۳	۲
۱۷	۳
۱۶	۴
۱۲	۵
۱۸	۶
۱۶	۷
۱۵	۸
۱۱	۹
۲۰	۱۰

برای محاسبه میانگین در Spss مراحل زیر را انجام دهید:

Analyze
 Descriptive Statistics
 Frequencies...
 Statistics...
 Mean

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:

Frequencies: Statistics

Percentile Values

☐ **Quartiles**

☐ **Cut points for:** 10 equal groups

☐ **Percentile(s):**

Central Tendency

☒ **Mean**

☐ **Median**

☐ **Mode**

☐ **Sum**

☐ **Values are group midpoints**

Dispersion

☐ **Std. deviation**

☐ **Variance**

☐ **Range**

☐ **Minimum**

☐ **Maximum**

☐ **S.E. mean**

Distribution

☐ **Skewness**

☐ **Kurtosis**

در پنجره ظاهر شده گزینه (Mean) را انتخاب نمایید و سپس (Continue) را کلیک کنید و سپس در پنجره (Frequencies) گزینه (OK) را کلیک نموده تا نتایج به صورت زیر آشکار گردد:

Frequencies

Statistics

math

N	Valid	10
	Missing	0
Mean		15.70

math

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid	11	1	10.0	10.0
	12	1	10.0	20.0
	13	1	10.0	30.0
	15	1	10.0	40.0
	16	2	20.0	60.0
	17	1	10.0	70.0
	18	1	10.0	80.0
	19	1	10.0	90.0
	20	1	10.0	100.0
Total	10	100.0	100.0	

مقایسه شاخص‌های مرکزی و کاربرد آنها

مقادیر گرایش مرکزی و کاربرد آنها در مورد داده‌های با مقیاس فاصله‌ای و نسبی تا اندازه زیادی با چگونگی توزیع داده‌ها بستگی دارد. توزیع داده‌ها ممکن است به صورت پراکندگی بهنجار (طبیعی) باشد و یا نباشد در صورتی که اندازه‌های گرایش مرکزی را برای توزیع فراوانی جدول زیر محاسبه کنیم ارزش‌های مختلفی برای مقایسه آنها وجود دارد.

۱۸	۱۷	۱۶
۲۰	۱۷	۱۴
۱۶	۱۶	۱۴
۱۶	۱۷	۱۶
۱۸	۱۶	۱۴
۱۷	۱۵	۱۲
۱۸	۱۶	۱۶
۱۸	۱۶	۱۲

برای بررسی اندازه‌های گرایش مرکزی در Spss مراحل زیر را انجام دهید:

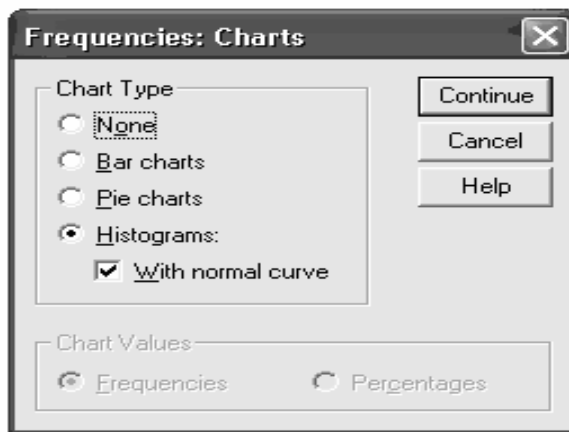
Analyze

Descriptive Statistics

Frequencies...

Charts...

Histograms



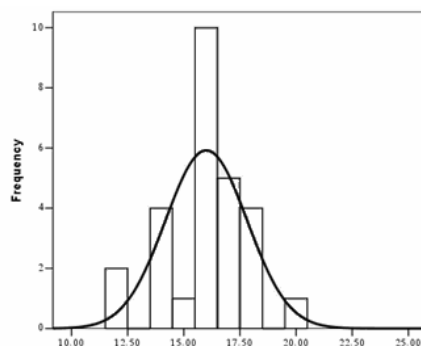
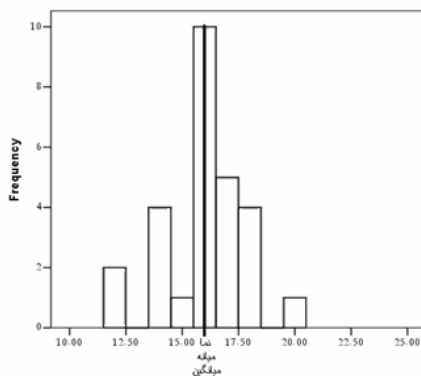
با انجام مراحل مذکور نتایج به صورت زیر آشکار می‌گردد:

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	12.00	2	7.4	7.4	7.4
	14.00	4	14.8	14.8	22.2
	15.00	1	3.7	3.7	25.9
	16.00	10	37.0	37.0	63.0
	17.00	5	18.5	18.5	81.5
	18.00	4	14.8	14.8	96.3
	20.00	1	3.7	3.7	100.0
Total		27	100.0	100.0	

Statistics

VAR00001

N	Valid	27
	Missing	7
Mean		16.0000
Median		16.0000
Mode		16.00



چنان‌که ملاحظه می‌شود در توزیع بهنجار میانگین، میانه و نما با هم برابرند بنابراین ارزش کاربرد هر سه آن‌ها به عنوان شاخص گرایش مرکزی یکسان است. اما اگر پراکندگی داده‌ها بهنجار نباشد در این صورت میانگین، میانه و نما برابر نخواهد بود و گفته می‌شود که منحنی دارای کجی یا کشیدگی است.

اندازه‌های گرایش مرکزی برای توزیع فراوانی جدول زیر یک کجی منفی، است میانگین کوچک‌تر از میانه و نما است. زیرا نمرات خیلی پایین، میانگین را به سمت پایین کشیده‌اند.

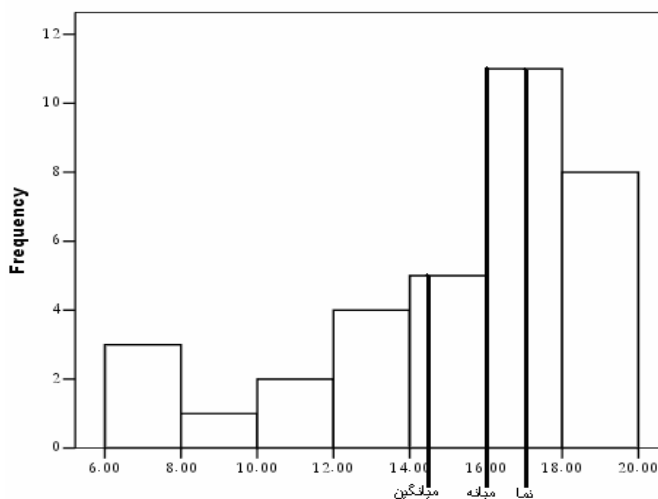
۲۰	۱۷	۱۴	۹
۲۰	۱۷	۱۴	۷
۲۰	۱۷	۱۴	۷
۱۸	۱۷	۱۴	۷
۱۸	۱۶	۱۲	
۱۸	۱۶	۱۲	
۱۸	۱۶	۱۲	
۱۸	۱۶	۱۲	
۱۷	۱۶	۱۱	
۱۷	۱۴	۱۰	

Statistics

VAR00001

N	Valid	34
	Missing	0
Mean		17.0809
Median		16.7500
Mode		16.00

Histogram



اندازه‌های گرایش مرکزی برای توزیع فراوانی جدول زیر یک کجی مثبت است
نما کوچک‌تر از میانه و میانگین است.

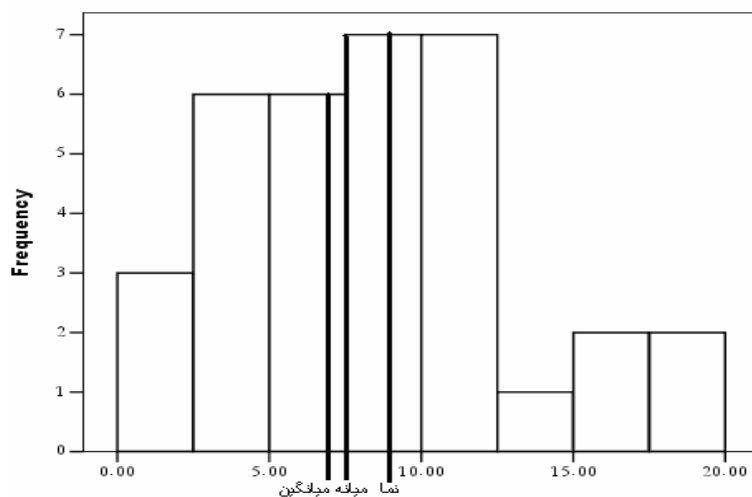
۱۷/۵	۰	۸	۹
۶	۱۲	۱۰	۹
۹	۹	۲	۷
۱۷	۱۲	۱۰	۷
۱۶	۳/۲۵	۴/۲۵	
۶	۱۱	۲/۵	
۷	۱۱	۸/۵	
۶	۳	۱۹/۲۵	
۱۴	۱۰	۹	
۲	۹	۳	

Statistics

VAR00003

N	Valid	34
	Missing	0
Mean		8.3309
Median		8.7500
Mode		9.00

Histogram



در زیر به طور خلاصه کاربرد هر یک از شاخص های مرکزی ذکر شده است:

الف) میانگین

۱. هنگامی که داده ها از مقیاس فاصله ای و نسبی هستند.
۲. توزیع داده ها بهنجار است.
۳. بخواهیم مرکز ثقل داده ها را بدانیم.
۴. هنگامی که معتبرترین اندازه گرایش مرکزی مورد نیاز است.
۵. زمانی که پژوهشگر علاقه مند باشد که ارزش عددی همه نمره ها در محاسبه دخالت داشته باشد.

ب) میانه

۱. داده ها دارای مقیاس رتبه ای باشند.
۲. پراکندگی داده ها دارای کجی زیاد است.
۳. وقتی به عددی که در وسط توزیع قرار دارد، نیاز باشد.
۴. ثبات میانه از میانگین کمتر ولی از نما بیشتر است.

ج) نما (مد)

۱. وقتی داده ها دارای مقیاس اسمی است.
 ۲. بخواهیم یک برآورد فوری از شاخص گرایش مرکزی داشته باشیم.
 ۳. وقتی که پژوهشگر علاقه مند باشد که عددی را که بیشترین تکرار دارد پیدا کند.
- (دلاور، ۱۳۸۰: ۱۴۳)

خود آزمایی

۱. سه شاخص مرکزی که دارای بیشترین استفاده است را نام ببرید
۲. نما را تعریف و نمای نمرات زیر را تعیین کنید.

۲۰	۱۷	۱۴	۹
۲۰	۱۷	۱۴	۷
۲۰	۱۷	۱۴	۷
۱۸	۱۷	۱۴	۷
۱۸	۱۶	۱۲	
۱۸	۱۶	۱۲	
۱۸	۱۶	۱۲	
۱۸	۱۶	۱۲	
۱۷	۱۶	۱۱	
۱۷	۱۴	۱۰	

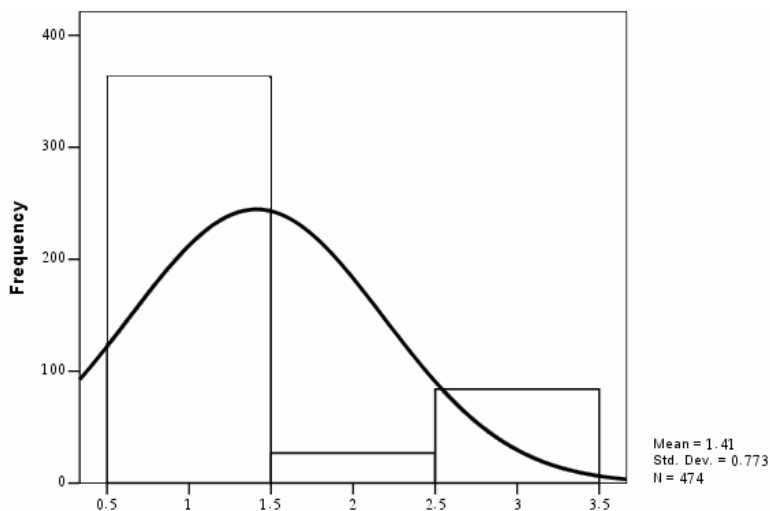
۳. میانه مجموعه اعداد زیر محاسبه کنید.

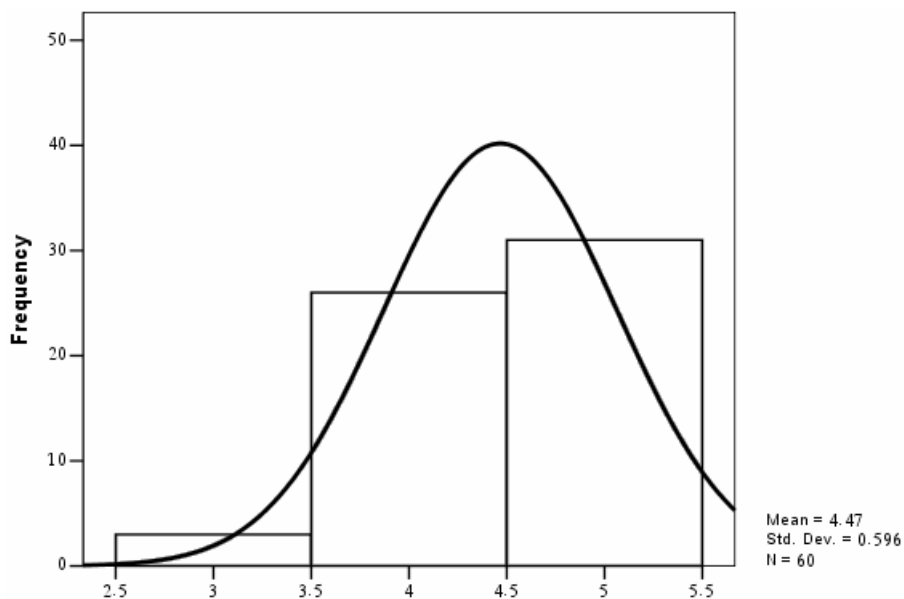
۱۶	۲۷	۱۴	۲۳	۱۶	۱۷	۱۵	۲۷	۲۵	۱۵	۲۴	۱۷
۲۳	۱۷	۱۶	۲۴	۲۷	۲۸	۲۵	۱۷	۱۵	۱۹	۲۷	۲۸
۲۴	۲۸	۲۷	۲۷	۱۷	۳۰	۲۶	۲۸	۱۹	۲۳	۱۷	۳۰
۲۷	۳۰	۱۷	۱۷	۲۸	۱۴	۱۶	۳۰	۲۳	۲۵	۲۸	۱۴
۱۷	۱۴	۲۸	۲۸	۳۰	۱۶	۲۳	۱۴	۲۸	۲۶	۳۰	۱۶
۲۸	۱۶	۳۰	۳۰	۱۴	۱۷	۲۴	۲۴	۳۰	۱۶	۱۴	۲۷

۴. در جدول زیر نمرات آمار دو گروه علوم تربیتی و روان‌شناسی ذکر شده است، میانگین دو گروه علوم تربیتی و روان‌شناسی را محاسبه و تعیین کنید.

روان‌شناسی		علوم تربیتی	
۱۰	۱۹	۱۵	۱۷
۱۲	۱۶/۲۵	۱۴	۱۹
۱۵	۱۴	۱۲	۱۶/۲۵
۱۶/۲۵	۲۰	۱۵	۱۰
۱۷	۲۰	۱۷	۱۲
۱۷	۱۲/۵	۱۸	۱۲/۵
۱۰	۱۴	۱۹	۱۷
۱۶	۱۲	۱۶	۱۸
۱۴	۱۹	۱۶/۲۵	۱۹
۲۰	۱۰	۲۰	۹

۵. ارتباط بین شاخص‌های مرکزی در توزیع‌های نامتقارن زیر بیان کنید.





۶. ویژگی‌های شاخص‌های مرکزی را بیان کنید.

فصل پنجم

شاخص‌های پراکندگی

هدف‌های یادگیری

- از دانشجو انتظار می‌رود که پس از خواندن فصل پنجم بتواند:
۱. شاخص‌های پراکندگی که دارای بیشترین استفاده است را نام ببرد.
 ۲. دامنه تغییرات را تعریف و تعیین کند.
 ۳. انحراف چارکی را در مجموعه‌ای از اعداد محاسبه و تعیین کند.
 ۴. واریانس را در مجموعه‌ای از اعداد محاسبه و تعیین کند.
 ۵. انحراف استاندارد را محاسبه و تفسیر کند.
 ۶. رتبه درصدی و دهک را محاسبه کند.
 ۷. ضریب تغییرپذیری (ضریب نسبی واریانس) را تعریف و محاسبه کند.
 ۸. کجی و کشیدگی را تعریف کند.
 ۹. ارتباط بین شاخص‌های مرکزی در توزیع‌های متقارن و نامتقارن بیان کند.

شاخص‌های پراکندگی

شاخص‌های مرکزی برای تعیین مقدار متوسط توزیع نمره‌ها به‌کار برده می‌شوند اما واقعی وجود دارد که نیازمند اطلاعات بیشتری هستیم. اگر بهره هوشی دو گروه با میانگین مساوی داشته باشیم و تنها اطلاع ما درباره میانگین آن‌ها باشد تنها مطلبی که درمورد این گروه‌ها می‌توانیم اظهار نظر کنیم، این است که، این دو گروه دارای میانگین مساوی هستند. اما پراکندگی آن‌ها متفاوت است مثلاً زمانی که معلم میانگین هوش کلاس را بداند کافی نیست، بلکه معلم باید پراکندگی بهره هوشی دانش‌آموزان که بین

۱۳۵-۸۵ یا بین ۱۱۰-۹۵ است را بدانند. زیرا روش تدریس و مواد آموزشی هر کدام از این دو گروه حتی اگر میانگین بهره هوشی مساوی داشته باشند، کاملاً فرق می‌کند بنابراین برای معرفی یک توزیع نه تنها به گرایش‌های مرکزی نیاز داریم بلکه در دست داشتن گرایش‌های پراکندگی نیز ضرورت دارد. (دلاور، ۱۳۸۰: ۱۵۴)

شاخص‌های پراکندگی، میزان پراکندگی یا تغییرات بین نمره‌های یک توزیع را نشان می‌دهد از مهم‌ترین شاخص‌های پراکندگی عبارتند از دامنه تغییرات، چارک‌ها، انحراف معیار و واریانس.

دامنه تغییرات^۱

که آن را با R مخفف Range نشان می‌دهند در واقع بیانگر تفاوت میان بزرگ‌ترین و کوچک‌ترین داده‌های حاصل در یک بررسی است:

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

مثال: فرض کنید دستمزد ۷ نفر از کارکنان دو سازمان که به‌طور تصادفی استفاده شده‌اند به شرح زیر است:

سازمان آموزش و پرورش	سازمان تأمین اجتماعی
۱۸۰,۵۰۰	۲۶۰,۰۰۰
۱۹۰,۰۰۰	۲۸۰,۰۰۰
۲۴۰,۰۰۰	۲۴۰,۰۰۰
۱۷۰,۵۰۰	۲۲۰,۰۰۰
۴۰۰,۵۰۰	۳۳۰,۰۰۰
۲۹۰,۰۰۰	۲۰۰,۰۰۰
۲۶۰,۵۰۰	۲۲۰,۰۰۰

برای محاسبه دامنه تغییرات در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies...
Statistics...
Range

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:

Frequencies: Statistics

Percentile Values

- ☐ Quartiles
- ☐ Cut points for: 10 equal groups
- ☐ Percentile(s):
- Add
- Change
- Remove

Central Tendency

- ☐ Mean
- ☐ Median
- ☐ Mode
- ☐ Sum

☐ Values are group midpoints

Dispersion

- ☐ Std. deviation
- ☐ Variance
- ☒ Range
- ☐ Minimum
- ☐ Maximum
- ☐ S.E. mean

Distribution

- ☐ Skewness
- ☐ Kurtosis

Continue
Cancel
Help

در پنجره ظاهر شده گزینه (Range) را انتخاب کنید و گزینه (Continue) را کلیک کنید سپس در پنجره (Frequencies) گزینه (OK) را کلیک نموده تا نتایج به صورت خروجی زیر مشاهده گردد:

Descriptive Statistics

	N	Range
تأمین اجتماعی	7	130000.00
آموزش و پرورش	7	230000.00
Valid N (listwise)	7	

ملاحظه می‌شود که پراکندگی دستمزدها در سازمان آموزش و پرورش بیشتر از سازمان تأمین اجتماعی است مقدار R هر قدر به صفر نزدیک باشد، همگنی بیشتر و هر قدر مقدار R بزرگ‌تر از صفر باشد بیانگر ناهمگنی دو گروه است.

چارک^۱

چارک‌ها نقاطی بر روی مقیاس اندازه‌گیری هستند که کلیه مشاهده‌ها یا نمره‌ها را به چهار قسمت مساوی تقسیم می‌کنند. به عبارت دیگر در صورتی که منحنی توزیع نمره‌ها را ترسیم و سطح زیر منحنی را به چهار قسمت مساوی تقسیم کنیم، هر قسمت یک

چهارم از نمره‌های توزیع را شامل خواهد شد. چارک اول نقطه‌ای در روی مقیاس اندازه‌گیری است که ۲۵ درصد نمره‌ها را از پایین جدا می‌کند. به چارک اول نقطه ۲۵ درصدی نیز گفته می‌شود. در محلی قرار دارد که درست یک چهارم اعداد از آن کوچک‌تر و سه‌چهارم اعداد از آن بزرگ‌تر هستند. چارک وسط برابر با میانه است یعنی نقطه‌ای که توزیع را به دو قسمت مساوی تقسیم می‌کند. چارک سوم نقطه‌ای است که سه‌چهارم نمره‌ها در زیر آن و بقیه در بالای آن واقع شده‌اند. به چارک سوم نقطه ۷۵ درصدی نیز گفته می‌شود. (دلاور، ۱۳۸۰: ۱۵۹)

محاسبه چارک اول (نقطه ۲۵ درصدی):

$$Q_1 = L + \frac{\frac{1N}{4} - cf}{f} (i)$$

محاسبه چارک دوم (نقطه ۵۰ درصدی):

$$Q_2 = L + \frac{\frac{2N}{4} - cf}{f} (i)$$

محاسبه چارک سوم (نقطه ۷۵ درصدی):

$$Q_3 = L + \frac{\frac{3N}{4} - cf}{f} (i)$$

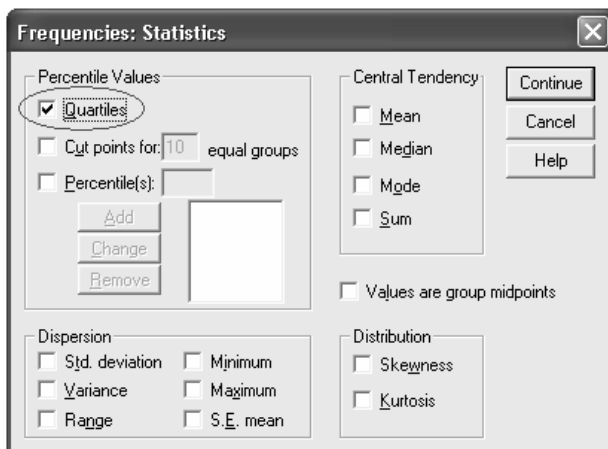
مثال: فرض کنید نمرات دو گروه از دانشجویان علوم تربیتی و روان‌شناسی در درس آمار در جدول زیر آمده است:

روان‌شناسی	علوم تربیتی
۱۰	۱۹
۱۲	۱۶/۲۵
۱۵	۱۴
۱۷	۲۰
۱۸	۱۲/۵۰
۱۹	۱۷
۱۶	۱۸
۱۴	۱۹
۲۰	۹

برای محاسبه دامنه تغییرات در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies...
Statistics...
Quartiles

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده گزینه (Range) را انتخاب کنید و گزینه (Continue) را کلیک کنید سپس در پنجره (Frequencies) گزینه (OK) را کلیک نموده تا نتایج به صورت خروجی زیر مشاهده گردد:

Statistics

		psychology	education
N	Valid	10	10
	Missing	0	0
Percentiles	25	13.50	13.6250
	50	15.50	17.0000
	75	18.25	19.0000

چنانچه در جدول فوق مشاهده می‌شود چارک‌های اول تا سوم برای رشته‌های روان‌شناسی و علوم تربیتی محاسبه شده است. در جدول زیر با توجه به اطلاعات به دست آمده به محاسبه دامنه تغییرات چارک‌ها و انحراف چارکی می‌پردازیم:

علوم تربیتی	روان شناسی	
۱۳/۶۲	۱۳/۵	چارک اول
۱۷	۱۵/۵	چارک دوم
۱۹	۱۸/۲۵	چارک سوم
$Q_3 - Q_1$ $۱۹ - ۱۳/۶۲ = ۵/۳۸$	$Q_3 - Q_1$ $۱۸/۵ - ۱۳/۵ = ۵$	دامنه تغییرات چارک‌ها
$Q = \frac{Q_3 - Q_1}{2}$ $\frac{۱۹ - ۱۳/۶۲}{2} ۲/۶۹$	$Q = \frac{Q_3 - Q_1}{2}$ $\frac{۱۸/۵ - ۱۳/۵}{2} = ۲/۵$	انحراف چارکی

واریانس و انحراف استاندارد (معیار)^۱

واریانس (پراکنش) و انحراف معیار بهترین اندازه‌های تغییرپذیری هستند واریانس یک شاخص پراکندگی است که از طریق انحراف نمره‌ها از میانگین محاسبه می‌شود و عبارتند از میانگین انحراف نمره‌ها از میانگین یا مجموع مجذورات نمره‌ها از میانگین تقسیم بر تعداد نمره‌ها. (دلاور، ۱۳۸۰: ۱۶۹).

محاسبه واریانس جامعه:

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

محاسبه واریانس نمونه:

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

واریانس، موارد استفاده فراوانی در آمار استنباطی دارد ولی استفاده از آن در آمار توصیفی محدود است زیرا با مجذور کردن انحراف نمره‌ها از میانگین، واریانس یا واحد اندازه‌گیری تغییر پیدا خواهد کرد. به این معنی که در صورتی که اندازه قد عده‌ای از دانشجویان را بر اساس سانتی‌متر اندازه‌گیری کنیم، واریانس محاسبه شده برحسب سانتی‌متر مربع خواهد بود همین مشکل اختلاف واحد اندازه‌گیری با واحد واریانس را می‌توان با جذر گرفتن از واریانس حل کرد این عمل جذر گرفتن از

واریانس موجب می‌شود که شاخص انحراف معیار (استاندارد) به‌دست آید. بنابراین انحراف استاندارد به‌عنوان ریشه دوم یا جذر و واریانس تعریف شده است. (دلاور، ۱۳۸۰: ۲۰۰)

محاسبه انحراف استاندارد جامعه:

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

محاسبه انحراف استاندارد نمونه:

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

مثال: در جدول زیر نمرات مربوط به آزمون آمار دو گروه دانشجویان علوم تربیتی و روان‌شناسی ذکر شده است، انحراف معیار و واریانس هر گروه چقدر است؟

روان‌شناسی	علوم تربیتی
۱۵	۱۷
۱۰	۱۹
۱۲	۱۶/۲۵
۱۵	۱۴
۱۷	۲۰
۱۸	۱۲/۵
۱۹	۱۷
۱۶	۱۸
۱۴	۱۹
۲۰	۹

برای محاسبه انحراف معیار و واریانس در Spss مراحل زیر را انجام دهید:

Analyze
 Descriptive Statistics
 Frequencies...
 Statistics...
 Std.deviation / Variance

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:

The screenshot shows the 'Frequencies: Statistics' dialog box. The 'Dispersion' section is highlighted with a circle, indicating that 'Std. deviation' and 'Variance' are selected. Other options in this section include 'Range', 'Minimum', 'Maximum', and 'S.E. mean'. The 'Central Tendency' section includes 'Mean', 'Median', 'Mode', and 'Sum'. The 'Distribution' section includes 'Skewness' and 'Kurtosis'. The 'Percentile Values' section includes 'Quartiles', 'Cut points for: 10 equal groups', and 'Percentile(s):'. The 'Values are group midpoints' checkbox is also present.

در پنجره ظاهر شده گزینه (Std.deviation/Variance) را انتخاب و (Continue) را فعال نمایید سپس در پنجره (Frequencies) گزینه (OK) را کلیک کنید تا نتایج به صورت خروجی زیر مشاهده گردد:

Frequencies

Statistics

		psychology	education
N	Valid	10	10
	Missing	0	0
Std. Deviation		3.098	3.41575
Variance		9.600	11.667

تفسیر انحراف استاندارد

این شاخص همانطور که از نامش مشخص است به عنوان معیار پراکندگی تلقی می‌شود هر قدر S (انحراف معیار) کوچک باشد نشانه این است که گروه مورد مطالعه از لحاظ ویژگی مورد سنجش متجانس بوده است و بالعکس مقدار بیشتر S بیانگر نامتجانس بودن گروه است.

در گروه (روانشناسی) انحراف استاندارد و واریانس نمره‌های روش تحقیق کوچک‌تر از انحراف استاندارد و واریانس همین آزمون در گروه (علوم تربیتی) می‌باشد. بنابراین می‌توان نتیجه گرفت که گروه (روانشناسی) توانش یادگیری کم و بیش یکسانی نسبت به گروه (علوم تربیتی) دارند. یعنی هرچه مقدار انحراف استاندارد کوچک‌تر باشد. گروه مورد مطالعه از نظر ویژگی مورد سنجش متجانس‌تر و همگون‌تر است. (پاشا شریفی، نجفی زند، ۱۳۷۵: ۱۰۳)

رتبه درصدی و دهک^۱

الف) محاسبه رتبه درصدی

رتبه درصدی یعنی نمره فرد از چند درصد نمرات بالاتر است، به عبارت دیگر چند درصد از نمرات زیر نمرات فرد قرار دارد رتبه ۵۰ درصدی یعنی ۵۰ درصد نمرات زیر نمرات این فرد قرار گرفته است، به عبارت دیگر رتبه درصدی وضعیت فرد در گروه را بر حسب کسانی که نمره پایین‌تر از او گرفته‌اند مشخص می‌کند هر نمره خام دارای یک رتبه درصدی است. مثلاً نمره خام ۱۵ یک دانش‌آموز دارای رتبه درصدی ۶۵ خواهد بود به شرط اینکه ۶۵ درصد دانش‌آموزان گروه نمره خام کمتر از ۱۵ گرفته باشند. به بیان دیگر نمره این دانش‌آموز (۱۵) از نمره ۶۵ درصد گروه بیشتر است.

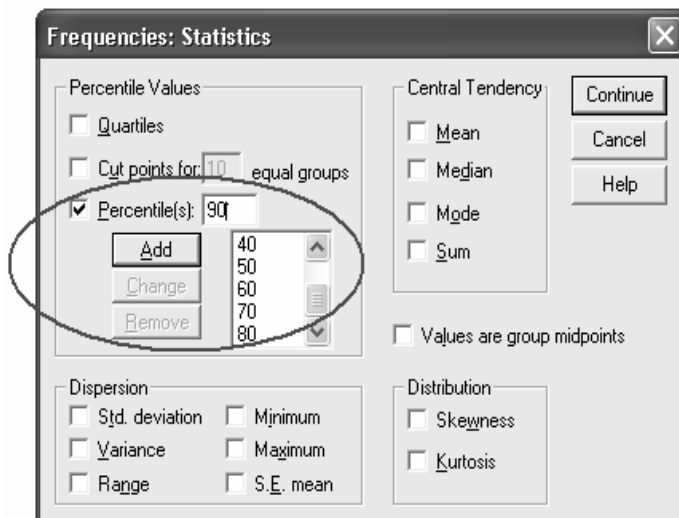
مثال: فرض کنید نمرات دو گروه از دانشجویان علوم تربیتی و روانشناسی در درس آمار در جدول زیر آمده است با استفاده از نمرات، رتبه درصدی را محاسبه کنید؟

روانشناسی	علوم تربیتی
۱۰	۱۹
۱۲	۱۶/۲۵
۱۵	۱۴
۱۷	۲۰
۱۸	۱۲/۵۰
۱۹	۱۷
۱۶	۱۸
۱۴	۱۹
۲۰	۹

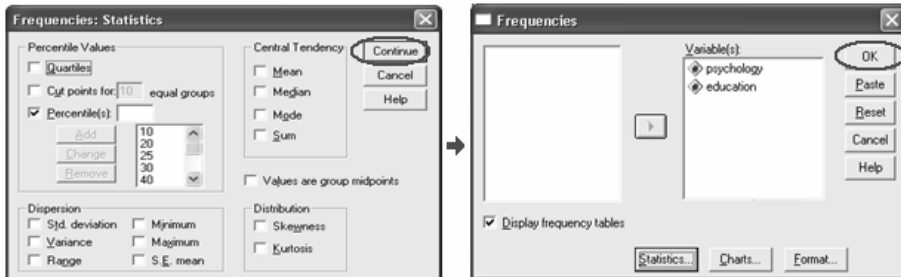
برای محاسبه رتبه درصدی در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies...
Statistics...
Percentile(S):

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده گزینه (Percentile(S):) را انتخاب کنید و رتبه‌هایی که قصد محاسبه آن را دارید نوشته و (Add) کنید سپس گزینه (Continue) را کلیک کنید، در پنجره (Frequencies) گزینه (OK) را کلیک نمایید:



تا نتایج به صورت خروجی زیر مشاهده گردد:

Statistics		psychology	education	
N	Valid	10	10	
	Missing	0	0	
Percentiles	10	10.20	9.3500	
	20	12.40	12.8000	
	25	13.50	13.6250	رتبه ۲۵ درصدی / چارک
	30	14.30	14.6750	
	40	15.00	16.5500	
	50	15.50	17.0000	رتبه ۵۰ درصدی / چارک
	60	16.60	17.6000	
	70	17.70	18.7000	
	75	18.25	19.0000	رتبه ۷۵ درصدی / چارک
	80	18.80	19.0000	
	90	19.90	19.9000	

رتبه درصدی ۵۰، همان میانه است که عملکرد متوسط دانش‌آموزان را نشان می‌دهد. رتبه درصدی کوچک‌تر از ۵۰ نشان دهنده عملکرد پایین‌تر از متوسط و رتبه درصدی‌های بزرگ‌تر از ۵۰ معرف عملکرد بالاتر از متوسط هستند. رتبه درصدی ۲۵ که به نقطه چارکی اول و رتبه درصدی ۷۵ که به نقطه چارکی سوم معروفند پایین‌ترین و بالاترین ربع توزیع را مشخص می‌کنند. این دو رتبه درصدی نیز مرزهای مشخصی برای توزیع نمره‌ها به حساب می‌آیند.

ب) محاسبه دهک

برای اینکه اندازه پراکندگی تنها از دو مقدار انتهایی ناشی نشود، می‌توانیم دامنه تغییر میان دهک‌ها را محاسبه کنیم. مثلاً دامنه تغییر میان ۱۰ و ۹۰ درصدی را حساب کنید. نقطه ۱۰ درصدی را دهک اول و نقطه ۹۰ درصدی را دهک نهم می‌نامند. (پاشا شریفی، نجفی زند، ۱۳۷۵: ۹۰) در محاسبه دهک همان مراحل‌هایی که در محاسبه رتبه درصدی انجام دادیم انجام دهید تا نتایج به صورت زیر آشکار شود:

Frequencies

Statistics		psychology	education	
N	Valid	10	10	
	Missing	0	0	
Percentiles	10	10.20	9.3500	دهک اول
	20	12.40	12.8000	
	30	14.30	14.6750	
	40	15.00	16.5500	
	50	15.50	17.0000	
	60	16.60	17.6000	
	70	17.70	18.7000	
	80	18.80	19.0000	
	90	19.90	19.9000	دهک نهم

ضریب تغییرپذیری (ضریب نسبی واریانس)

یکی دیگر از مقادیر پراکندگی، ضریب تغییر پذیری^۱ (CV) است که می‌توان آن را به

صورت نسبت یا درصد محاسبه کرد:

محاسبه ضریب نسبی واریانس به صورت نسبت:

$$CV = \frac{S}{\bar{X}}$$

محاسبه ضریب نسبی واریانس به صورت درصد:

$$CV = \frac{S}{\bar{X}} \times 100$$

مثال: فرض کنید دستمزد ۷ نفر از کارکنان دو سازمان که به طور تصادفی استفاده

شده‌اند به شرح زیر است:

سازمان آموزش و پرورش	سازمان تأمین اجتماعی
۱۸۰,۵۰۰	۲۶۰,۰۰۰
۱۹۰,۰۰۰	۲۸۰,۰۰۰
۲۴۰,۰۰۰	۲۴۰,۰۰۰
۱۷۰,۵۰۰	۲۲۰,۰۰۰
۴۰۰,۵۰۰	۳۳۰,۰۰۰
۲۹۰,۰۰۰	۲۰۰,۰۰۰
۲۶۰,۵۰۰	۲۲۰,۰۰۰

Statistics

	سازمان.الف	سازمان.ب
N Valid	7	7
Missing	0	0
Mean	250000.0	247428.6
Std. Deviation	44347.12	80792.65

محاسبه ضریب نسبی واریانس سازمان (الف) به صورت نسبت:

$$CV = \frac{44347.12}{250000} = 0.17738$$

محاسبه ضریب نسبی واریانس سازمان (ب) به صورت نسبت:

$$CV = \frac{80792.65}{247428.6} = 0.32652$$

ملاحظه می‌شود که نسبت پراکندگی حقوق سازمان (ب) بیشتر از نسبت پراکندگی حقوق سازمان (الف) است. هر قدر که مقدار ضریب نسبی واریانس به صفر نزدیک‌تر باشد، بیانگر همگونی و هر چه مقدار آن به یک نزدیک‌تر باشد، نشانگر ناهمگونی در داده‌ها است.

شاخص پراکندگی که در مورد داده‌های اسمی و رتبه‌ای، مورد استفاده قرار می‌گیرد معمولاً نسبت همگنی یا ناهمگنی طبقات ثبت شده و یا رتبه‌ها است که با نسبت تغییرپذیری^۱ (VR) مشخص می‌شود.

$$VR = 1 - \frac{f_m}{n}$$

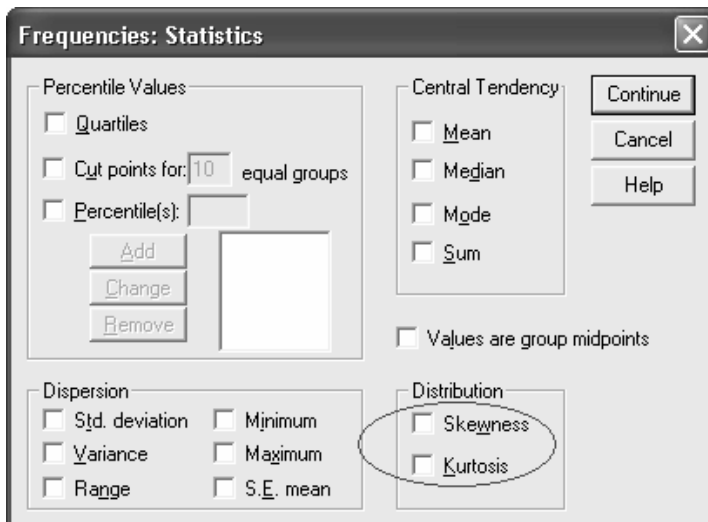
در این فرمول f_m فراوانی طبقه نمایی یعنی طبقه‌ای که بیشترین فراوانی را دارد. هر قدر مقدار VR به صفر نزدیک‌تر باشد بیانگر همگونی و بر عکس هر چه مقدار آن نزدیک به یک باشد، نشانگر ناهمگونی در داده‌هاست.

کجی و کشیدگی

میانگین و انحراف استاندارد، به دسته‌ای از شاخص‌های آمار توصیفی که گشتاورها نامیده می‌شوند گشتاور رتبه اول (میانگین انحراف هر یک از نمره‌ها از میانگین) صفر است و گشتاور رتبه دوم برابر واریانس نمونه است. گشتاور رتبه سوم به منظور محاسبه اندازه کجی و گشتاور رتبه چهارم برای محاسبه کشیدگی منحنی بکار برده می‌شود. (دلاور، ۱۳۸۰: ۲۰۲)

برای محاسبه کجی و کشیدگی در Spss مراحل زیر را انجام دهید:

Analyze
Descriptive Statistics
Frequencies...
Statistics...
Skewness / Kurtosis



الف) کجی^۱

در صورتی که منحنی متقارن باشد کجی برابر صفر است بنابراین کجی یعنی، انحراف یک منحنی از حالت تقارن. (دلاور، ۱۳۸۰: ۲۰۳)

۴	۲	۵	۳	۱
۱	۳	۴	۲	۴
۳	۱	۵	۲	۳
۱	۴	۵	۴	۵
۲	۳	۵	۱	۲

برون‌داد محاسبه شده کجی مربوط به نمرات فوق

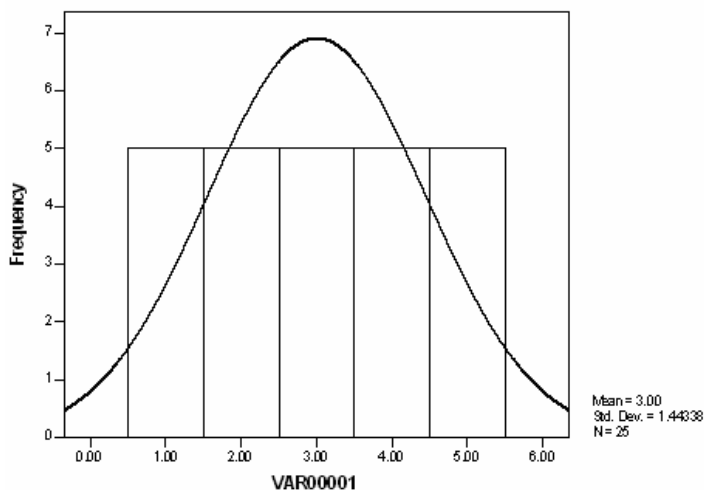
Frequencies

Statistics

VAR00001

N	Valid	25
	Missing	0
Skewness		.000
Std. Error of Skewness		.464

Histogram



اما در صورتی که توزیع نمره دارای کجی منفی باشد، مکعب مجموع مجذور انحراف نمره‌ها از میانگین عدد منفی خواهد شد (دلاور، ۱۳۸۰: ۲۰۴)

Frequencies

Statistics

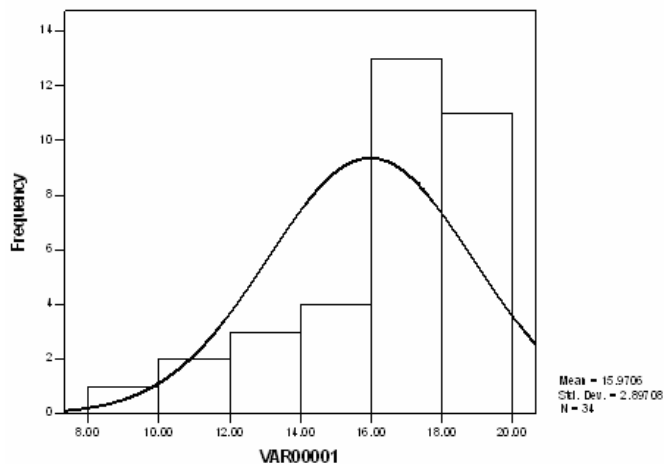
VAR00001

N	Valid	34
	Missing	0
Skewness		-.724
Std. Error of Skewness		.403

VAR00001

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 9.00	1	2.9	2.9	2.9
10.00	1	2.9	2.9	5.9
11.00	1	2.9	2.9	8.8
12.00	3	8.8	8.8	17.6
14.00	4	11.8	11.8	29.4
16.00	7	20.6	20.6	50.0
17.00	6	17.6	17.6	67.6
18.00	6	17.6	17.6	85.3
19.00	1	2.9	2.9	88.2
20.00	4	11.8	11.8	100.0
Total	34	100.0	100.0	

Histogram



به همین ترتیب هرگاه توزیع دارای کجی مثبت باشد، مکعب مجموع مجذور انحراف نمره‌ها از میانگین مثبت خواهد شد (دلاور، ۱۳۸۰: ۲۰۴)

Frequencies

Statistics

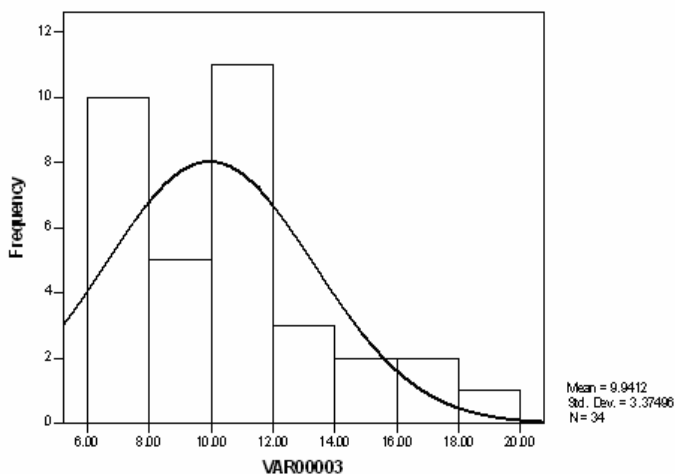
VAR00003

N	Valid	34
	Missing	0
Skewness		1.044
Std. Error of Skewness		.403

VAR00003

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 6.00	7	20.6	20.6	20.6
7.00	3	8.8	8.8	29.4
9.00	5	14.7	14.7	44.1
10.00	8	23.5	23.5	67.6
11.00	3	8.8	8.8	76.5
12.00	3	8.8	8.8	85.3
14.00	2	5.9	5.9	91.2
16.00	1	2.9	2.9	94.1
17.00	1	2.9	2.9	97.1
20.00	1	2.9	2.9	100.0
Total	34	100.0	100.0	

Histogram



ب) کشیدگی^۱

هنگامی که مقدار کشیدگی برابر صفر باشد توزیع نمره‌ها طبیعی است ولی در صورتی که کشیدگی منفی باشد منحنی توزیع نمره‌ها در نقطه اوج خوابیده است. (دلاور، ۱۳۸۰: ۲۰۵)

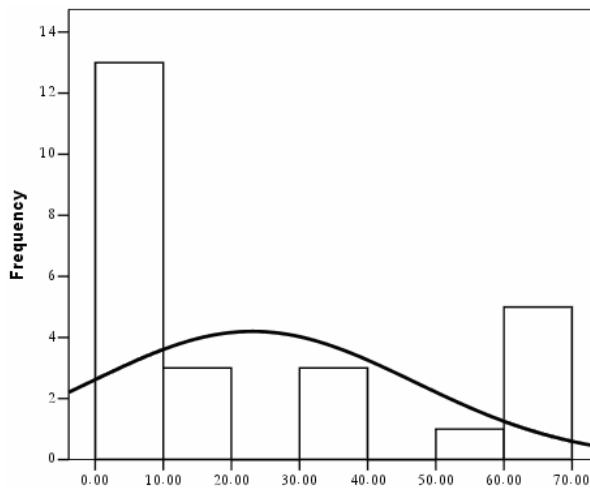
Frequencies

Statistics

N	Valid	25
	Missing	0
Kurtosis		-.537
Std. Error of Kurtosis		.902

VAR00002

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 1.00	3	12.0	12.0	12.0
6.00	4	16.0	16.0	28.0
9.00	6	24.0	24.0	52.0
12.00	3	12.0	12.0	64.0
30.00	3	12.0	12.0	76.0
50.00	1	4.0	4.0	80.0
60.00	3	12.0	12.0	92.0
70.00	2	8.0	8.0	100.0
Total	25	100.0	100.0	



اما بالاخره زمانی که کشیدگی مثبت باشد، برآمدگی منحنی توزیع نمره‌ها در نقطه اوج قرار خواهد گرفت. (دلاور، ۱۳۸۰: ۲۰۵)

Frequencies

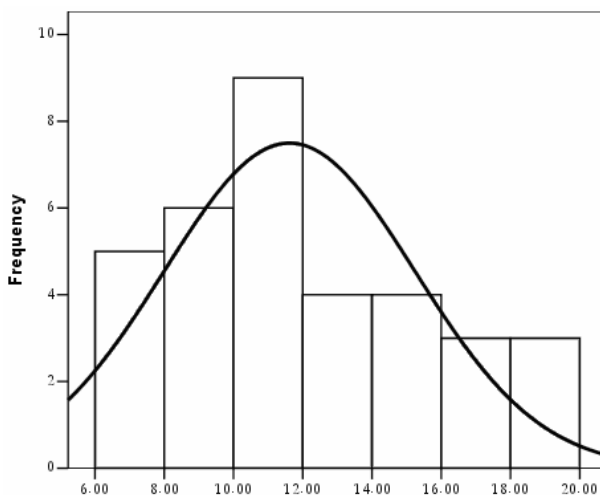
Statistics

VAR00002

N	Valid	34
	Missing	0
Kurtosis		.062
Std. Error of Kurtosis		.788

VAR00002

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 7.00	5	14.7	14.7	14.7
9.00	6	17.6	17.6	32.4
10.00	4	11.8	11.8	44.1
11.00	5	14.7	14.7	58.8
12.00	4	11.8	11.8	70.6
14.00	4	11.8	11.8	82.4
16.00	2	5.9	5.9	88.2
17.00	1	2.9	2.9	91.2
18.00	1	2.9	2.9	94.1
20.00	2	5.9	5.9	100.0
Total	34	100.0	100.0	



خود آزمایی

۱. شاخص‌های پراکندگی که دارای بیشترین استفاده هستند، را نام ببرید.
۲. دامنه تغییرات را تعریف و دامنه تغییرات اعداد زیر را تعیین کنید.
۳. انحراف چارکی را در مجموعه‌ای از اعداد محاسبه و تعیین کنید.
۴. واریانس و انحراف استاندارد نمره‌های زیر را محاسبه کنید:

۱۵	۱۷	۱۴	۱۱	۱۱	۶	۱۹	۱۵
۱۳	۱۴	۱۶	۱۷	۱۷	۱۵	۲۰	۱۷
۱۲	۱۶	۳	۲۰	۹	۲	۱۵	۱۴
۱۷	۲	۷	۱۵	۱۶	۱۹	۱۴	۱۲
۱۴	۹	۱۸	۱۶	۱۸	۱۸	۱۷	۱۳
۶	۸	۱۹	۱۸	۲۰	۵	۱۶	۱۰
۲	۱۷	۲۰	۱۰	۶	۱۲	۱۲	۱۱
۱۷	۱۹	۳	۱۱	۴	۱۰	۱۳	۹

۵. اعداد زیر، نمرات امتحانی ۲۴ دانشجوی روان‌شناسی را در درس آمار توصیفی نشان می‌دهد:

۱۴	۱۱/۲۵	۱۴	۱۲	۷	۱۲
۱۳	۱۶	۵	۱۱/۵	۱۱	۱۴
۲۰	۱۷	۱۵/۷۵	۱۰	۲۰	۱۵
۱۹	۱۰	۶	۱۹	۱۹	۱۳
۱۴	۶/۷۵	۱۷	۱۶/۲۵	۱۳	۱۰

- الف) انحراف دهک‌ها را محاسبه کنید.
- ب) دامنه چارک‌ها را محاسبه کنید.
- ج) انحراف چارک‌ها را محاسبه کنید.
۶. آزمون واحدی با دو گروه (الف) و (ب) اجرا شده و نتایج آن به صورت زیر است:

گروه (الف)	گروه (ب)
۱۵	۱۴
۱۶	۱۷
۱۶/۵	۱۷
۱۰	۱۱
۱۰/۵	۱۱
۱۷	۱۳
۱۷/۲۵	۱۲
۱۱	۱۵
۱۲	۱۴
۱۳	۱۴
۱۳/۲۵	۱۵
۱۴	۱۴
۱۶	۱۷

۸. ضریب تغییرپذیری (ضریب نسبی واریانس) را محاسبه کنید و آن را تفسیر نمایید.
۸. کجی و کشیدگی را تعریف کنید

فصل ششم

پیش فرض‌های آزمون پارامتریک و ناپارامتریک

هدف‌های یادگیری

- از دانشجو انتظار می‌رود که پس از خواندن فصل ششم بتواند:
۱. مهم‌ترین پیش فرض‌های آزمون‌های پارامتریک را نام ببرد.
 ۲. دربارهٔ نرمال بودن توزیع جامعه تصمیم‌گیری کند.
 ۳. کاربرد آزمون علامت‌ها را بیان کند.

پیش فرض‌های آزمون پارامتریک و ناپارامتریک

به آزمون‌های آماری کلاسیک نظیر آزمون t و تحلیل واریانس که فرضیه‌های مربوط به پارامترها را در جامعه مورد آزمایش قرار می‌دهند، آزمون‌های پارامتریک گفته می‌شود. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۲۱۵) معنی‌دار بودن نتایج یک آزمون آماری پارامتری بستگی به اعتبار مفروضات و شرایطی دارد که در زیر به آن می‌پردازیم:

۱. مشاهده‌های ما باید مستقل از یکدیگر باشد. یعنی انتخاب هر موردی از میان جامعه آماری برای قرار دادن آن مورد در نمونه ما نباید اختلالی در شانس انتخاب شدن هر یک از موارد دیگر در نمونه ما ایجاد کند.

۲. مشاهده‌های ما باید از جامعه‌های آماری که دارای توزیع نرمال هستند بیرون کشیده شده باشند.

۳. جامعه‌های آماری مورد مطالعه باید دارای واریانس برابر باشند.

۴. متغیرهای مربوط باید حداقل با مقیاس فاصله‌ای اندازه‌گیری شده باشند. (سیگل

ترجمه کریمی، ۱۳۸۰: ۲۴)

وقتی دلایلی موجود باشد که نشان بدهد اطلاعات جمع‌آوری شده ما واجد شرایط مذکور باشد در آن صورت یقیناً باید از آزمون‌های پارامتری برای تجزیه و تحلیل داده‌ها استفاده شود.

مهم‌ترین پیش‌فرض‌های آزمون‌های پارامتریک عبارتند از:

۱. نرمال بودن توزیع جامعه

۲. استقلال داده‌ها (تصادفی بودن داده‌ها) (حسینی، ۱۳۸۲: ۳۴)

نرمال بودن توزیع جامعه

برای اثبات نرمال بودن توزیع یک متغیر معمولاً از آزمون کلموگروف-اسمیرنوف که به آزمون‌های نیکویی برازش معروف است، استفاده می‌شود و جز آزمون‌های ناپارامتریک محسوب می‌شود، به عبارت دیگر برای آزمون نرمال بودن توزیع، می‌بایست از آزمون ناپارامتریک استفاده شود (حسینی، ۱۳۸۲: ۳۴) برای اجرای این آزمون، مراحل زیر را اجرا کنید:

Analyze

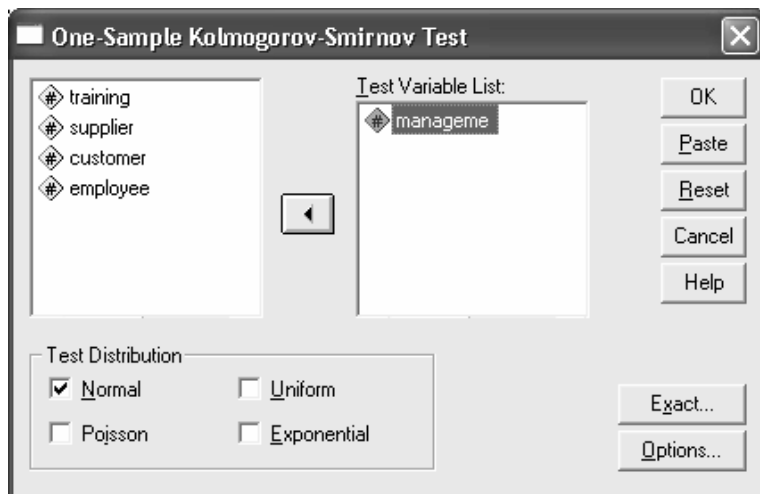
Nonparametric test

1- Sample k-s.....

The screenshot shows the SPSS 13.0 software interface. The 'Analyze' menu is open, and the path 'Nonparametric Tests' > '1-Sample K-S...' is highlighted. The background shows a data view with columns 'training', 'supplier', and 'custo' and rows of numerical data.

	training	supplier	custo
1	18.00	22.00	33
2	15.75	25.00	34
3	15.00	15.00	29
4	16.50	18.00	26
5	16.00	17.00	27
6	15.93	17.00	30
7	17.00	24.00	25
8	17.61	20.00	26
9	17.00	16.00	28
10	16.58	19.00	28
11	14.00	14.00	27
12	16.60	22.00	28
13	15.28	21.00	23
14	17.00	13.00	29.00
15	18.75	23.00	27.00

بعد از اجرای این فرمان، پنجره گفتگوی زیر ظاهر می‌شود.



در پنجره ظاهر شده، متغیر مورد نظر را به قسمت (test variable list) منتقل و Spss به‌عنوان پیش فرض، قسمت (test distribution) گزینه نرمال را انتخاب کرده و سپس (ok) کنید. تا نتایج به‌صورت زیر آشکار شود:

NPar Tests

One-Sample Kolmogorov-Smirnov Test

		manageme
N		325
Normal Parameters ^{a, b}	Mean	89.6154
	Std. Deviation	10.35585
Most Extreme Differences	Absolute	.061
	Positive	.049
	Negative	-.061
Kolmogorov-Smirnov Z		1.097
Asymp. Sig. (2-tailed)		.180

a. Test distribution is Normal.

b. Calculated from data.

برای تصمیم‌گیری می‌توان به دو روش عمل کرد:
روش اول: در خروجی Spss به مقدار (Kolmogorov-Smirnov z) توجه نمایید.
چنانچه مقدار آن بین $-1/96$ و $1/96$ باشد با ۹۵ درصد اطمینان می‌توان به نرمال بودن

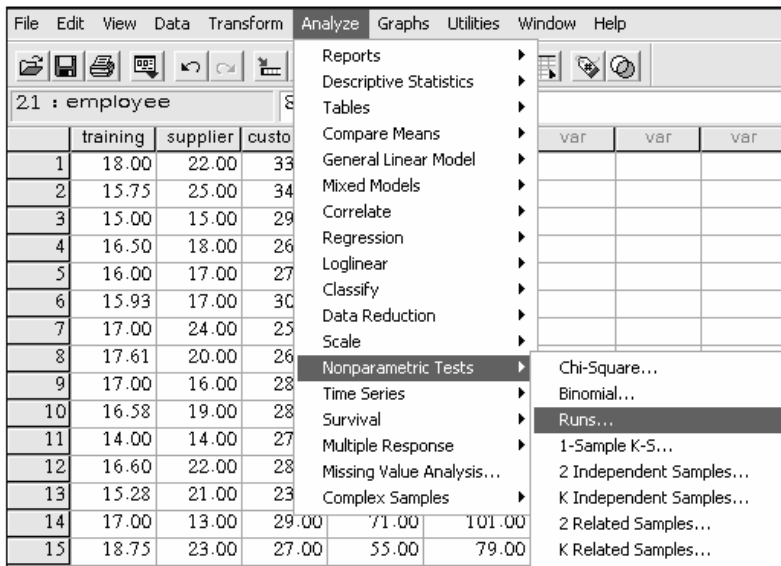
توزیع حکم کرد و چنانچه مقدار بزرگ‌تر از $1/96+$ و یا کوچک‌تر از $1/96-$ باشد توزیع جامعه نرمال نمی‌باشد.

روش دوم: در جدول خروجی Spss به قسمت Asymp.sig(2-tailed) توجه کنید و آن را بر دو تقسیم نمایید، اگر مقدار آن بیشتر از $2/5$ درصد بوده، توزیع نرمال است و اگر کمتر از $2/5$ باشد نرمال نبودن توزیع را نتیجه‌گیری کنید (حسینی، ۱۳۸۲: ۷۰)

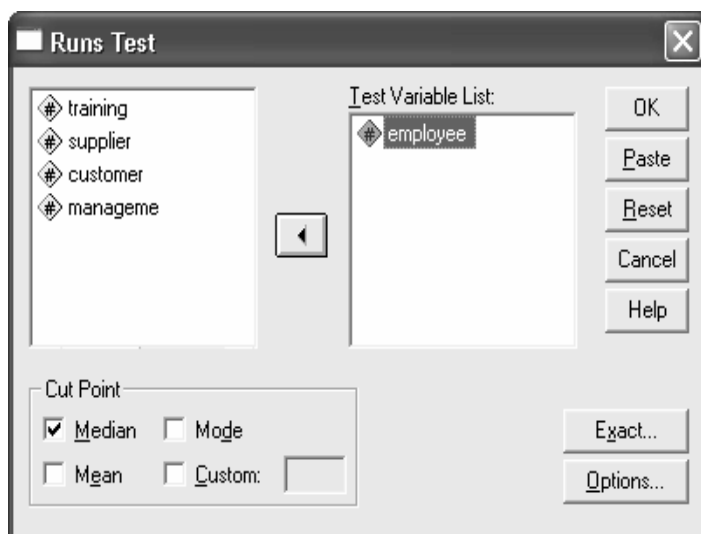
استقلال داده‌ها (تصادفی بودن داده‌ها)

یکی دیگر از پیش فرض‌های آزمون پارامتریک، استقلال داده‌ها (تصادفی بودن داده‌ها) است. برای مشخص کردن این پیش فرض، از آزمون علامت استفاده می‌شود. آزمون علامت‌ها، از این حقیقت ناشی می‌شود که در این آزمون، به جای مقادیر عددی از (+) و (-) به‌عنوان یافته‌های آن استفاده می‌شود و برای همین منظور با مبنا قرار دادن میانه برای تعیین علامت مثبت و منفی استفاده می‌شود. (سیگل ترجمه کریمی، ۱۳۸۰: ۸۸) البته شایان ذکر است که این آزمون قابلیت آن را خواهد داشت که به جزء این مبنا، با هر مبنا دلخواه دیگر مانند نما یا میانگین نیز اجرا می‌شود. برای اجرای این آزمون مراحل زیر را اجرا کنید:

Analyze
Nonparametric test
Runs.....



بعد از اجرای این فرمان، پنجره گفتگوی زیر ظاهر می‌شود:



در پنجره ظاهر شده، متغیر مورد نظر را به قسمت (test variable list) منتقل نمایید و سپس در قسمت (Cut point) می‌توانید هر کدام از مبنایها را انتخاب نمایید. در این قسمت (median) به معنای میانه، خود به عنوان پیش فرض انتخاب نرم‌افزار می‌باشد. نما (mode) و میانگین (mean) نیز می‌تواند انتخاب شود. بر حسب نیاز در قسمت (Custom) نیز محقق می‌تواند، مبنا را هر عدد دلخواهی قرار دهد. با انتخاب گزینه میانه، گزینه (ok) را کلیک کنید تا نتایج زیر ظاهر شود:

NPar Tests

Runs Test

	employee
Test Value ^a	76.00
Cases < Test Value	161
Cases ≥ Test Value	164
Total Cases	325
Number of Runs	167
Z	.390
Asymp. Sig. (2-tailed)	.696

a. Median

نحوه قضاوت بدین‌گونه است که سطح پوشش دوسویه را بر دو تقسیم کنید و چنانچه نتیجه کوچک‌تر از ۲/۵ درصد باشد، عدم استقلال داده‌ها (تصادفی نبودن

داده‌ها) و چنانچه مقدار از ۲/۵ درصد بیشتر باشد، استقلال داده‌ها را می‌توان نتیجه‌گیری کرد.

خودآزمایی

۱. مهم‌ترین پیش فرض‌های آزمون‌های پارامتریک را نام ببرید.
۲. کاربرد آزمون علامت و کلموگروف – اسمیرنوف یک نمونه‌ای را بیان کنید.
۳. درباره نرمال بودن اعداد زیر تصمیم‌گیری کند.

۱۶	۲۷	۱۴	۲۳	۱۶	۱۷	۱۵	۲۷	۲۵	۱۵	۲۴	۱۷
۲۳	۱۷	۱۶	۲۴	۲۷	۲۸	۲۵	۱۷	۱۵	۱۹	۲۷	۲۸
۲۴	۲۸	۲۷	۲۷	۱۷	۳۰	۲۶	۲۸	۱۹	۲۳	۱۷	۳۰
۲۷	۳۰	۱۷	۱۷	۲۸	۱۴	۱۶	۳۰	۲۳	۲۵	۲۸	۱۴
۱۷	۱۴	۲۸	۲۸	۳۰	۱۶	۲۳	۱۴	۲۸	۲۶	۳۰	۱۶
۲۸	۱۶	۳۰	۳۰	۱۴	۱۷	۲۴	۲۴	۳۰	۱۶	۱۴	۲۷

۴. درباره استقلال داده‌ها (تصادفی بودن داده‌ها) نمرات زیر اظهار نظر کنید

۲۵	۳۰	۲۰	۱۴	۱۵	۱۵	۱۱	۳۰	۳۵	۱۰	۱۲	۲۳
۳۰	۲۹	۱۰	۱۰	۱۵	۱۸	۱۷	۲۳	۱۶	۱۱	۳۵	۲۵
۳۱	۲۴	۱۱	۱۲	۱۶	۱۰	۱۵	۲۵	۲۵	۱۰	۲۶	۱۳
۳۶	۲۶	۱۳	۱۰	۱۲	۲۵	۱۶	۱۲	۲۵	۱۲	۲۵	۲۶
۳۰	۲۷	۱۲	۳۰	۲۳	۲۴	۱۹	۲۱	۲۱	۳۵	۲۴	۲۹
۲۵	۲۳	۱۳	۳۵	۲۷	۲۴	۱۴	۱۰	۲۴	۲۵	۲۸	۲۸
۲۵	۲۶	۱۵	۳۰	۳۱	۲۷	۲۴	۲۰	۲۸	۲۴	۱۷	۲۷
۱۴	۱۳	۲۰	۳۴	۱۳	۲۰	۲۷	۳۰	۲۹	۱۸	۱۴	۱۳

فصل هفتم

آزمون‌های آماری پارامتریک

هدف‌های یادگیری

از دانشجو انتظار می‌رود که پس از خواندن فصل هفتم بتواند:

۱. آزمون پارامتریک را تعریف کند.
۲. آزمون t یک نمونه‌ای را محاسبه کند.
۳. پیش‌نیازها و فرض‌های آزمون t گروه‌های وابسته را بیان کند.
۴. آزمون t برای گروه‌های وابسته را محاسبه کند.
۵. پیش‌نیازها و فرض‌های آزمون t دو نمونه مستقل را بیان کند.
۶. آزمون t دو نمونه مستقل را محاسبه کند.
۷. پیش‌نیازها و فرض‌های آزمون تحلیل واریانس را بیان کند.
۸. آزمون تحلیل واریانس را محاسبه کند.
۹. کاربرد، مقایسه‌های بعد از تجربه (شفه و توکی) را بیان کند.
۱۰. آزمون شفه را محاسبه کند.
۱۱. آزمون توکی را محاسبه کند.

آزمون‌های آماری پارامتریک

آزمون پارامتریک، آزمونی است که مدل آماری آن برخی شرایط معین را درباره پارامتر جامعه‌ای که نمونه ما از آن گرفته شده است، وضع می‌کند. (سیگل ترجمه کریمی، ۱۳۸۰: ۴۱) معروف‌ترین آزمون‌های آماری پارامتریک عبارتند از:

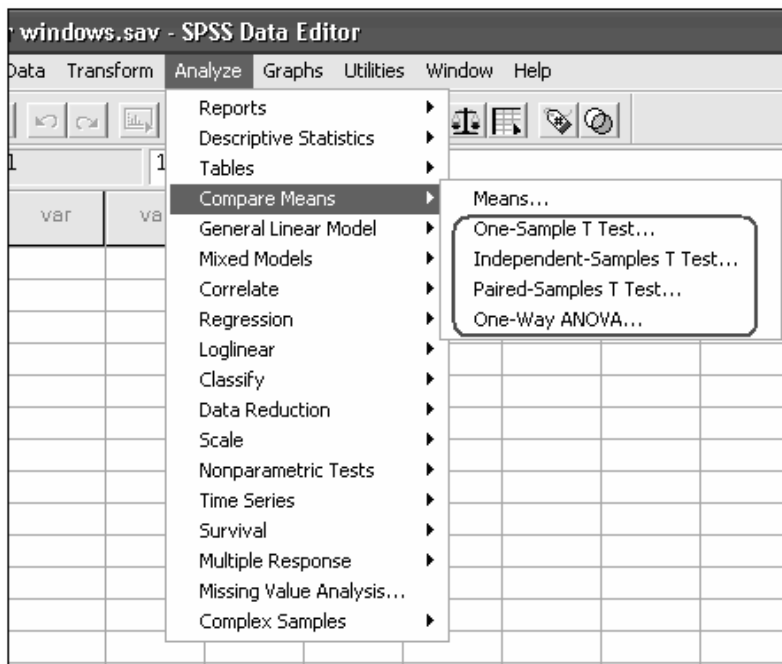
۱. آزمون t یک نمونه‌ای

۲. آزمون t برای گروه‌های وابسته

۳. آزمون t دو نمونه مستقل

۴. آزمون تحلیل واریانس

زیر منوی (Compare Means) برای اجرای آزمون‌های پارامتری در منوی (Statistics) گنجانده شده است، شکل زیر فرمان‌های مختلف این زیر منو را نشان می‌دهد:



آزمون t یک نمونه

هدف از اجرای آزمون t آن است که به سؤال (آیا میانگین نمونه از جامعه نظری با میانگین μ انتخاب شده یا از جامعه دیگر) پاسخ داده شود.

فرض‌ها

۱. نمره‌ها به صورت نمونه‌گیری تصادفی از جامعه مورد نظر انتخاب شده باشند.

۲. شکل توزیع پراکندگی نمرات در جامعه نرمال باشد (شیولسون ترجمه کیامنش، ۱۳۸۲: ۱۸۱).

$$\tau = \frac{\bar{X} - \mu}{S_{\bar{X}}}$$

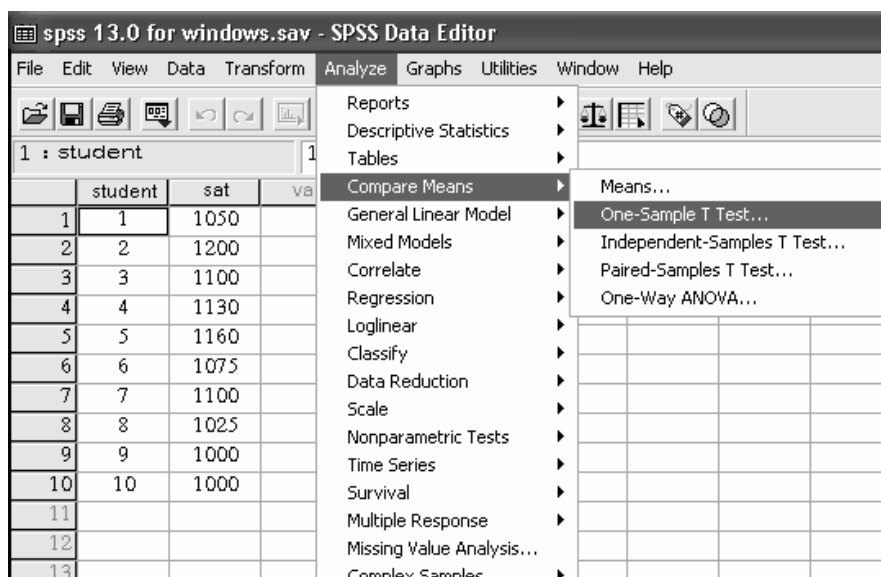
فرض کنید یک نمونه تصادفی با حجم ۱۰ نفر از میان جامعه دانشجویان سال اول دانشگاه انتخاب شده که اطلاعات نظری مربوط به نمرات آزمون استعداد تحصیلی در جدول زیر ارائه گردیده است.

SAT	دانش‌آموز
۱۰۵۰	۱
۱۲۰۰	۲
۱۱۰۰	۳
۱۱۳۰	۴
۱۱۶۰	۵
۱۰۷۵	۶
۱۱۰۰	۷
۱۰۲۵	۸
۱۰۰۰	۹
۱۰۰۰	۱۰

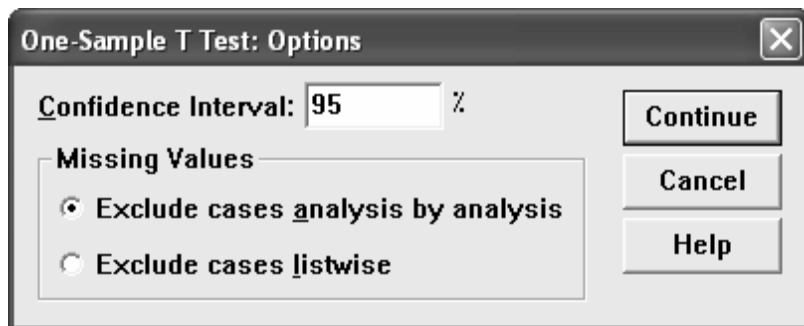
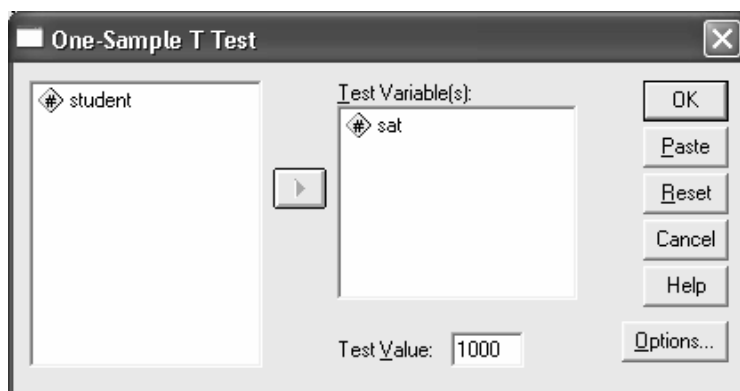
با توجه به این مسأله که میانگین نمرات برای جامعه دانش‌آموزان داوطلب ورود به دانشگاه ۱۰۰۰ است سؤال مورد پرسش این است که آیا نمره دانشجویان سال اول دانشگاه به‌طور متوسط از جامعه دانش‌آموزان داوطلب ورود به دانشگاه بیشتر است؟ (شیولسون ترجمه کیامنش، ۱۳۸۲: ۱۸۶).

مراحل اجرای آزمون t یک نمونه‌ای در Spss به شرح زیر می‌باشد:

Analyze
Compare Means
One-Sample t test...



بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده متغیرها را به (test Variable(s): منتقل کنید و در قسمت (Test Value:) عدد ۱۰۰ (ملاک جهانی) تایپ نمایید و در قسمت (Options...) می‌توانید سطح معنی‌داری را تعیین کنید.

پس از تعیین سطح معنی‌داری، با فعال کردن گزینه (Continue) دوباره به پنجره (One-Sample t test) باز می‌گردید و (OK) را کلیک نمایید تا نتایج به صورت خروجی زیر مشاهده گردد:

T-Test

One-Sample Test

	Test Value = 1000				
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference
					Lower Upper
sat	3.951	9	.003	84.000	35.90 132.10

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
sat	10	1084.00	67.239	21.263

در جدول نتایج مقدار آزمون معادل ۳/۹۵۱ و درجه آزادی ۹ را نشان می‌دهد که این مقدار در سطح ۰/۰۵ معنی‌دار است بنابراین فرضیه صفر باید رد گردد. یعنی میانگین نمره‌های استعداد تحصیلی دانشجویان دانشگاه مورد بحث که گروه نمونه به صورت تصادفی از میان آنان انتخاب شده بود از میانگین نمره‌های دانش‌آموزان فارغ‌التحصیل دبیرستانی شرکت کننده در آزمون پیشرفت تحصیلی بیشتر است.

آزمون t برای گروه‌های وابسته

هدف آزمون t برای گروه‌های وابسته آن است که مشخص کند تفاوت بین دو میانگین نمونه حاصل، عامل شانس است یا تفاوت واقعی بین میانگین‌های جامعه. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۱۰).

پیش‌نیازهای طرح

۱. وجود یک متغیر مستقل با دو سطح
۲. مشاهده یک گروه آزمودنی در هر دو حالت عمل آزمایشی
۳. وجود تفاوت کمی و کیفی بین سطوح‌های متغیر مستقل

فرض‌ها

۱. نمره آزمودنی‌ها از یکدیگر مستقل بوده و از جامعه مورد نظر به صورت تصادفی انتخاب شده‌اند.
۲. شکل توزیع نمره‌ها در جامعه مورد نظر نرمال باشد.
۳. واریانس نمره‌ها در جامعه‌های مورد نظر با یکدیگر مساوی‌اند (همگنی واریانس) (شیولسون ترجمه کیامنش، ۱۳۸۲: ۱۳).

$$t = \frac{\sum D}{\sqrt{\frac{n\sum D^2 - (\sum D)^2}{(n-1)}}}$$

مثال: معلم یک کلاس، یک آزمون استاندارد شده را پیش از شروع تدریس در کلاس اجرا می‌کند. در پایان دوره آموزش نیز همان آزمون را درباره همان دانش‌آموزان اجرا می‌کند. هدف وی از تحقیق این است که میزان افزایش آموخته‌های شاگردان را در دوره آموزش تعیین کند که نتایج دو آزمون پیش آزمون و پس آزمون در جدول زیر آمده است (پاشا شریفی و نجفی زند، ۱۳۷۵: ۲۱۶)

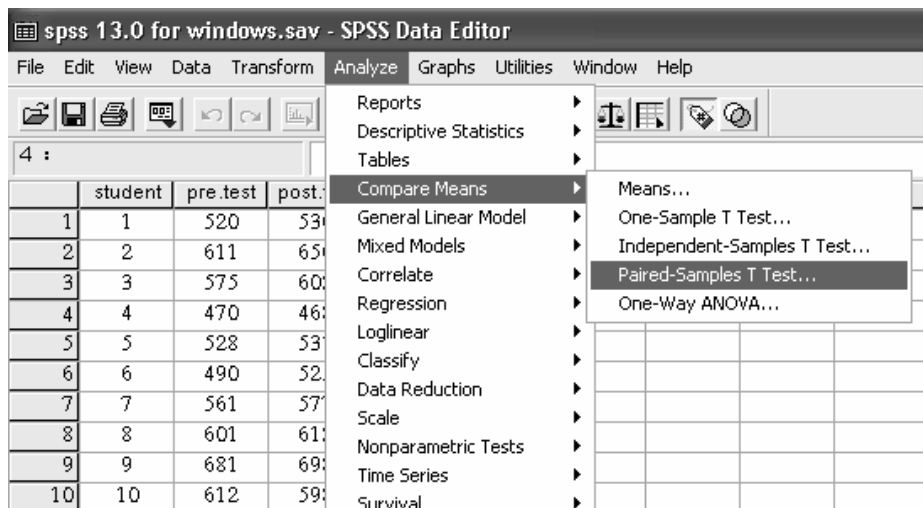
پس آزمون	پیش آزمون	دانش آموز
۵۳۰	۵۲۰	۱
۶۵۰	۶۱۱	۲
۶۰۲	۵۷۵	۳
۴۶۳	۴۷۰	۴
۵۳۱	۵۲۸	۵
۵۲۵	۴۹۰	۶
۵۷۷	۵۶۱	۷
۶۱۲	۶۰۱	۸
۶۹۸	۶۸۱	۹
۵۹۸	۶۱۲	۱۰

مراحل اجرای آزمون t وابسته در Spss به شرح زیر در نرم‌افزار اجرا می‌شود:

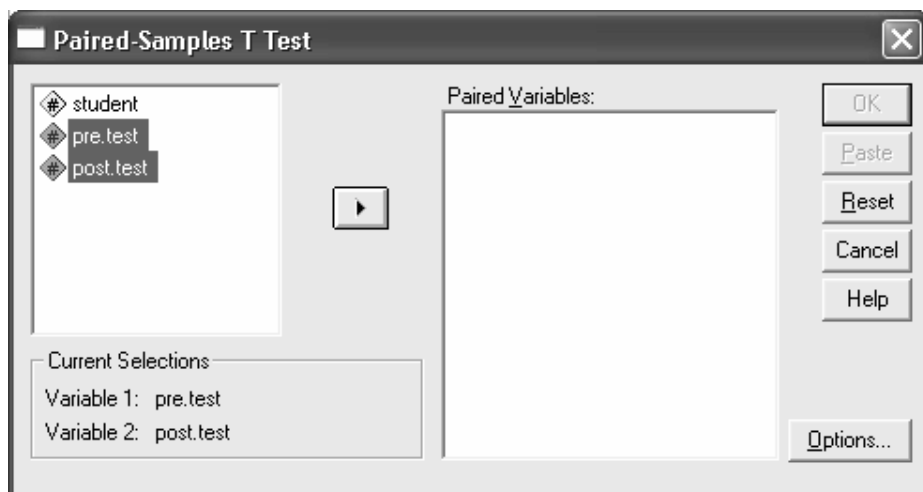
Analyze

Compare Means

Paired-Samples t test



بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده متغیرها را به قسمت (Paired Variables) وارد کنید در

قسمت (Options...) می‌توانید سطح معنی‌داری را تعیین کنید.

Paired-Samples T Test: Options

Confidence Interval: 95 %

Missing Values

☒ Exclude cases analysis by analysis

☐ Exclude cases listwise

Continue

Cancel

Help

پس از تعیین سطح معنی داری با فعال کردن گزینه (Continue)، دوباره به پنجره (Paired-Sample t test) باز می گردید و سپس (OK) را کلیک نموده تا نتایج به صورت خروجی زیر مشاهده گردد:

T-Test

Paired Samples Statistics

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	post.test	578.6000	10	68.52445	21.66933
	pre.test	564.9000	10	64.32288	20.34068

Paired Samples Correlations

	N	Correlation	Sig.
Pair 1 post.test & pre.test	10	.969	.000

Paired Samples Test

		Paired Differences					t	df	Sig. (2-tailed)
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower	Upper			
Pair 1	post.test – pre.test	13.70000	17.05579	5.39351	1.49902	25.90098	2.540	9	.032

جدول نتایج مقدار آزمون را معادل ۲/۵۴۰ و درجه آزادی ۹ نشان می دهد که این مقدار در سطح ۰/۰۵ معنی دار است یعنی بین دو گروه مورد مقایسه تفاوت وجود دارد.

آزمون t دو نمونه مستقل

هدف از آزمون t در بررسی دو میانگین مستقل آن است که محقق را در زمینه تصمیم گیری یاری نماید، محقق باید تصمیم بگیرد که آیا تفاوت مشاهده شده، بین دو

میانگین نمونه در اثر عوامل شانس به وجود آمده یا تفاوت مشاهده شده بیانگر تفاوت واقعی بین دو جامعه است.

پیش‌نیازهای طرح

۱. وجود یک متغیر مستقل با دو سطح
۲. قرار گرفتن هر فرد آزمودنی فقط در یکی از سطح‌های متغیر مستقل
۳. وجود تفاوت کمی و کیفی بین سطوح متغیر مستقل

فرض‌ها

۱. نمره آزمودنی‌های هر یک از گروه‌ها به صورت نمونه‌گیری تصادفی از جامعه مورد نظر انتخاب شده و انتخاب هر آزمونی مستقل از آزمونی دیگر صورت گرفته باشد.
۲. شکل توزیع نمره‌ها در جامعه‌های مورد نظر نرمال باشد.
۳. واریانس نمره‌های دو جامعه با هم مساوی باشد. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۱۹۳).

$$t = \frac{\overline{X}_1 - \overline{X}_2}{\sqrt{\frac{\sum X_1^2 + \sum X_2^2}{n + n - 2} \left(\frac{1}{n_1} \right) \left(\frac{1}{n_2} \right)}}$$

چنانچه ذکر شد یکی از فرض‌های آزمون t ، همگنی واریانس است و زمانی که واریانس‌ها با هم برابر نباشند، آزمون t با استفاده از فرمول زیر محاسبه می‌شود:

$$t = \frac{\overline{X}_1 - \overline{X}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$$

در صورتی که تعداد (حجم نمونه) در هر دو گروه مساوی باشد، آزمون t با استفاده از فرمول زیر محاسبه می‌شود:

$$t = \frac{\overline{X}_1 - \overline{X}_2}{\sqrt{\frac{s_1^2 + s_2^2}{n}}}$$

مثال: پژوهشگری علاقمند است تأثیر دو روش مختلف تدریس را در یادگیری ریاضی دانش‌آموزان در کلاس اول مورد آزمون قرار دهد. به همین منظور از بین دانش‌آموزان کلاس اول به‌طور تصادفی دو نمونه (۱۵ نفری) انتخاب کرد. برای یک گروه روش تدریس بحث گروهی (A) و برای نمونه دوم روش تدریس حل مسئله (B) را به‌کار برد. در پایان برای اینکه مشخص کند که کدام روش تدریس مؤثرتر بوده آزمونی به عمل آورد که نتایج آن در جدول زیر ارائه شده است:

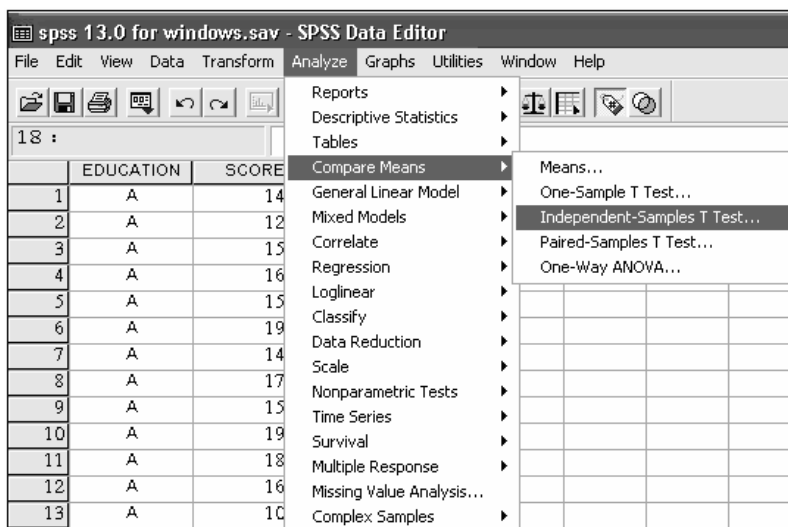
۱۷	۱۶	۱۰	۱۶	۱۸	۱۹	۱۵	۱۷	۱۴	۱۹	۱۵	۱۶	۱۵	۱۲	۱۴	روش تدریس (A)
۱۰	۱۰	۱۰	۱۳	۱۲	۱۷	۱۶	۱۴	۱۸	۱۵	۱۴	۱۴	۱۷	۱۶	۱۲	روش تدریس (B)

مراحل اجرای آزمون t مستقل در Spss به شرح زیر می‌باشد:

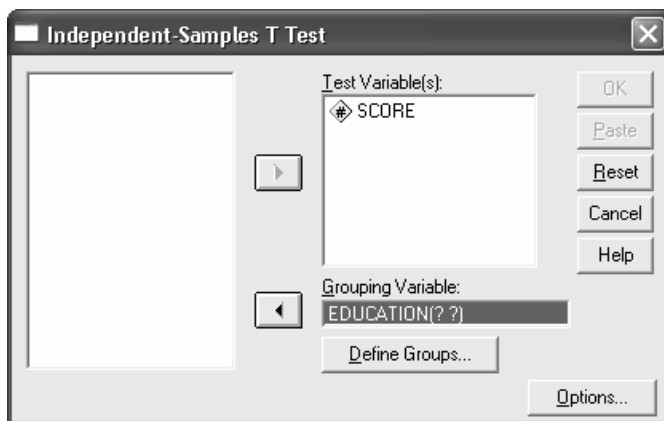
Analyze

Compare Means

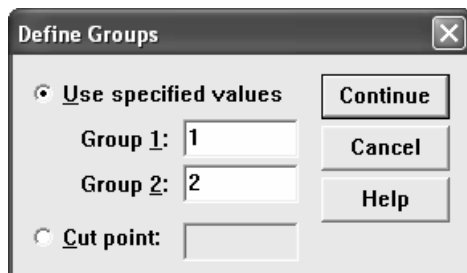
Independent- Samples t test



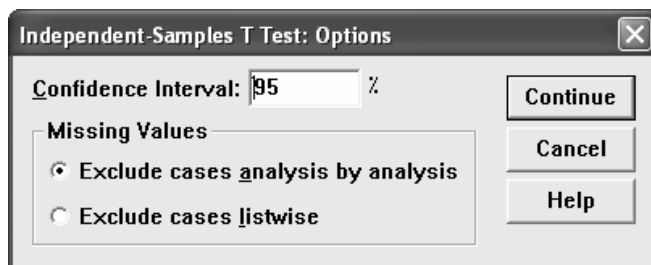
بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده متغیر (معدل) را به قسمت (Test Variable(s)) و متغیر (جنسیت) را به قسمت (Grouping Variable) منتقل کنید در جلوی جنسیت [؟؟] قرار دارد که باید جنسیت را تعریف کرد. برای تعریف دو گروه دانش‌آموزان دختر و پسر بر روی گزینه (Define Grouping...) کلیک کنید تا پنجره زیر آشکار شود:



که در این مثال، چون به بررسی ۲ گروه پرداخته‌اید در قسمت (Define Grouping) قسمت (Group 1) عدد یک و در قسمت (Group 2) عدد دو را قرار دهید و (Continue) را فعال کنید. با کلیک بر روی گزینه (Continue) به پنجره گفتگوی (Independent Samples t tests) باز گردید و در این قسمت برای تعیین سطح معنی‌داری گزینه (Options) را کلیک نمایید تا پنجره زیر آشکار شود. در این پنجره سطح معنی‌داری را انتخاب کنید:



بعد از انتخاب سطح معنی‌داری گزینه (Continue) فعال نماید تا به پنجره (Independent Samples t tests) باز گردد و گزینه (OK) را کلیک کنید و نهایت، نتایج به صورت خروجی زیر ظاهر خواهد شد:

T-Test

Group Statistics

	EDUCATION	N	Mean	Std. Deviation	Std. Error Mean
SCORE	A	15	15.5333	2.44560	.63145
	B	15	13.8667	2.66905	.68914

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
									Lower	Upper
SCORE	Equal variances assumed	.359	.554	1.783	28	.085	1.66667	.93469	-.24796	3.58130
	Equal variances not assumed			1.783	27.789	.085	1.66667	.93469	-.24862	3.58195

این خروجی با زیر جدولی آغاز می‌شود (Group Statistics) که آماره‌هایی برای دو نمونه ارائه می‌کند از جمله این آماره‌ها میانگین (۱۵/۵۳ و ۱۳/۸۶) و انحراف استاندارد (۲/۴۴ و ۲/۶۶) خطای معیار میانگین (۰/۶۳۱ و ۰/۶۸۹) نشان داده شده است. زیر جدول دوم (Independent Samples test) با ارائه نمودن مقدار (P-Value) و مقدار t آن Sig(2- tailed) و فاصله اطمینان ۹۵٪ برای اختلاف به این سؤال پاسخ می‌دهد. تمامی این مقادیر در دو حالت با فرض برابری واریانس‌ها (Equal Variance assumed) و بدون برابری واریانس‌ها (Equal Variance not assumed) ارائه شده است.

در ستون اول که عنوان (Leven's test for Equality of Variance) دارا می‌باشند برای بررسی فرض یکنواختی واریانس‌ها در آزمون t است، اگر آزمون (Leven's test...) معنی‌دار نباشد واریانس‌ها را می‌توان یکنواخت فرض کرد و می‌توان از مقادیر t ارائه شده در خطی که واریانس‌ها برابر فرض شده‌اند، استفاده کرد بنابراین:

- اگر F محاسبه شده از جدول F کوچک‌تر باشد چنین استنباط می‌شود که واریانس‌ها یکسانند و تفاوت آن‌ها از نظر آماری معنی‌دار نیست و آزمون t نرمال با واریانس‌های برابر گزارش می‌شود. (Equal Variance assumed)
- اگر F محاسبه شده از F جدول بزرگ‌تر باشد چنین استنباط می‌شود که واریانس‌ها برابر نیستند و آزمون t نرمال با واریانس‌های برابر گزارش می‌شود.

(Equal Variance not assumed)

در این آزمون، F لوین معنی‌دار نیست بنابراین مقدار t محاسبه شده با فرض برابری واریانس‌ها گزارش می‌شود. مقدار $t = (1/78)$ با درجه آزادی $df = 28$ و P -Value برابر با $0/085$ تفاوت معنی‌دار میانگین‌ها را نشان نمی‌دهد.

تحلیل واریانس

آزمون t مستقل برای آزمون معنی‌دار بودن تفاوت میانگین دو گروه به کار می‌رود. اما اگر ما با بیش از دو گروه میانگین سروکار داشته باشیم برای مقایسه میانگین‌ها از تحلیل واریانس استفاده می‌شود.

هدف از تحلیل واریانس دو متغیری آن است که از طریق مقایسه حداقل بیش از دو میانگین به دست آمده در یک طرح چند عاملی، مشخص کند که آیا تفاوت‌های موجود بین میانگین‌ها حاصل شانس و تصادف است یا ناشی از تأثیر عمل آزمایش؟

پیش‌نیازهای طرح

شرط لازم برای انجام تحلیل واریانس آن است که داده‌ها به صورت زیر جمع‌آوری شده باشند:

۱. وجود دو متغیر مستقل، با دو یا بیش از دو سطح

۲. وجود تفاوت کمی یا کیفی بین سطوح‌های هر یک از متغیرها مستقل
۳. هر فرد آزمودنی فقط در یکی از خانه‌های طرح مورد مشاهده قرار می‌گیرد و آزمودنی‌ها بیانگر یک نمونه تصادفی از جامعه مشخص شده در هر خانه جدول می‌باشند.

فرض‌ها

۱. برای استفاده از تحلیل واریانس باید فرض‌های زیر برقرار باشد:
 ۱. نمره هر فرد آزمودنی مستقل از نمره افراد دیگر است.
۲. نمره‌های موجود در هر یک از خانه‌های طرح از جامعه‌ای انتخاب شده‌اند که توزیع نمره‌های در آن جامعه شکل نرمال دارد به عبارت دیگر فرض شده نمره‌های هر یک از خانه‌های طرح به صورت نمونه‌گیری تصادفی از جامعه‌ای که دارای توزیع پراکندگی نرمال است انتخاب شده‌اند.
۳. واریانس نمره‌های خانه‌های جدول با هم یکسان (همگن) باشند. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۲۲۱).

$$F = \frac{MS_B}{MS_w}$$

$$MS_B = \frac{SS_B}{df_B} (df_B = k - 1)$$

$$MS_w = \frac{SS_w}{df_{Bw}} (df_w = n - k)$$

واریانس داخل گروه‌ها شاخصی از واریانس خطای نمونه‌گیری و واریانس بین گروهی، یعنی واریانس بین میانگین‌های ناشی از تأثیر متغیر آزمایشی را نشان می‌دهد. اگر واریانس درون گروهی بزرگ‌تر از واریانس بین گروهی باشد، می‌توان نتیجه گرفت که این تفاوت ناشی از خطاست و بالعکس.

مثال: دریک بررسی برای پی بردن به رابطه بین مدت زمان مطالعه و میزان یادگیری، از میان دانشجویان رشته روان‌شناسی به طور تصادفی تعدادی انتخاب و در ۴ گروه که از مدت زمان مطالعه با هم فرق داشتند به دانشجویان مطالبی در زمینه هیجان داده شد تا مطالعه کنند. در پایان مدت، از آن‌ها آزمون پیشرفت به عمل آمد. داده‌های مربوط به بررسی میزان یادگیری و مدت زمان مطالعه (۱۰ دقیقه، ۲۰ دقیقه، ۴۵ دقیقه و ۷۵ دقیقه) در جدول زیر ذکر شده است: (شیولسون ترجمه کیامتش، ۱۳۸۲: ۲۳۴).

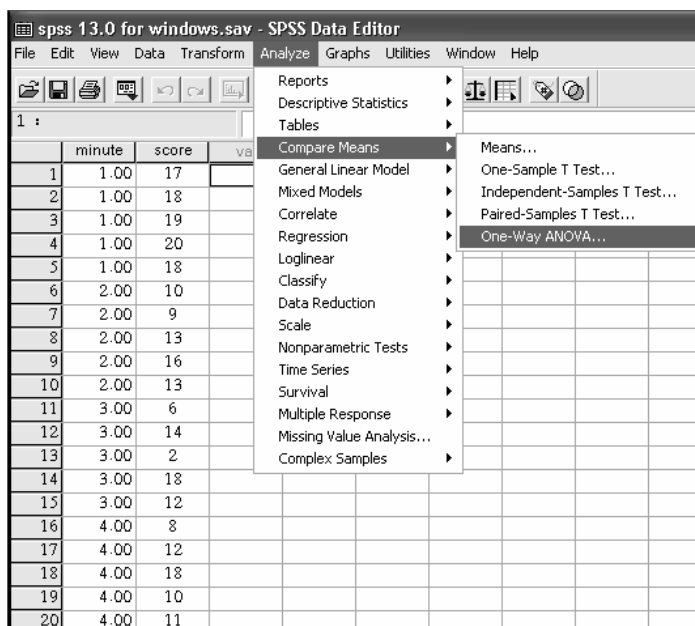
۱۰ دقیقه	۲۰ دقیقه	۴۵ دقیقه	۷۵ دقیقه
۱۷	۱۰	۶	۸
۱۸	۹	۱۴	۱۲
۱۹	۱۳	۲	۱۸
۲۰	۱۶	۱۸	۱۰
۱۸	۱۳	۱۲	۱۱

برای اجرای آزمون فوق در Spss مراحل زیر را انجام دهید:

Analyze

Compare Means

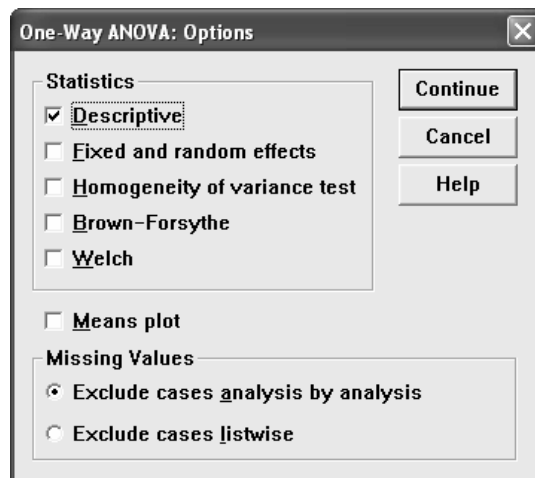
One- Way ANOVA



بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره (One-Way ANOVA) متغیر وابسته (نمره) را به قسمت (Dependent List) و گروه را به قسمت (Factor) منتقل کنید برای محاسبه آمار توصیفی، گزینه (Options...) را فعال نمایید تا پنجره زیر آشکار شود:



در پنجره: (One-Way ANOVA Options) گزینه (Descriptive) را انتخاب کرده و گزینه (Continue) را فعال نمایید تا به پنجره (One-Way ANOVA) باز گردد سپس (OK) را کلیک نموده تا نتایج به صورت خروجی زیر مشاهده گردد:

Oneway

Descriptives

score								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
10 minute	5	18.40	1.140	.510	16.98	19.82	17	20
20 minute	5	12.20	2.775	1.241	8.75	15.65	9	16
45 minute	5	10.40	6.387	2.857	2.47	18.33	2	18
75 minute	5	11.80	3.768	1.685	7.12	16.48	8	18
Total	20	13.20	4.841	1.082	10.93	15.47	2	20

ANOVA

score					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	189.200	3	63.067	3.942	.028
Within Groups	256.000	16	16.000		
Total	445.200	19			

چون مقدار F مشاهده شده (۳/۹۲۴) از مقدار F بحرانی بزرگ‌تر است. بنابراین نتیجه می‌گیریم که بین مدت زمان مطالعه و میزان یادگیری تفاوت معنی‌داری مشاهده می‌شود که این تفاوت ناشی از روش نمونه‌گیری نمی‌باشد بلکه این تفاوت ناشی از تأثیر متغیر مستقل است.

مقایسه‌های بعد از تجربه

مقایسه‌های بعد از تجربه به پژوهشگر امکان می‌دهد تا میانگین‌های مختلف را با یکدیگر مقایسه نماید و علت معنی‌دار بودن مقدار F را در تحلیل واریانس پیدا کند. مقایسه‌های بعد از تجربه در واقع روش‌های کشف محل قرار گرفتن تفاوت یا تفاوت‌های موجود بین میانگین‌ها می‌باشد، تنها شرط لازم برای استفاده از روش مقایسه‌های پسین (بعد از تجربه) به عنوان یک آزمون آماری، معنی‌دار بودن مقدار F مشاهده شده در تحلیل واریانس است و چنانچه آزمون معنی‌دار باشد، مقایسه‌های پسین این امکان را فراهم می‌آورند تا آزمون‌های متعددی در سطح معنی‌داری تعیین شده، انجام گیرد (شیولسون ترجمه کیامنش، ۲۴۶:۱۳۸۲).

برای انجام مقایسه‌های بعد از تجربه روش‌های مختلفی مانند شفه، توکی (HSD)، دونت، نیومن کولس، (LSD) و.. استفاده می‌شود و چون آزمون شفه و توکی بیشتر از سایر آزمون‌ها به کار گرفته می‌شود، بنابراین به معرفی این دو آزمون می‌پردازیم:

آزمون شفه

آزمون شفه جهت مقایسه یک به یک میانگین‌ها یا مقایسه ترکیب چندتایی میانگین‌ها استفاده می‌شود. (شیولسون ترجمه کیامش، ۱۳۸۲: ۲۵۳).

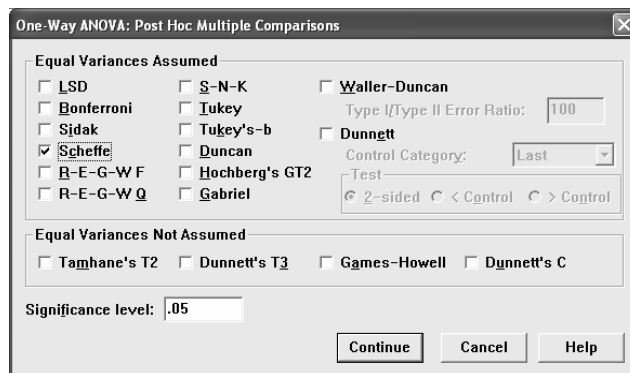
$$\tau = \frac{\hat{c}}{\sqrt{\frac{MS_w}{N} (w_1^2 + w_2^2 + \dots + w_k^2)}}$$

مطالعه مربوط به «مدت زمان صرف شده جهت مطالعه» در مثال تحلیل واریانس را در نظر بگیرید اکنون با استفاده از آزمون شفه به بررسی آن می‌پردازیم. (شیولسون ترجمه کیامش، ۱۳۸۲: ۲۵۷).

برای اجرای آزمون شفه در Spss مراحل زیر را انجام دهید:

Analyze
One-Way ANOVA
Compare Means
Post Hoc....

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره فوق گزینه (scheffe) را انتخاب کنید و نتایج به صورت خروجی زیر ظاهر خواهد شد:

Post Hoc tests

Multiple Comparisons

Dependent Variable: score

Scheffe

(I) minute	(J) minute	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
10 minute	20 minute	6.200	2.530	.154	-1.69	14.09
	45 minute	8.000*	2.530	.046	.11	15.89
	75 minute	6.600	2.530	.120	-1.29	14.49
20 minute	10 minute	-6.200	2.530	.154	-14.09	1.69
	45 minute	1.800	2.530	.916	-6.09	9.69
	75 minute	.400	2.530	.999	-7.49	8.29
45 minute	10 minute	-8.000*	2.530	.046	-15.89	-1.11
	20 minute	-1.800	2.530	.916	-9.69	6.09
	75 minute	-1.400	2.530	.958	-9.29	6.49
75 minute	10 minute	-6.600	2.530	.120	-14.49	1.29
	20 minute	-.400	2.530	.999	-8.29	7.49
	45 minute	1.400	2.530	.958	-6.49	9.29

*. The mean difference is significant at the .05 level.

خروجی آزمون شفه تفاوت بین میانگین‌ها، خطاهای معیار آن‌ها، sig و فاصله اطمینان هر جفت را نشان می‌دهد. با مشاهده نتایج، می‌توان گفت بین میانگین گروهی که ۱۰ دقیقه مطالعه کرده با میانگین گروهی که ۴۵ دقیقه مطالعه کرده‌اند، تفاوت معنی‌داری وجود دارد.

آزمون توکی (HSD)

آزمون توکی (HSD) از آزمون تفاوت معنی‌دار حقیقی برای مقایسه تمام حالت‌های ممکن یک به یک میانگین‌ها در سطح معنی‌دار استفاده می‌شود. آزمون توکی (HSD) از طریق مقایسه مقدار تفاوت بین هر جفت از میانگین‌ها با مقدار (HSD) انجام می‌شود. مقدار توکی با استفاده از فرمول زیر محاسبه می‌شود (شیولسون ترجمه کیامنش، ۱۳۸۲: ۲۵۸).

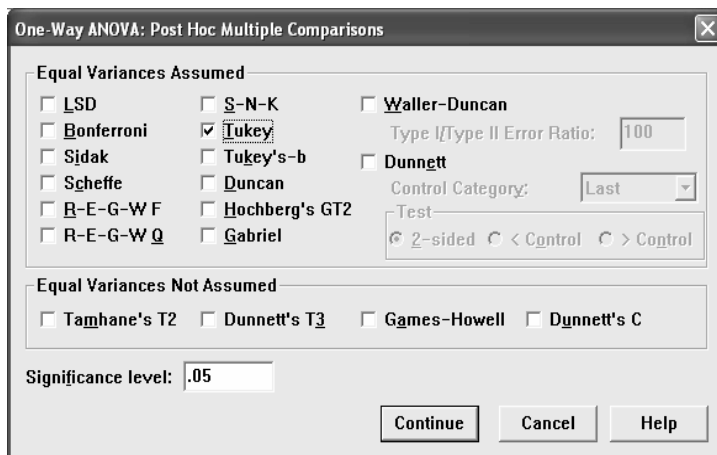
$$HSD = q \sqrt{\frac{MS_w}{N}}$$

مثال: با استفاده از داده‌های بررسی شده «مدت زمان صرف شده جهت مطالعه» در مثال تحلیل واریانس از طریق آزمون توکی (HSD) تمام حالت‌های ممکن مقایسه یک به یک میانگین‌ها در سطح ۰/۰۵ را انجام دهید.

برای اجرای آزمون توکی فوق در Spss مراحل زیر را انجام دهید:

Analyze
One- Way ANOVA
Compare Means
Post Hoc...

بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره فوق گزینه (Tukey) را انتخاب کنید و نتایج به صورت خروجی زیر

ظاهر خواهد شد:

Post Hoc tests Multiple Comparisons

Dependent Variable: score
Tukey HSD

(I) minute	(J) minute	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
10 minute	20 minute	6.200	2.530	.107	-1.04	13.44
	45 minute	8.000*	2.530	.028	.76	15.24
	75 minute	6.600	2.530	.080	-.64	13.84
20 minute	10 minute	-6.200	2.530	.107	-13.44	1.04
	45 minute	1.800	2.530	.891	-5.44	9.04
	75 minute	.400	2.530	.999	-6.84	7.64
45 minute	10 minute	-8.000*	2.530	.028	-15.24	-.76
	20 minute	-1.800	2.530	.891	-9.04	5.44
	75 minute	-1.400	2.530	.944	-8.64	5.84
75 minute	10 minute	-6.600	2.530	.080	-13.84	.64
	20 minute	-.400	2.530	.999	-7.64	6.84
	45 minute	1.400	2.530	.944	-5.84	8.64

*. The mean difference is significant at the .05 level.

اگر مقدار تفاوت از مقدار (HSD) بزرگ‌تر باشد فرضیه صفر را رد کنید، برای تفاوت‌های بزرگ‌تر از (HSD) نتیجه می‌گیریم که تفاوت بین میانگین‌هایی که در جدول نتایج با علامت (*) مشخص شده‌اند در سطح ۰/۰۵ معنی‌دار است. در واقع تفاوت واقعی بین گروهی است که ۴۵ دقیقه مطالعه داشته، با سایر گروه‌ها با مدت زمان‌های مختلف می‌باشد. این تفاوت‌ها نشان می‌دهد که ۴۵ دقیقه مطالعه، مناسب‌ترین مدت مطالعه می‌باشد.

خودآزمایی

۱. با استفاده از روش نمونه‌گیری تصادفی، یک نمونه تصادفی از میان دانشجویان یک دانشگاه انتخاب کرده ایم. نمرات آزمون SAT این گروه در زیر ارائه گردیده است
۱۲۰۰-۱۱۶۰-۱۲۰۰-۱۱۸۵-۱۱۵۰-۱۱۵۰-۱۰۸۰-۱۱۲۰-۱۱۰۰-۱۱۵۰

بر اساس این داده‌ها، آیا دانشجویان دانشگاه فوق‌الذکر نماینده جامعه افراد شرکت‌کننده در آزمون SAT هستند یا با جامعه تفاوت دارند؟

۲. ۴۰ نفر دانش‌آموز به‌صورت تصادفی در معرض چهار روش مختلف تدریس قرار داده شده است. میزان یادگیری چهار گروه پس از آموزش مورد اندازه‌گیری قرار گرفت. با استفاده از یک آزمون آماری مناسب، تعیین کنید که اختلاف معنی‌داری بین میانگین چهار گروه وجود دارد یا خیر؟ (دلاور، ۱۳۸۰: ۴۷۷)

گروه چهار	گروه سه	گروه دو	گروه یک
۱۰	۱۰	۱۵	۱۴
۹	۱۰	۱۳	۱۱
۶	۱۰	۱۲	۸
۷	۱۰	۱۰	۷
۶	۶	۷	۶
۵	۶	۶	۴
۳	۵	۶	۳
۲	۵	۵	۲
۲	۵	۳	۱
۱	۵	۲	۱

۳. پژوهشگری علاقمند است تأثیر دو روش مختلف تقویت را در یادگیری ریاضی دانش‌آموزان در کلاس سوم مورد آزمون قرار دهد. به همین منظور از بین دانش‌آموزان کلاس سوم به صورت تصادفی، دو نمونه انتخاب می‌کند. برای یکی از نمونه‌ها روش تقویت A و برای نمونه دوم روش B را به کار می‌برد در پایان آزمایش، میزان یادگیری هر دو نمونه را به وسیله آزمون مورد اندازه‌گیری قرار می‌دهد. جدول زیر نتایج اجرای آزمون را برای دو نمونه نشان می‌دهد. (دلاور، ۱۳۸۰: ۴۰۴)

گروه ۱ روش A	گروه ۲ روش B
۱۶	۱۰
۱۲	۷
۱۱	۷
۸	۶
۹	۵
۴	

۴. به منظور آزمون این فرضیه که (روش تن آرامی، اضطراب امتحان دانش‌آموزان را کاهش می‌دهد) ۵ دانش‌آموز به صورت تصادفی انتخاب شده و اضطراب امتحان آن‌ها به وسیله یک وسیله معتبر قبل و بعد از روش تن آرامی اندازه‌گیری شده است. نمره اضطراب قبل و بعد از تن آرامی این ۵ نفر در جدول زیر نمایان است، با استفاده از یک آزمون آماری مناسب و در سطح معنی‌داری ۰/۰۱ فرضیه پژوهش را آزمون کنید؟

نمره اضطراب بعد از تن آرامی	نمره اضطراب قبل از تن آرامی
۱۶	۲۵
۱۵	۳۰
۲۰	۴۰
۳۶	۵۰
۳۰	۳۸

۵. روان‌شناسی علاقمند است تأثیر یک روش آموزش را در بالا بردن بهره هوشی دانش‌آموزان عقب مانده مورد آزمون قرار دهد. به همین منظور ۱۲ نفر دانش‌آموز عقب مانده را به صورت تصادفی انتخاب می‌کند. ابتدا بهره هوشی آن‌ها را

اندازه‌گیری می‌کند سپس روش آموزشی خود را اجرا می‌کند. پس از اتمام آموزش، مجدداً بهره‌هوشی آن‌ها را مورد اندازه‌گیری قرار می‌دهد اطلاعات جمع‌آوری شده به شرح زیر در دست است. با یک آزمون آماری مناسب و با احتمال $0/01$ درصد خطا فرضیه صفر را آزمون کنید. (دلاور، ۱۳۸۰: ۴۲۶)

قبل از آزمون	بعد از آزمون
۹۶	۱۰۱
۸۹	۹۴
۸۱	۸۸
۸۵	۸۵
۹۶	۱۰۲
۹۵	۱۰۰
۸۷	۹۰
۷۹	۸۹
۸۰	۸۵
۹۰	۹۸
۹۲	۱۰۰
۹۹	۱۰۵

فصل هشتم

آزمون‌های آماری ناپارامتریک

هدف‌های یادگیری

از دانشجو انتظار می‌رود که پس از خواندن فصل هشتم بتواند:

۱. تفاوت آزمون‌های آماری پارامتریک و ناپارامتریک را بیان کند.
۲. محاسن و معایب آزمون‌های ناپارامتریک را توضیح دهد.
۳. آزمون کای دو را محاسبه کند.
۴. کاربرد، محاسن و معایب آزمون Uمان - ویتنی را بیان کند.
۵. آزمون Uمان - ویتنی را محاسبه کند.
۶. کاربرد، محاسن و معایب آزمون کروسکال - والیس را بیان کند.
۷. آزمون کروسکال - والیس را محاسبه کند.
۸. کاربرد، محاسن و معایب آزمون ویلکاکسون را بیان کند.
۹. آزمون ویلکاکسون را محاسبه کند.
۱۰. کاربرد، محاسن و معایب آزمون فریدمن را بیان کند.
۱۱. آزمون فریدمن را محاسبه کند.

آزمون‌های آماری ناپارامتریک

در پژوهش غالباً داده‌هایی جمع‌آوری می‌شود که مقیاس اندازه‌گیری آن‌ها اسمی یا رتبه‌ای است یا گاهی متغیری مورد اندازه‌گیری قرار می‌گیرد که مقیاس آن فاصله‌ای است ولی توزیع آن در جامعه نرمال نیست. در چنین شرایطی پژوهشگر ملزم به استفاده از آزمون‌های ناپارامتریک است. استفاده از آزمون‌های ناپارامتریک به هیچ پیش‌فرضی

مشخصی نیاز ندارد ولی انتخاب آزمون‌های پارامتریک به پیش‌فرض‌های مشخصی نیاز دارد و در صورتی که این پیش‌فرض‌ها تحقق نیابد نمی‌توان از آزمون پارامتریک استفاده کرد. مهم‌ترین پیش‌فرض‌های آزمون پارامتریک چنانچه قبلاً ذکر شد عبارتند از:

۱. مقیاس حداقل فاصله‌ای باشد.

۲. توزیع نمره‌ها در هر دو جامعه نرمال باشد.

۳. واریانس‌ها برابر باشد.

اگر واریانس‌ها برابر نباشد اما تعداد نمونه مساوی باشد از آزمون‌های پارامتریک می‌توان استفاده کرد. اما وقتی واریانس‌ها همگن نباشد و تعداد افراد نمونه هم مساوی نباشد. نمی‌توان از آزمون‌های پارامتریک استفاده کرد و حتماً باید از آزمون ناپارامتریک استفاده شود. (شیولسون ترجمه کیامنش جلد دوم ۲۲۳)

در حجم‌های کوچک ($n < 10$) نمونه کمتر از ۱۰ نفر استفاده از آزمون‌های ناپارامتریک، منطقی‌تر و بهتر است. (شیولسون ترجمه کیامنش، ۲۱۶:۱۳۸۲).

محاسن آزمون‌های آماری ناپارامتریک

۱. گزاره‌های احتمالاتی که در بیشتر آزمون‌های ناپارامتریک به‌دست می‌آید «احتمالاً قطعی» هستند.

۲. اگر حجم نمونه به کوچکی ($N=6$) باشد، راهی به جزء استفاده از آزمون‌های ناپارامتریک نیست.

۳. تعدادی از آزمون‌های آماری ناپارامتریک وجود دارند که برای تجزیه و تحلیل نمونه‌هایی مناسب هستند که از مشاهده چندین جامعه مختلف حاصل شده است.

این در حالی است که هیچ‌یک از آزمون‌های پارامتریک نمی‌توانند داده‌های حاصل از این نوع را تجزیه و تحلیل کنند، مگر اینکه یک رشته مفروضات غیر واقعی ارائه کنند.

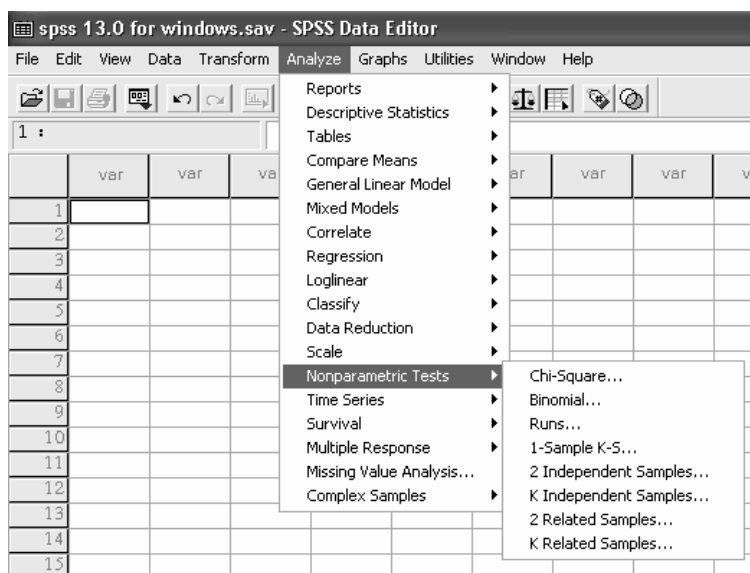
۴. آزمون‌های ناپارامتریک گوناگونی وجود دارند که برای تجزیه و تحلیل داده‌هایی که به‌صورت رتبه‌ای هستند و همچنین داده‌هایی که نمرات به ظاهر شمارشی آن‌ها رتبه را دارا هستند، مناسبند بدین معنی که محقق ممکن است فقط بتواند بگوید که آزمونی‌های تحقیق او، فلان خصیصه را کم یا زیاد دارند بدون اینکه بتواند بگوید چقدر کمتر یا زیادتر می‌باشد.

۵. آزمون‌های ناپارامتریک از نظر یادگیری و کاربرد به مراتب ساده‌تر و آسان‌تر از آزمون‌های پارامتریک است.

معایب آزمون‌های آماری ناپارامتریک

۱. اگر تمام مفروضات مدل آماری پارامتریک برآورده شوند و اگر مقیاس سنجش ما به حد لزوم یعنی به حد فاصله‌ای برسند، در آن صورت به‌کاربردن آزمون‌های ناپارامتریک اتلاف مقداری از داده‌ها را به همراه دارد. میزان اتلاف به‌وسیله توان - کارآمدی آزمون‌های ناپارامتریک بیان می‌شود.
۲. هنوز روش‌های ناپارامتریک که برای آزمون تعامل‌ها در مدل تحلیل واریانس مناسب باشد ابداع نشده است.
۳. اعتراض دیگر به آزمون‌های ناپارامتریک این است که جدول‌ها مقادیر معنی‌دار مربوط به آزمون‌های ناپارامتریک در منابع مختلف، پراکنده است و اکثر این منابع بسیار تخصصی هستند. بنابراین چنین منابعی به سختی در دسترس علاقه‌مندان قرار می‌گیرد. (سیگل ترجمه کریمی، ۱۳۸۰: ۴۳)

زیر منوی (Nonparametric Test) برای اجرای آزمون‌های ناپارامتری در منوی (Statistics) گنجانده شده است. شکل زیر فرمان‌های مختلف این زیر منو را نشان می‌دهد:



آزمون ناپارامتریک خی دو (Chi-Square)

آزمون χ^2 یکی از آزمون‌های ناپارامتریک مهم است در آزمون‌های ناپارامتریک χ^2 هم برای مقایسه تک نمونه‌ای (نرمال بودن جامعه)، مقایسه دو گروه مستقل و یک متغیر آزمایشی (به جای U مان-وینتنی)، برای مقایسه سه گروه یا بیشتر (k گروه) با یک متغیر آزمایشی (به جای کروسکال والیس) به کار می‌رود.

$$\chi^2 = \sum_{i=1}^r \sum_{j=1}^c \frac{(O_{ij} - E_{ij})^2}{E_{ij}}$$

کاربرد

در بعضی تحقیقات، پژوهشگران با مشاهده و شمارش تعداد آزمونی‌ها (یعنی با فراوانی) سرو کار دارند. مثلاً تعداد اعضای احزاب مختلف یا تعداد رأی‌دهندگان زمانی که پژوهشگران با نظر و علاقه و نگرش افراد سرو کار دارد. علاوه بر موارد کاربرد χ^2 که ذکر شد، عمده ترین ملاک که باید به آن توجه شود، مقیاس است. وقتی متغیرها در سطح مقیاس نامی (اسمی) و رتبه‌ای باشند از χ^2 استفاده کنید. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۳۵۴).

معایب

آزمون χ^2 زمانی به کار برده می‌شود که داده‌ها به صورت مقیاس اسمی باشد. هنگامی که مقیاس اندازه‌گیری ترتیبی یا فاصله‌ای است بهتر است از آزمون دیگری استفاده شود. (دلاور، ۱۳۸۰: ۴۹۸)

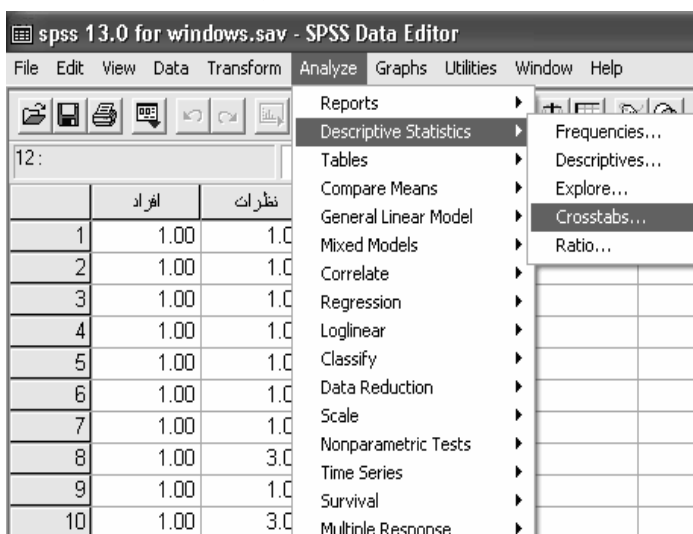
مثال: فرض کنید در یکی از کشورهای اروپایی، پژوهشگری علاقه‌مند است به چنین پرسشی پاسخ دهد: «آیا در زمینه تهیه بمب اتمی بین افرادی که در ارتش خدمت کرده اند و افرادی که در ارتش خدمت نکرده اند تفاوت عقیده وجود دارد؟» فرض کنید از طریق تصادفی از ۵۰ نفر پرسیده شده که:

۱. آیا در ارتش خدمت کرده اید؟ بلی، خیر

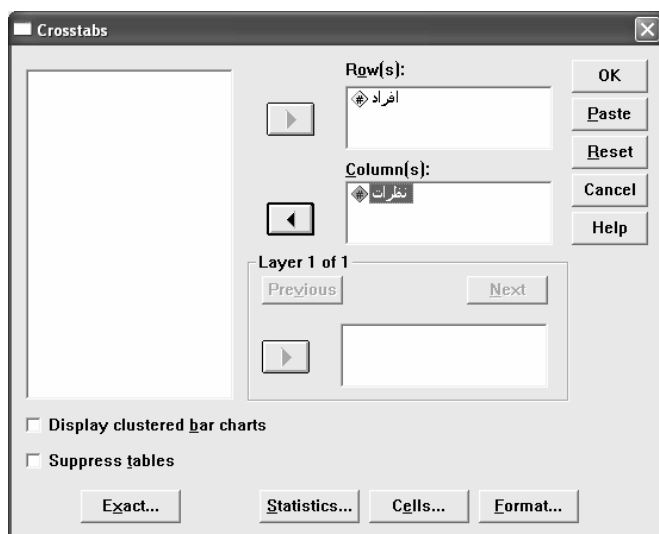
۲. آیا با تهیه بمب اتمی موافق هستید؟ بلی، خیر، بی نظر (شیولسون ترجمه کیامنش، ۳۵۵:۱۳۸۲).

نظرات افراد	بلی	خیر	بی نظر
نظامیان	۱۹	۶	۵
غیرنظامیان	۵	۸	۷

برای اجرای آزمون فوق در Spss مراحل زیر را انجام دهید



بعد از اجرای این فرمان شکل زیر ظاهر می شود:



در پنجره ظاهر شده یکی از متغیرها را به قسمت (Row(s)) و متغیر دیگر را به قسمت (Column(s)) وارد کنید، سپس گزینه (Statistics...) را کلیک کنید تا پنجره زیر باز شود:

در پنجره بالا گزینه (Chi-square) را انتخاب نمایید سپس گزینه (Continue) را فعال نمایید تا به پنجره (Crosstabs) بازگردید در این پنجره گزینه (OK) را کلیک کنید و نتایج به صورت خروجی زیر ظاهر خواهد شد:

Cross tabs

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
افراد * نظرات	50	100.0%	0	.0%	50	100.0%

Crosstabulation افراد * نظرات

			نظرات			Total
			بی نظری	خیر	بی نظری	
افراد	در ارتش خدمت کرده	Count	19	6	5	30
		Expected Count	14.4	8.4	7.2	30.0
	در ارتش خدمت نکرده	Count	5	8	7	20
		Expected Count	9.6	5.6	4.8	20.0
Total	Count	24	14	12	50	
	Expected Count	24.0	14.0	12.0	50.0	

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	7.068 ^a	2	.029
Likelihood Ratio	7.316	2	.026
Linear-by-Linear Association	5.701	1	.017
N of Valid Cases	50		

a.

در جدول (Crosstabulation) تعداد افراد پاسخگو و در جدول (Chi-square Tests) سطح معنی‌داری ذکر شده است که در این مثال، مقدار کای‌دو پیرسون برابر ۷/۰۶ است که این مقدار در سطح ۰/۰۵ معنی‌دار است. بنابراین بین نظرات نظامیان و غیر نظامیان در استفاده از بمب اتمی تفاوت معنی‌داری وجود دارد. و با مقایسه نظرات افراد می‌توان گفت که اکثر افرادی که در ارتش خدمت کرده اند، موافق استفاده از بمب اتمی هستند و اکثر افرادی که در ارتش خدمت نکرده اند مخالف استفاده از بمب اتمی هستند.

آزمون‌های ناپارامتریک مربوط به دو گروه مستقل (2 Independent samples)

آزمون t ویلکاکسون (من-ویتی) برای دو گروه مستقل

یکی از آزمون ناپارامتریک مربوط به دو گروه مستقل است که برای مقایسه میانگین دو جامعه، وقتی که مقیاس اندازه‌گیری رتبه‌ای است، مطلوب می‌باشد و در مقابل آزمون پارامتری t بوجود آمده است. زمانی که محقق قصد مقایسه دو گروه مستقل را دارد و فرض نرمال بودن و یکسانی واریانس‌ها برقرار نیست از آزمون U من-ویتی استفاده می‌شود (شیولسون ترجمه کیامنش، ۱۳۸۲: ۲۱۷).

$$Z = \frac{2}{\sqrt{\frac{n_1 n_2}{n(n-1)} \left(\frac{n^2 - n}{12} - \sum T \right)}}$$

کاربرد

این آزمون برای تشخیص تفاوت بین دو جامعه (مقایسه میانگین دو جامعه) که با استفاده از نمونه‌های تصادفی از همان جامعه انتخاب شده است و در صورتی که مقیاس

اندازه‌گیری صفت متغیر مورد مطالعه به صورت مقیاس ترتیبی باشد، به کار برده می‌شود.

محاسن

آزمون U من- ویتنی در مقایسه با سایر آزمون‌های ناپارامتریک مربوط به دو گروه مستقل قویتر است و می‌توان آن را در آزمون‌های یک دامنه‌ای و دو دامنه‌ای برای آزمایش فرض صفر بکار برد.

معایب

عیب U من- ویتنی این است که تنها می‌توانیم به مقایسه دو میانگین بپردازیم و اگر بیش از دو باشد، نمی‌توان از این آزمون استفاده کرد. بلکه باید از آزمون کروسکال والیس استفاده نماییم.

مثال: دو گروه از دانش‌آموزان کلاس درس بهداشت در دوره اول دبیرستان انتخاب شدند و در یک گروه طرز به هوش آوردن مصنوعی بیمار با روش بحث و سخنرانی (Lecture) و گروه دیگر با نمایش فیلم استریپ (Film) آموزش داده شد. در پایان دوره آموزش هر دو گروه با آزمون واحدی مورد آزمایش قرار گرفتند، اکنون با استفاده از آزمون U من- ویتنی در مورد تفاوت دو روش تدریس مورد بحث اظهار نظر کنید (پاشا شریفی نجفی زند، ۱۳۷۵: ۲۷۳)

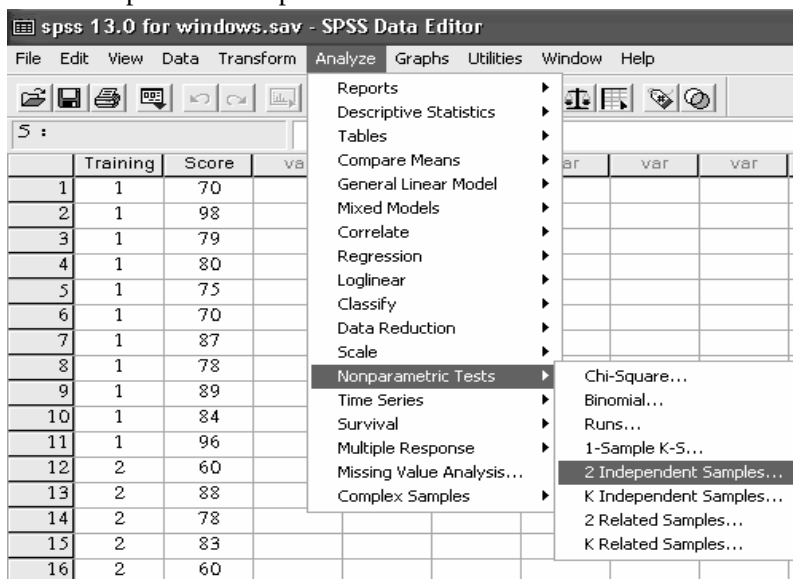
سخنرانی	نمایش فیلم
۶۰	۷۰
۸۸	۹۸
۷۸	۷۹
۸۳	۸۰
۶۰	۷۵
۵۴	۷۰
۸۰	۸۷
۵۴	۷۸
۵۰	۸۹
۸۰	۸۴
۶۰	۹۶
۶۵	

از آنجایی که پیش فرض‌های آزمون t برقرار نیست بنابراین پژوهشگر از آزمون ناپارامتریک U -من- ویتنی استفاده نموده است.
برای اجرای آزمون U -من- ویتنی مراحل زیر را انجام دهید:

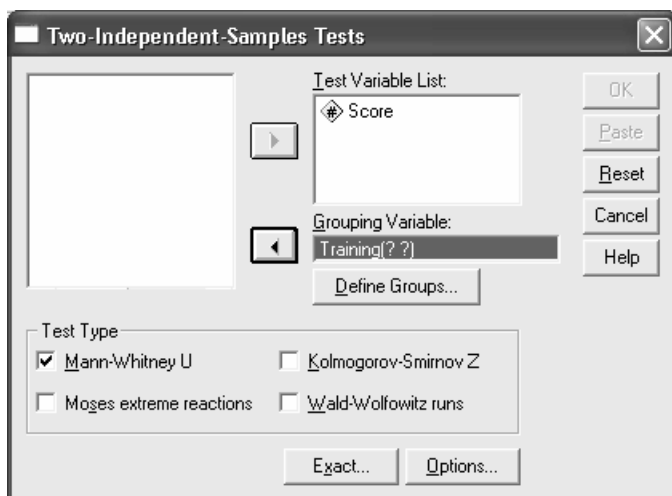
Analyze

Nonparametric test

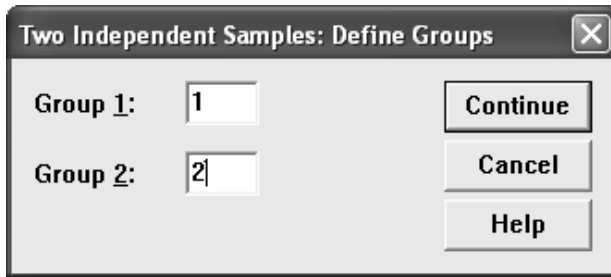
2 Independent Samples...



بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:

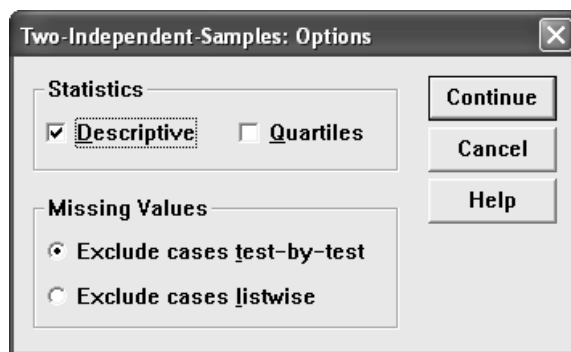


در پنجره ظاهر شده متغیر (رتبه) را به قسمت (Test Variable List) و متغیر آزمون (Examin) را به قسمت (Grouping Variable) منتقل کنید در جلوی آزمون [؟؟] قرار دارد که باید نظامیان را تعریف کرد برای تعریف دو گروه نظامیان بر روی گزینه (Define Grouping...) کلیک نمایید تا پنجره زیر آشکار شود:



در این مثال پژوهشگر چون به بررسی ۲ گروه پرداخته است در قسمت (Two Independent Samples: Define Grouping) در قسمت (Minimum) عدد یک و در قسمت (Maximum) عدد دو قرار می‌دهید سپس (Continue) را کلیک کنید.

با فعال کردن گزینه (Continue) به پنجره (Two Independent Samples Tests) باز می‌گردید و در این قسمت برای محاسبه آمار توصیفی گزینه (Options) را کلیک نمایید تا به پنجره زیر برسید در پنجره مذکور گزینه (Descriptive) را انتخاب کنید.



سپس گزینه (Continue) را فعال نمایید تا به پنجره گفتگوی (Two Independent Samples Tests) باز گردید در پنجره گفتگوی مذکور گزینه (Mann-Whitney U) را کلیک نمایید:

Test Type

☒ **Mann-Whitney U**

☐ **Kolmogorov-Smirnov Z**

☐ **Moses extreme reactions**

☐ **Wald-Wolfowitz runs**

سپس (OK) را کلیک نموده تا نتایج به صورت خروجی زیر مشاهده شود:

Mann-Whitney Test

Ranks

	Training	N	Mean Rank	Sum of Ranks
Score	Film	11	15.41	169.50
	Lecture	12	8.88	106.50
	Total	23		

Test Statistics^b

	Score
Mann-Whitney U	28.500
Wilcoxon W	106.500
Z	-2.314
Asymp. Sig. (2-tailed)	.021
Exact Sig.	
[2*(1-tailed Sig.)]	.019 ^a

a. Not corrected for ties.

b. Grouping Variable: Training

جدول نخست، تعداد نمونه در هر گروه، میانگین رتبه‌ها و جمع رتبه‌ها را نشان می‌دهد. در جدول دوم نیز میزان U ، Z و سطح پوشش دو سویه آورده شده است. مقدار U محاسبه شده برابر ۲۸ است که این مقدار در سطح ۰/۰۵ معنی‌دار نیست.

نکته: اگر حجم نمونه بزرگ‌تر از ۲۰ باشد توزیع نمونه‌گیری U تقریباً نرمال خواهد بود (با توجه به قضیه حد مرکزی با افزایش حجم نمونه، شکل توزیع نمونه‌گیری به نرمال نزدیک می‌شود) اگر مقدار U را به نمره Z معیار تبدیل کنیم، می‌توان از طریق ZU نیز

در مورد رد یا عدم رد فرضیه صفر قضاوت کنیم (شیولسون ترجمه کیامنش، ۱۳۸۲: ۲۲۴). نحوه قضاوت بدین گونه می باشد که چنانچه سطح پوشش دوسویه بر دو تقسیم شده و مقدار آن از ۲/۵ درصد کمتر باشد، تفاوت دو گروه را نتیجه گیری نمایید و در صورتی که این میزان از ۲/۵ درصد بیشتر باشد، نبود تفاوت بین دو گروه را نتیجه گیری کنید. که در این مثال مقدار z تقسیم بر ۲ برابر ۱/۱۵۷ است و چون این مقدار کمتر از ۲/۵ است، بنابراین نتیجه می گیریم که بین دو گروه تفاوت معنی داری وجود دارد.

آزمون های ناپارامتریک مربوط به دو گروه وابسته (2 Related samples)

آزمون T ویلکاکسون برای نمونه های وابسته (همبسته) مشابه آزمون t وابسته است. با این تفاوت که با استفاده از آزمون ناپارامتریک مربوط به دو گروه وابسته، می توان توزیع دومتغیر وابسته را بررسی کرد. آزمون ویلکاکسون یکی از آزمون ناپارامتریک مربوط به دو گروه وابسته است که در آزمون ویلکاکسون علاوه بر جهت تغییرات، اندازه نسبی این تغییرات (تفاوت ها) نیز مورد توجه است. یعنی با آن می توان درباره اینکه کدام یک از اجزاء از دیگری بزرگ تر است، قضاوت کرد.

$$z = \frac{T - \frac{n(n+1)}{4}}{\sqrt{\frac{n(n+1)(2n+1)}{24}}}$$

کاربرد

این آزمون نظیر آزمون U مان-ویتنی است، وقتی بکار برده می شود که داده ها دست کم در مقیاس ترتیبی و به صورت علامت گذاری شده (جفت های جور شده) باشند، توان آزمون ویلکاکسون در داده های پارامتری تقریباً با آزمون t استودنت وابسته متناظر است و توان آن برابر ۹۵/۵٪ آزمون t می باشد. همچنین اگر بخواهیم علاوه بر جهت تفاوت های موجود بین گروه های همتا، اندازه نسبی این تفاوت ها را نیز بدست آوریم، از این آزمون استفاده می کنیم. (سیگل ترجمه کریمی، ۱۳۸۰: ۱۰۷)

محاسن

آزمون ویلکاکسون بیانگر این پاسخ است که کدام جزء از جزء دیگر بزرگ‌تر است و همچنین تفاوت‌ها را به ترتیب قدر مطلق آن‌ها رتبه بندی می‌کند علاوه بر آن جفت‌هایی که بین اجزای آن‌ها تفاوت بزرگ‌تری وجود دارد را نشان می‌دهند که این از امتیاز آزمون ویلکاکسون می‌باشد و قابل ذکر است که در این آزمون علاوه بر جهت تفاوت‌ها، اندازه نسبی این تفاوت‌ها نیز منظور می‌شود.

معایب

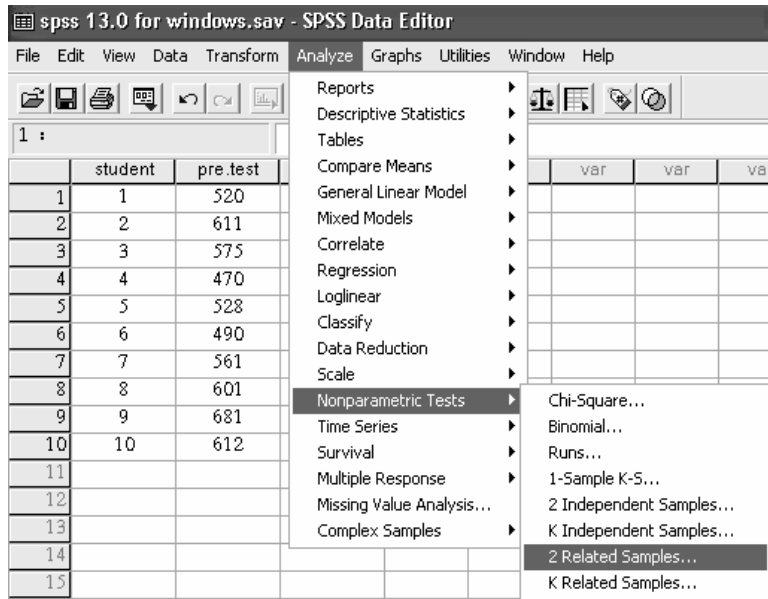
در این آزمون جفت‌هایی که نمره یکسانی دارند یعنی تفاوت آن‌ها صفر است از حجم نمونه حذف می‌شوند که این از معایب می‌باشد.

مثال: معلم یک کلاس، یک آزمون استاندارد شده را پیش از شروع تدریس در کلاس اجرا می‌کند. در پایان دوره آموزش نیز همان آزمون را درباره همان دانش‌آموزان اجرا می‌کند. هدف وی از تحقیق این است که میزان افزایش آموخته‌های شاگردان را در دوره آموزش تعیین کند که نتایج دو آزمون پیش آزمون و پس آزمون در جدول زیر آمده است:

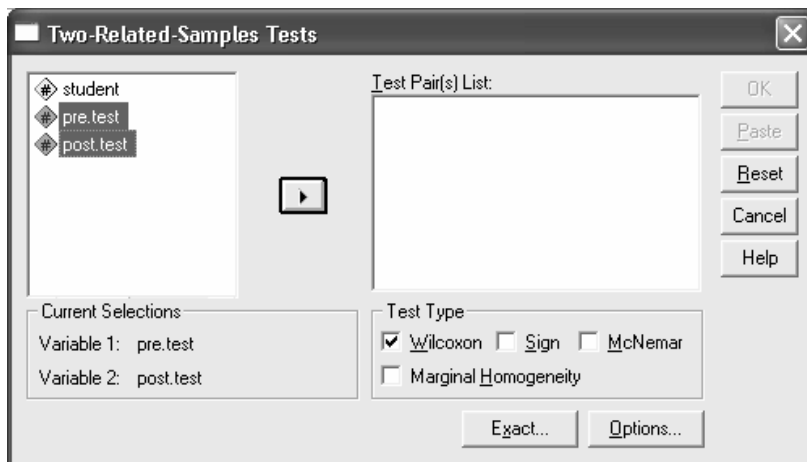
پس آزمون	پیش آزمون	دانش آموز
۵۳۰	۵۲۰	۱
۶۵۰	۶۱۱	۲
۶۰۲	۵۷۵	۳
۴۶۸	۴۷۰	۴
۵۳۱	۵۲۸	۵
۵۲۵	۴۹۰	۶
۵۷۷	۵۶۱	۷
۶۱۲	۶۰۱	۸
۶۹۸	۶۸۱	۹
۵۹۸	۶۱۲	۱۰

مراحل اجرای آزمون ویلکاکسون در Spss به شرح زیر می‌باشد:

Analyze
Nonparametric test
2 Related Samples



بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده متغیرها را به قسمت (Test Pair(s) List) منتقل کنید و گزینه (Wilcoxon) را انتخاب نمایید و سپس (OK) را کلیک کنید و نتایج به صورت خروجی زیر ظاهر خواهد شد:

NPar Tests
Wilcoxon Signed Ranks Test

Ranks

		N	Mean Rank	Sum of Ranks
post.test – pre.test	Negative Ranks	2 ^a	3.00	6.00
	Positive Ranks	8 ^b	6.13	49.00
	Ties	0 ^c		
	Total	10		

a. post.test < pre.test

b. post.test > pre.test

c. post.test = pre.test

Test Statistics^b

	post.test – pre.test
Z	-2.191 ^a
Asymp. Sig. (2-tailed)	.028

a. Based on negative ranks.

b. Wilcoxon Signed Ranks Test

در جدول نخست تعداد و میانگین رتبه افرادی که نمرهای آنان در دو وضعیت افزایش، کاهش و یا ثابت مانده است آورده شده است.

a. post.test < pre.test

b. post.test > pre.test

c. post.test = pre.test

(a) بیانگر این است که دو نفر از دانش‌آموزان نمره پیش‌آزمون آن‌ها بزرگ‌تر از پس‌آزمون است.

(b) بیانگر این است که هشت نفر از دانش‌آموزان نمره پیش‌آزمون آن‌ها کوچک‌تر از پس‌آزمون است.

(c) بیانگر این است که هیچ‌کدام از دانش‌آموزان نمره پیش‌آزمون و پس‌آزمون آن‌ها مساوی نیست.

در جدول دوم مقدار Z و سطح پوشش آماره آزمون محاسبه شده است و چنانکه قدر مطلق Z محاسبه شده بیشتر از ۱/۹۶ باشد تفاوت معنی‌دار در پیش آزمون و پس آزمون را نتیجه‌گیری نمایید و چنانچه این مقدار کمتر از ۱/۹۶ باشد، تساوی دو وضعیت را نتیجه‌گیری کنید. از روی مقدار sig نیز می‌توان درباره معنی‌دار بودن تفاوت‌ها اظهار نظر کرد که در این مثال مقدار sig برابر ۰/۰۲۸ است که این مقدار در سطح ۰/۰۵ خطا معنی‌دار است و با ۹۵٪ اطمینان می‌توان گفت که بین نمرات پیش آزمون و پس آزمون تفاوت معنی‌دار وجود دارد.

آزمون‌های ناپارامتریک مربوط به k گروه مستقل (K Independent samples)

آزمون‌های تحلیل واریانس یک متغیری داده‌های رتبه‌بندی کروسکال - والیس (H یا T) یکی از آزمون‌های ناپارامتری مربوط به K گروه مستقل است که متمم آزمون U-من-ویتی می‌باشد و در مقابل آزمون F بوجود آمده و توزیع آن با χ^2 تقریباً یکی است و برای مقایسه میانگین رتبه‌های دو گروه یا بیشتر به کار می‌رود:

$$H = \frac{12}{N(N+1)} \sum_{i=1}^k \frac{R_i^2}{n_i} - 3(N+1)$$

کاربرد

این آزمون در مواردی بکار برده می‌شود که همانند آزمون F تعداد گروه‌ها دو یا بیش از دو باشد و علاوه بر آن مقیاس اندازه‌گیری دست‌کم ترتیبی باشد.

محاسن

برای مقایسه بیش از دو گروه بکار می‌رود و نیز در صورتی که فرض یکسان بودن واریانس‌ها برقرار نباشد، آنگاه می‌توان به جای آزمون F از آزمون کروسکال والیس استفاده کرد. که البته محدودیت آن نسبت به آزمون F کمتر است.

معایب

کارایی این آزمون در حدود ۹۵ درصد آزمون f است و فرضیه‌های مخالف در این آزمون همواره بدون جهت است. یعنی نشان می‌دهد که، گروه‌ها با یکدیگر تفاوت دارند. اما هرگز بیان نمی‌کند که کدام گروه بالاتر و کدام گروه پایین‌تر است. (سیگل ترجمه کریمی، ۱۳۸۰: ۲۴۲)

مثال: پژوهشگری علاقمند است که به بررسی جو خانوادگی و تأثیر آن بر پیشرفت تحصیلی درس ریاضی دانش‌آموزان ابتدایی بپردازد. به همین منظور ۲۱ نفر از دانش‌آموزان در سه گروه، خانواده‌های دارای جو دموکراسی (X)، دیکتاتوری (Y) و آزادی مطلق (Z) انتخاب و آزمون ریاضی اجرا کرد که نتایج آن به شرح زیر است:

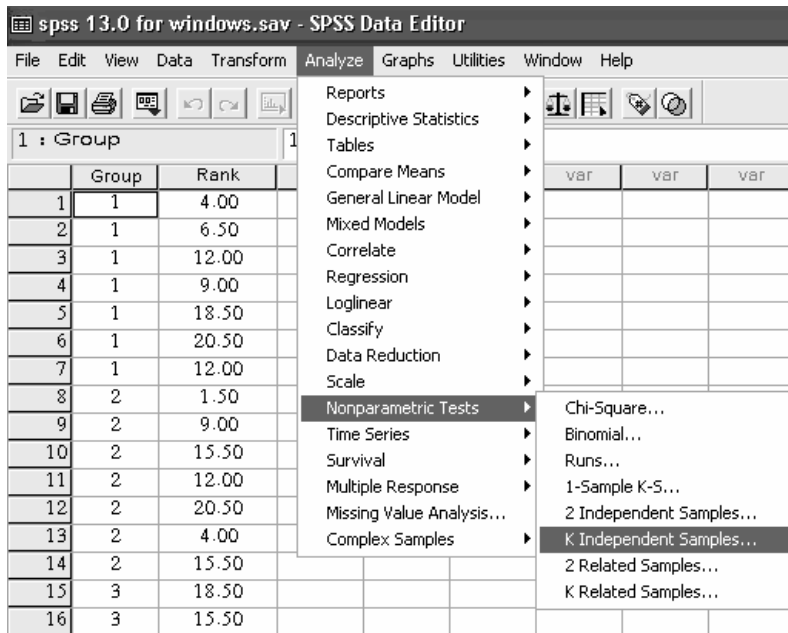
رتبه	Z	رتبه	Y	رتبه	X	دانش‌آموز
۱۸/۵	۱۸	۱/۵	۱۲	۴	۱۳	۱
۱۵/۵	۱۷	۹	۱۵	۶/۵	۱۴	۲
۴	۱۳	۱۵/۵	۱۷	۱۲	۱۶	۳
۹	۱۵	۱۲	۱۶	۹	۱۵	۴
۱۵/۵	۱۷	۲۰/۵	۱۹	۱۸/۵	۱۸	۵
۱/۵	۱۲	۴	۱۳	۲۰/۵	۱۹	۶
۶/۵	۱۴	۱۷	۱۷	۱۶	۱۶	۷

مراحل اجرای آزمون کروسکال والیس در Spss به شرح زیر می‌باشد:

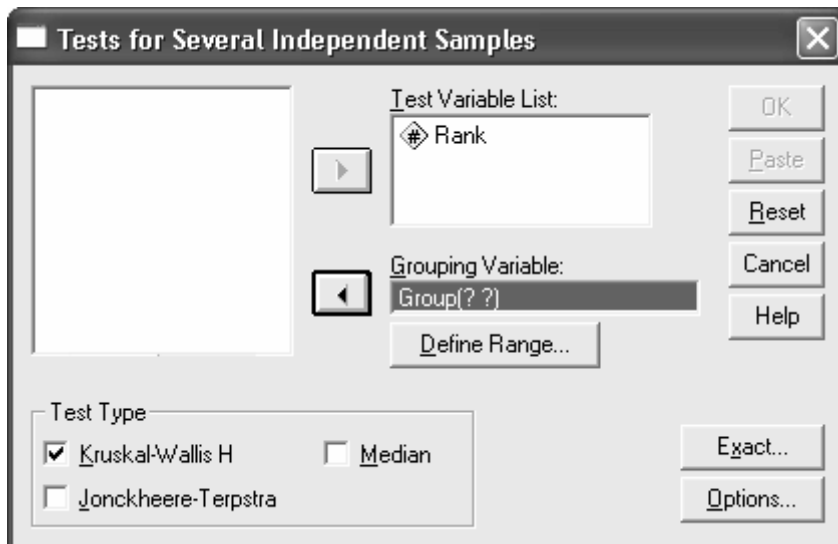
Analyze

Nonparametric test

K Independent Samples...

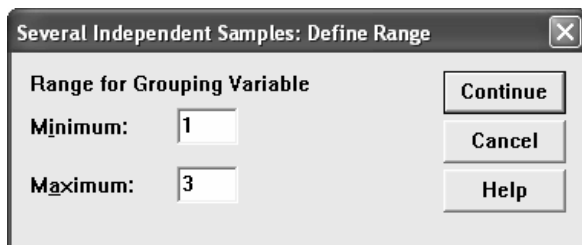


بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده متغیر از قسمت (Test Variable List)، رتبه (Rank) را انتخاب کرده و از قسمت (Grouping Variable) (Grouping Variable) را در نظر می‌گیریم

باید گروه را تعریف کنید. در این مثال، پژوهشگر چون به بررسی ۳ گروه پرداخته است در قسمت (Several Independent Samples: Define Range) در قسمت (Minimum) عدد یک و در قسمت (Maximum) عدد سه قرار دهید. سپس (Continue) را فعال نمائید:



Several Independent Samples: Define Range

Range for Grouping Variable

Minimum: 1

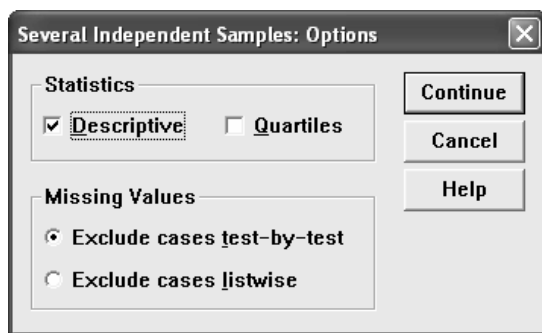
Maximum: 3

Continue

Cancel

Help

تا به پنجره گفتگوی (Test for Several Independent Samples) باز گردید و چنانچه بخواهید آمار توصیفی را محاسبه کنید، گزینه (Descriptive) را کلیک نمایید، تا پنجره گفتگوی (Option...) حاصل شود و گزینه (Descriptive) را انتخاب کنید:



Several Independent Samples: Options

Statistics

☒ Descriptive ☐ Quartiles

Missing Values

☒ Exclude cases test-by-test

☐ Exclude cases listwise

Continue

Cancel

Help

سپس گزینه (Continue) کلیک نمایید تا دوباره به پنجره گفتگوی (Test for Several Independent Samples) باز گردید در این پنجره گزینه (Kruskal-Wallis H) را انتخاب نمایید:



Test Type

☒ **Kruskal-Wallis H** ☐ **Median**

☐ **Jonckheere-Terpstra**

و سپس (OK) را کلیک نمایید تا نتایج به صورت خروجی زیر مشاهده گردد:

NPar Tests

Kruskal-Wallis Test

Ranks

	Group	N	Mean Rank
Rank	X	7	11.79
	Y	7	11.14
	Z	7	10.07
	Total	21	

Test Statistics^{a,b}

	Rank
Chi-Square	.277
df	2
Asymp. Sig.	.870

a. Kruskal Wallis Test

b. Grouping Variable: Group

در جدول مربوط به آزمون کروسکال والیس همانطور که مشاهده می‌کنید میانگین گروه‌ها به ترتیب از دانشکده اول تا سوم، ۱۱/۷۹ و ۱۱/۱۴ و ۱۰/۰۷ می‌باشد. همچنین مقدار به دست آمده در جدول که برابر ۰/۸۷۰ است که این مقدار از مقدار جدول کوچک‌تر است. بنابراین می‌توان نتیجه گرفت، نمرات دانش‌آموزان گروه‌های مختلف (دانش‌آموزان دارای جو خانوادگی دموکراسی، دانش‌آموزان دارای جو خانوادگی دیکتاتوری و دانش‌آموزان دارای جو خانوادگی آزادی مطلق) در درس ریاضی با یکدیگر تفاوت معنی‌داری ندارند.

مقایسه‌های بعد از تجربه زوجی (یک به یک)

همان‌طور که معنی‌دار بودن آزمون F در ANOVA یک متغیری، نشان می‌دهد که میان میانگین K جامعه تفاوت معنی‌دار وجود دارد، معنی‌دار بودن آزمون H در

KWANOVA نیز نشان می‌دهد که، تفاوت معنی‌داری بین K جامعه وجود دارد. اگر فرضیه صفر را رد کنیم، قدم بعدی آن است که تفاوت میان میانگین رتبه‌ها را به صورت آماری بررسی کنیم. در اینجا شیوه مقایسه بعد از تجربه زوجی ارائه می‌شود. روش همان روش HSD توکی است که قبلاً مورد بحث قرار گرفته است. در KWANOVA از آزمون توکی برای مقایسه میانگین رتبه‌ها استفاده می‌شود. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۲۴۹).

آزمون‌های ناپارامتریک مربوط به K گروه وابسته (K Related samples)

آزمون‌های ناپارامتریک مربوط به K گروه وابسته به مقایسه دو یا چند متغیر وابسته را می‌پردازد که از بین این آزمون‌ها، آزمون فریدمن مورد بررسی قرار می‌گیرد. آزمون فریدمن، یکی از آزمون‌های ناپارامتری است که در مقابل آزمون F بوجود آمده و مشابه آزمون کروسکال والیس است، با این شرط که k گروه نمونه جور شده باشند. این آزمون معلوم می‌کند که آیا میانگین رتبه‌ها یا حاصل جمع رتبه‌ها به طور معنی‌داری با یکدیگر تفاوت دارند یا خیر؟

$$T = \frac{(k-1) \left[\sum_{j=1}^k R_j^2 - n^2(k)(k+1)/\epsilon \right]}{\frac{nk(k+1)(2k+1)}{6} - \frac{nk(k+1)^2}{4}}$$

کاربرد

این آزمون زمانی بکار برده می‌شود که مقیاس اندازه‌گیری دست کم در سطح مقیاس ترتیبی باشد. بنابراین وقتی که داده‌ها به صورت رتبه‌ای باشند آزمون فریدمن بر آزمون f ترجیح داده می‌شود. همچنین وقتی این آزمون را به کار می‌بریم که حجم گروه‌ها تقریباً یکسان باشد.

محاسن

۱. توان آزمون فریدمن برابر با قوی‌ترین آزمون پارامتری برای K گروه مستقل (یعنی آزمون F) است.

۲. یکی از شرایط اصلی کاربرد آزمون f، همگنی واریانس است که در آزمون فریدمن نیاز به رعایت این شرط نیست.
۳. آزمون فریدمن را می‌توان با گروه‌هایی که تعداد آزمودنی‌ها خیلی کوچک است به‌کار برد. (سیگل، ترجمه کریمی، ۱۳۸۰: ۲۱۰)

معایب

مشکل آزمون فریدمن این است که هر یک از گروه‌ها حجم‌شان باید تقریباً یکی باشد، یعنی تعداد افراد و نمونه‌ها در گروه‌های مختلف با هم برابر باشند.

مثال: پژوهشگری علاقمند است تأثیر اوقات مختلف روز (ساعت ۸ صبح، ساعت ۱۰ صبح و ساعت ۲ بعد از ظهر) را در عملکرد گروهی دانش‌آموزان در یک مدرسه مورد مطالعه قرار دهد. به همین منظور ۱۶ آزمودنی را انتخاب می‌کند و عملکرد آنان را در اوقات تعیین شده اندازه‌گیری می‌کند با استفاده از آزمون فریدمن و با احتمال ۰/۰۵ خطا تعیین کنید که آیا اوقات مختلف روز در عملکرد آزمودنی‌ها مؤثر است یا خیر؟ (حسینی، ۱۳۸۲: ۱۵۸)

رتبه	۲ بعدظهر	رتبه	۱۰ صبح	رتبه	۸ صبح	دانش‌آموز
۳	۱۴	۲	۱۲	۱	۱۰	۱
۳	۱۷	۲	۱۶	۱	۱۵/۵	۲
۲	۱۲	۱	۱۰	۳	۱۴	۳
۳	۱۷	۲	۱۵	۱	۱۲	۴
۳	۱۵	۱	۱۰	۲	۱۱	۵
۱	۱۳	۲	۱۵	۳	۱۷	۶
۲	۱۸	۳	۱۹	۱	۱۸/۲۵	۷
۱	۱۴	۳	۱۷	۲	۱۷	۸
۱	۱۲	۲	۱۴	۳	۱۶	۹
۱	۹	۲	۱۰	۳	۱۱	۱۰
۳	۱۹	۲	۱۷	۱	۱۲	۱۱
۳	۱۴	۲	۱۴	۱	۱۳	۱۲
۳	۱۶	۲	۱۵/۵	۱	۱۵	۱۳
۳	۱۸	۲	۱۷	۱	۱۶	۱۴
۱	۱۱	۱	۱۴	۳	۱۶	۱۵
۲	۱۰/۵	۱	۱۰	۳	۱۶	۱۶

اکنون با توجه به نمرات می‌توانید نمره هر فرد را نسبت به خودش در هر وضعیت به دست آورید به عنوان مثال، فرد شماره یک در سه آزمون نمرات ۱۰-۱۲-۱۴ را کسب نموده، می‌توانید رتبه یک را به نمره ۱۰ رتبه دوم را به ۱۲ و رتبه سوم را به نمره ۱۴ اختصاص دهید و برای فرد دوم با توجه به نمرات رتبه‌های ۱-۲-۳ را اختصاص دهید.

با استفاده از آزمون مناسب بیان کنید که آیا تفاوت معنی‌داری میان گروه‌ها (افراد) از لحاظ انواع روش‌های تدریس بیان شده وجود دارد یا خیر؟
برای اجرای آزمون فوق در Spss مراحل زیر را انجام دهید

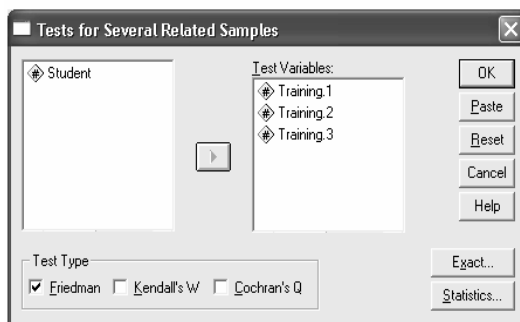
Analyze

Nonparametric test

K Related Samples...



بعد از اجرای این فرمان شکل زیر ظاهر می‌شود:



در پنجره ظاهر شده متغیرها را به قسمت (Test Variable) منتقل کنید و گزینه (Friedman) را انتخاب نمایید و سپس (OK) را کلیک نموده تا نتایج به صورت خروجی زیر آشکار شود:

NPar Tests
Friedman Test

Ranks

	Mean Rank
Training.1	1.88
Training.2	1.94
Training.3	2.19

Test Statistics^a

N	16
Chi-Square	.875
df	2
Asymp. Sig.	.646

a. Friedman Test

جدول بالا مربوط به آزمون فریدمن است، درجدول مذکور به بررسی سه گروه پرداخته شده است، میانگین رتبه‌های سه گروه عبارتند از: ۱/۱،۹۱/۸۸ و ۲/۲۲ همچنین χ^2 مربوط به سه گروه با درجه آزادی ۲ برابر با ۸۷۵/ است که این مقدار در سطح ۰/۰۵ معنی‌دار نیست پس تفاوت معنی‌داری وجود ندارد. و عملکرد دانش‌آموزان در هر سه موقع یکسان بوده است

خودآزمایی

۱. پژوهشگری نمونه‌ای از کارکنان یک منطقه آموزش و پرورش را انتخاب و آن‌ها را در سطح آمادگی، دبستان، راهنمایی و دبیرستان تقسیم بندی نمود. سپس نظر آنان را در مورد تشکیل انجمن معلمان سؤال و به صورت جدول زیر تنظیم کرد. با یک

آزمون آماری مناسب و در سطح $0/01$ درصد تعیین کنید بین فراوانی های جمع‌آوری شده اختلاف معنی‌داری وجود دارد یا خیر؟ (دلاور، ۱۳۸۰: ۵۰۱)

	آمادگی	دبستان	راهنمایی	دبیرستان
موافق	۶	۲۰	۲۵	۲۰
مخالف	۱۶	۱۵	۱۴	۱۰

۲. به منظور بررسی رابطه بین سواد زنان و نگرش آنان نسبت به تنبیه فرزندان، نظر زنان باسواد و بی‌سواد را درباره ضرورت تنبیه فرزندان مورد سؤال قرار داده‌ایم. جدول زیر نظر موافق و مخالف آن‌ها را درباره تنبیه نشان می‌دهد. با یک روش آماری مناسب و در سطح $0/05$ این فرضیه را که بین سواد و نگرش زنان درباره تنبیه رابطه وجود دارد، آزمون کنید.

	موافق	مخالف
باسواد	۱۵	۴۵
بی سواد	۳۵	۵

۳. از ۴۴ بیمار یک مرکز پزشکی که به‌طور تصادفی انتخاب شده بودند، پرسیده شد که از پزشک خود راضی هستند یا خیر. مؤسسه تحقیقاتی علاقه‌مند است بداند که آیا نظر بیماران در مورد پزشک معالج با شدت بیماری آن‌ها ارتباط دارد یا خیر؟ اطلاعات به‌دست آمده در جدول زیر ارائه شده است. (شیولسون ترجمه کیامنش، ۱۳۸۲: ۳۹۸).

	راضی	ناراضی	جمع
بیماری سخت	۸	۱۴	۲۲
بیماری خفیف	۱۲	۱۰	۲۲
جمع	۲۰	۲۰	۴۴

آزمون مجذور خی را در سطح $0/05$ انجام دهید.
فرضیه‌های صفر و خلاف را بنویسید.
اکنون نشان دهید که فرض‌ها برقرار هستند و شرایط لازم برای اجرای آزمون موجود است.

۴. یک گروه ۱۴ نفری از دانشجویان سال اول دانشکده به طریقی تصادفی به دو گروه گواه و آزمایشی تقسیم شدند. هدف این مطالعه، آن بود که تأثیر برنامه های مشاوره ای در مورد دانشجویان گروه آزمایشی در معدل آنها مورد بررسی قرار گیرد با استفاده از آزمون U مان- ویتنی فرضیه صفر را که بین پیشرفت تحصیلی دو گروه آزمایشی و گواه تفاوت معنی داری وجود ندارد را آزمون کنید.

گروه کنترل	گروه آزمایش
۱۷	۱۴
۱۸	۱۷
۱۴	۱۶
۱۵	۱۳
۱۴	۱۷
۱۳	۱۸
۱۱	۱۶

۵. محقق می خواهد تفاوت تأثیر سه روش درمانی A,B,C را در کاهش میزان افسردگی سه گروه از بیماران افسرده که به طور تصادفی انتخاب کرده است، مورد بررسی قرار دهند. فرض کنید نمره های افراد مورد مطالعه در پایان دوره درمان به شرح زیر تعیین شده باشد. با استفاده از آزمون کروسکال - والیس فرضیه صفر را که بین روش های درمانی A,B,C در کاهش میزان افسردگی بیماران مؤثر نیست را آزمون کنید:

روش درمانی A	روش درمانی B	روش درمانی C
۱۰	۲۵	۳
۱۵	۱۷	۷
۲۰	۳۵	۱۰
۱۱	۴۰	۱۴
۱۷	۲۰	۱۵
	۱۰	۲۰
		۲۵

۶. پژوهشگری علاقمند است که به بررسی تأثیر روش تدریس حل مسأله بر پیشرفت تحصیلی دانش آموزان بپردازد به همین منظور پیش از شروع آموزش، یک پیش آزمون به عمل می آورد و سپس متغیر مستقل (روش تدریس حل مسأله ارائه کرده

است) از دانش‌آموزان پس از آزمون به عمل می‌آورد و با استفاده از آزمون ویلکاکسون، فرضیه صفر را در سطح ۰/۰۱ تفاوت، آزمون کنید.

پس آزمون	پیش آزمون
۱۴	۱۵
۱۷	۱۶
۱۷	۱۶/۵
۱۱	۱۰
۱۱	۱۰/۵
۱۳	۱۷
۱۲	۱۷/۲۵
۱۵	۱۱
۱۴	۱۲
۱۴	۱۳
۱۵	۱۳/۲۵
۱۴	۱۴
۱۷	۱۶

۷. فرض نمایید که نمره امتحان یک درس با سه روش تدریس (بحث گروهی، حل مساله، اکتشافی) ارزیابی شود. محقق در این زمینه بر روی یک گروه، سه روش تدریس را اجرا و سپس نتایج امتحان در هر یک از روش‌ها را یادداشت می‌کند. نتایج امتحان این ۱۶ نفر در سه آزمون در جدول زیر آورده شده است، اکنون با استفاده از آزمون فریدمن فرضیه صفر را آزمون کنید. (حسینی، ۱۳۸۲: ۱۵۷)

نمره امتحان پس از روش تدریس بحث گروهی	نمره امتحان پس از روش تدریس حل مساله	نمره امتحان پس از روش تدریس اکتشافی
۱۰	۱۲	۱۴
۱۵/۵	۱۶	۱۷
۱۴	۱۰	۱۲
۱۲	۱۵	۱۷
۱۱	۱۰	۱۵
۱۷	۱۵	۱۳
۱۸/۲۵	۱۹	۱۸/۵
۱۷	۱۷/۵	۱۴
۱۶	۱۴	۱۲
۱۱	۱۰	۹
۱۲	۱۷	۱۹
۱۳/۵	۱۴	۱۴/۵
۱۵	۱۵/۵	۱۶
۱۶	۱۷	۱۸
۱۶	۱۴	۱۱
۱۶	۱۰	۱۰/۵

فصل نهم

آزمون‌های همبستگی

هدف‌های یادگیری

از دانشجو انتظار می‌رود که پس از خواندن فصل نهم بتواند:

۱. مفهوم همبستگی را بیان کند.
۲. نمودار پراکنش را توضیح دهد.
۳. تفاوت بین همبستگی مثبت و منفی را توصیف کند.
۴. ضریب همبستگی اسپیرمن را محاسبه نماید.
۵. ضریب همبستگی پیرسون را محاسبه نماید.
۶. ضریب همبستگی جزئی را محاسبه نماید.

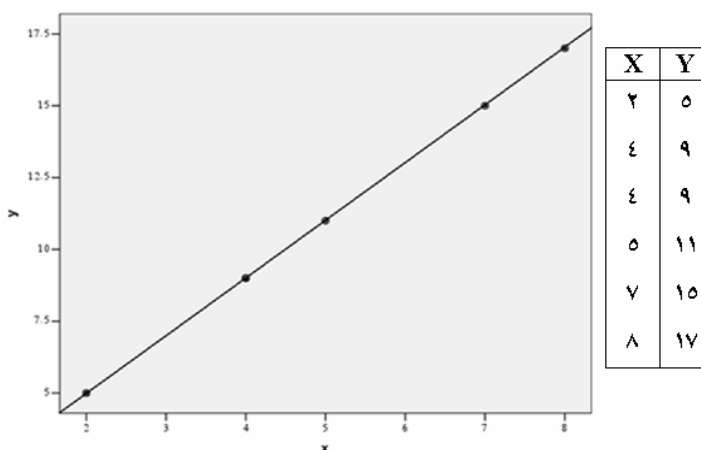
مفهوم همبستگی

ارتباط بین دو متغیر را همبستگی می‌نامند و این نوع ارتباط مورد توجه خاص محققین و پژوهشگران است. بین دو متغیر ممکن است همبستگی مستقیم مثبت یا همبستگی مستقیم منفی و یا اصلاً همبستگی وجود نداشته باشد. دو متغیر زمانی دارای همبستگی مثبت است که وقتی که مقدار یکی از متغیرها ازدیاد پیدا کند، دیگری هم روندی در جهت ازدیاد به‌طور یکنواخت و متشابه با تغییرات متغیر اول را نشان می‌دهد. دو متغیر را در حالت کلی غیر همبسته می‌گویند که وقتی یکی از آن‌ها تغییر کند در دیگری هیچ‌گونه تغییری مشاهده نگردد و بالاخره دو متغیر را همبسته منفی می‌گویند که وقتی یکی از آن‌ها تغییر کند مثلاً ازدیاد پیدا کند تغییر متغیر دیگری در جهت عکس تغییر اولی باشد، مثلاً تنزیل یابد.

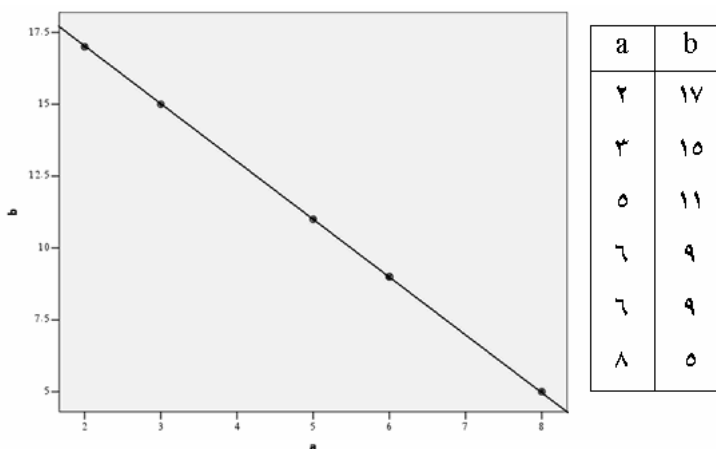
اندازه همبستگی بین متغیرها ضریب همبستگی نامیده می‌شود که معمولاً از صفر تا +۱ و از صفر تا -۱ تغییر می‌کند. ضریب همبستگی +۱ را همبستگی مثبت کامل و ضریب همبستگی -۱ را همبستگی منفی کامل می‌نامند. ضریب همبستگی صفر نشانگر آن است که بین دو متغیر هیچ نوع همبستگی وجود ندارد. البته باید دانست که به ندرت ممکن است بین دو دسته متغیر در علوم رفتاری همبستگی کامل موجود باشد. در شکل‌های زیر نمایش ترسیمی همبستگی‌ها در حالت‌های مختلف نشان داده شده است

(پاشا شریفی و نجفی زند، ۱۳۷۵: ۱۱۴)

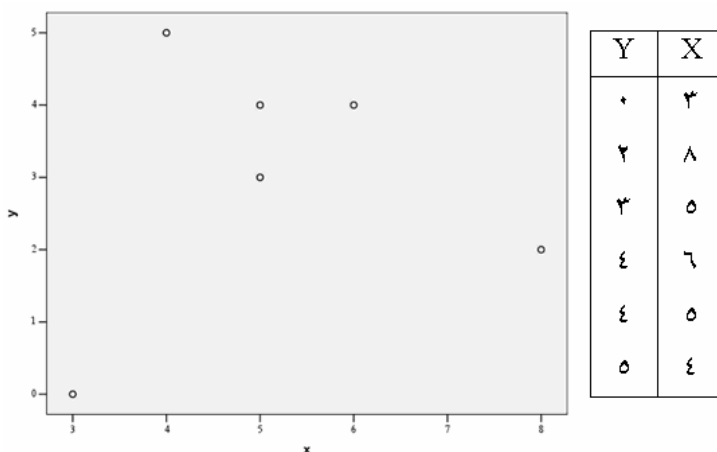
الف) همبستگی +۱ (همبستگی مثبت کامل)



ب) همبستگی -۱ (همبستگی منفی کامل)



ج) همبستگی صفر (بین دو متغیر هیچ نوع همبستگی وجود ندارد)



نمودار پراکنش

برای تعیین نوع رابطه میان دو متغیر ابتدا باید میزان وابستگی آن‌ها را اندازه گرفت یکی از روش‌های بررسی این رابطه، رسم نمودار تغییرات این دو متغیر در ارتباط با یکدیگر است، که آن را نمودار پراکنش می‌نامند. (دلاور، ۱۳۸۰: ۲۷۷) این نمودار یک نمایش ترسیمی است که از طریق آن ارزش‌های دو متغیر نشان داده می‌شود.

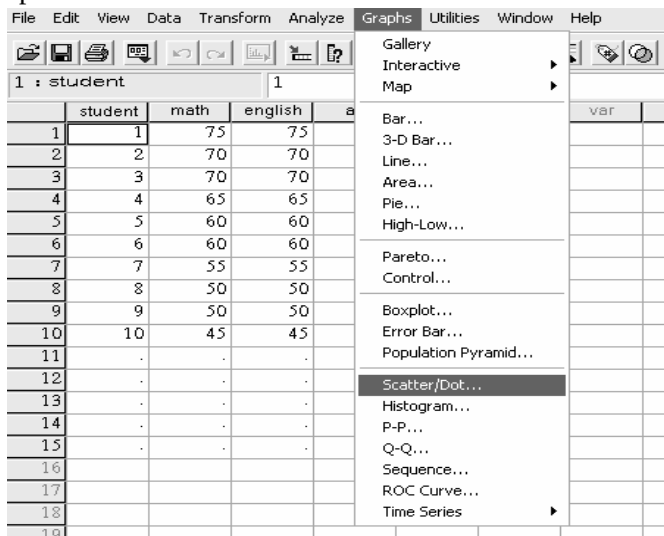
مثال: در جدول زیر نتایج آزمون ۱۰ نفر دانش‌آموز در چهار درس نشان داده شده است:

دانش‌آموزان	ریاضی	هنر	انگلیسی	شیمی
۱	۷۵	۴۵	۷۵	۷۱
۲	۷۰	۵۰	۷۰	۴۰
۳	۷۰	۵۰	۷۰	۵۶
۴	۶۵	۵۵	۶۵	۵۰
۵	۶۰	۶۰	۶۰	۶۰
۶	۶۰	۶۰	۶۰	۷۰
۷	۵۵	۶۵	۵۵	۷۰
۸	۵۰	۷۰	۵۰	۵۰
۹	۵۰	۷۰	۵۰	۶۵
۱۰	۴۵	۷۵	۴۵	۵۱

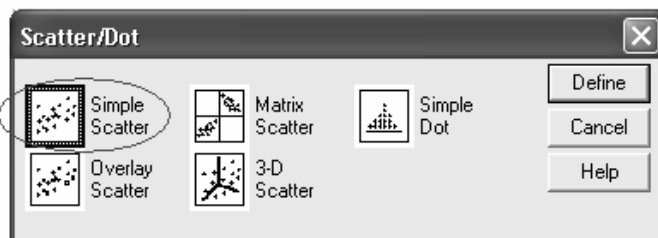
همان‌طور که ملاحظه می‌شود نمره‌های هر دانش‌آموز در آزمون‌های ریاضی و انگلیسی یکسان است. بدست آوردن چنین نمره‌هایی در عمل، غیر ممکن است در اینجا صرفاً به خاطر توضیح همبستگی از این نمرات استفاده شده است. (دلاور، ۱۳۸۰: ۲۷۷)

برای رسم نمودار پراکنش مراحل زیر را انجام دهید:

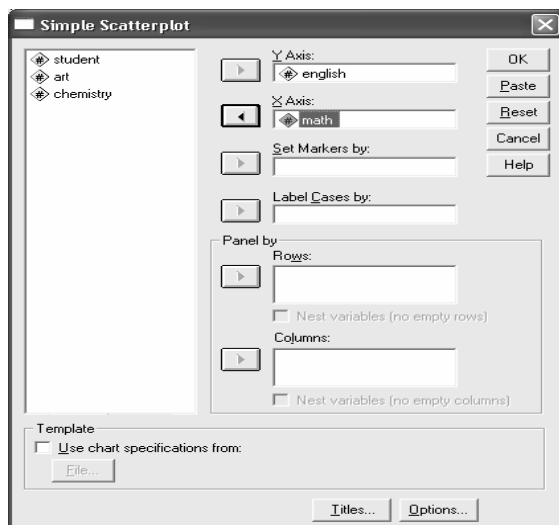
Graph Scatter plot...



با اجرای فرمان فوق پنجره گفتگوی زیر ظاهر می‌شود:

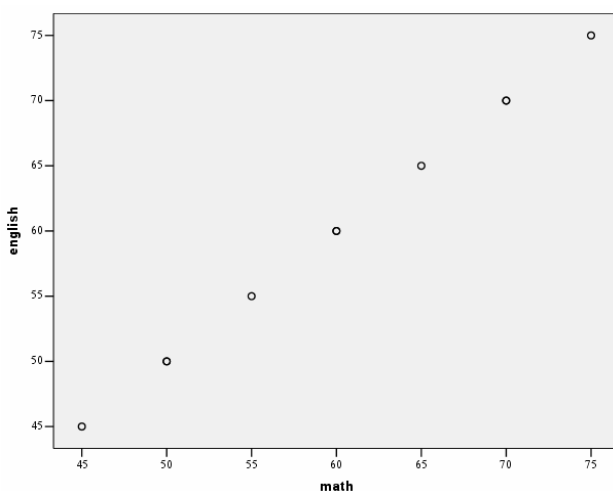


در پنجره (Scatter plot) گزینه (Simple) نمودار پراکنش ساده دو متغیر را نشان می‌دهد که در تحلیل همبستگی مورد نیاز است. با کلیک بر روی گزینه (Simple) برای معرفی متغیرهایی که مقادیرشان در نمودار رسم می‌شود کلید (Define) را فعال کنید تا پنجره گفتگوی زیر (Simple Scatterplot) ظاهر شود:

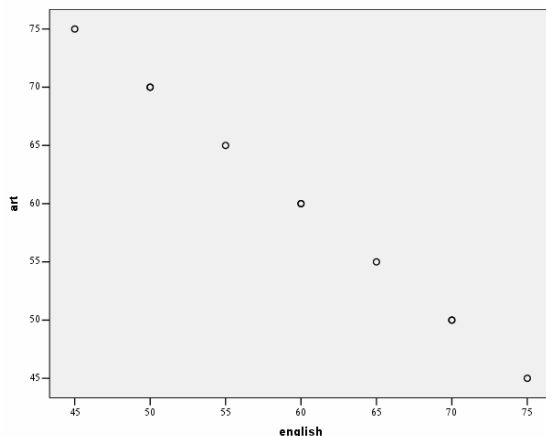


در این پنجره جعبه (Y Axis) متغیر محور عمودی و در جعبه (X Axis) متغیر محور افقی قرار می‌گیرد اگر بخواهیم نمودار بر حسب رده‌های مختلف متغیر دیگری علامت گذاری شود متغیر مورد نظر را به جعبه (Set Markers by) منتقل می‌کنیم و در صورتی که بخواهیم مقادیر زوج متغیرها بر حسب خاصی داشته باشند، متغیری که برچسب‌ها در آن جای دارد به جعبه (Label Cases by) منتقل می‌کنیم با اجرای دستورات فوق نمودار پراکنش زیر ظاهر می‌شود.

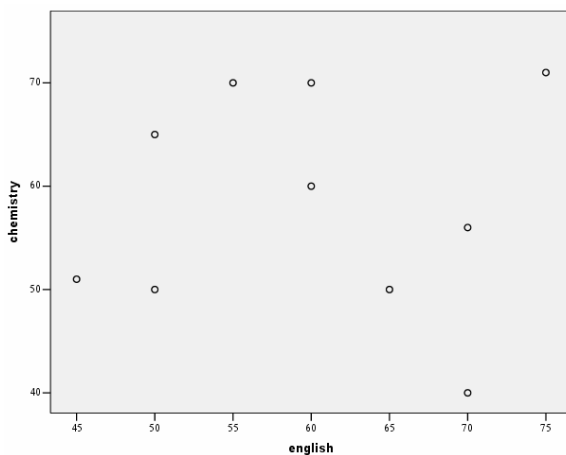
نمودار پراکنندگی نمرات ریاضی و نمرات انگلیسی همبستگی بین آن‌ها را نشان می‌دهد که این همبستگی مثبت مستقیم و کامل می‌باشد.



نموار پراکندگی زیر همبستگی بین دو متغیر انگلیسی و هنر را نشان می‌دهد. این نمودار برعکس نمودار قبل است. همانطور که ملاحظه می‌شود افزایش نمره در یک درس همراه با کاهش نمره در درس دیگر است این نمودار همبستگی منفی و معکوس (بین انگلیسی و هنر) را نشان می‌دهد:



و نمودار پراکندگی بین نمرات انگلیسی و شیمی نشان می‌دهد که رابطه روشنی بین نمره‌های شیمی و انگلیسی ملاحظه نمی‌شود. به این معنی که دانش‌آموزانی که در انگلیسی نمره خوب گرفته‌اند نمره آن‌ها در شیمی هم خوب و هم بد شده است. به عبارت دیگر نمی‌توان از طریق نمره انگلیسی نمره شیمی را پیش‌بینی کرد. نمودار مورد بحث همبستگی بین این دو متغیر را صفر یا نزدیک به صفر نشان می‌دهد. بنابراین می‌توان نتیجه گرفت که دو متغیر مستقل هستند.



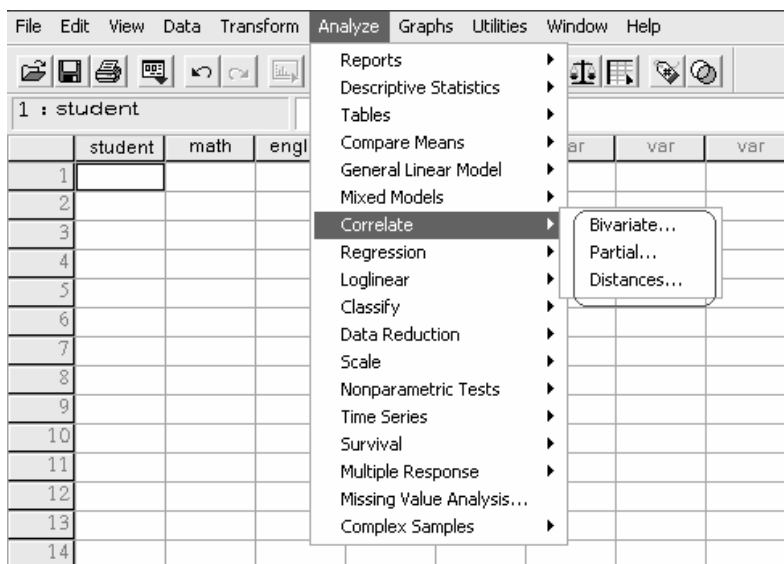
محاسبه ضریب همبستگی

برای محاسبه ضریب همبستگی بین دو یا چند متغیر مراحل زیر را انجام دهید:

Analyze

Correlate

Bivariate.../Partial.../Distances...



فرمان (Bivariate) برای محاسبه ضریب همبستگی خطی دو متغیر، فرمان (Partial) برای محاسبه ضریب همبستگی جزئی زوج متغیرها به فرض ثابت بودن مقادیر سایر متغیرها و فرمان (Distances) برای محاسبه ضرائب همبستگی متغیرهایی که در مقیاس فاصله‌ای تغییر می‌کنند، به کار می‌روند.

ضریب همبستگی پیرسون

این آزمون یکی از متداول‌ترین آزمون‌های تعیین ضریب همبستگی بین متغیرهای دارای اندازه‌های فاصله‌ای و نسبی است (پاشا شریفی و نجفی زند، ۱۳۷۵: ۱۱۸) و برای محاسبه ضرایب آن از فرمول زیر استفاده می‌شود:

$$r_{xy} = \frac{N \sum XY - (\sum X)(\sum Y)}{\sqrt{[N \sum X^2 - (\sum X)^2][N \sum Y^2 - (\sum Y)^2]}}$$

پس از محاسبه ضریب همبستگی، محقق معنی‌دار بودن یا نبودن آن را مورد بررسی قرار می‌دهد. برای این کار ابتدا لازم است درجه آزادی محاسبه شود سپس با در نظر گرفتن درجه آزادی و سطح احتمال مورد نظر (۰/۰۱ یا ۰/۰۵ خطا) محقق ضریب همبستگی محاسبه شده مساوی یا بزرگ‌تر از عدد جدول باشد، می‌توان وجود ضریب همبستگی را پذیرد.

مثال: فرض کنید دو آزمون متفاوت، یکی ریاضی و دیگری آزمون هوش روی گروهی از دانش‌آموزان اجرا شده و نمره‌های به‌دست آمده به شرح زیر است:

دانش‌آموزان	هوش بهر	آزمون ریاضی
۱	۱۲۸	۱۹
۲	۱۲۴	۱۳
۳	۱۱۰	۱۷
۴	۱۱۹	۱۶
۵	۱۰۵	۱۲
۶	۱۲۴	۱۸
۷	۱۱۳	۱۶
۸	۱۱۳	۱۵
۹	۱۰۳	۱۱
۱۰	۱۲۸	۲۰

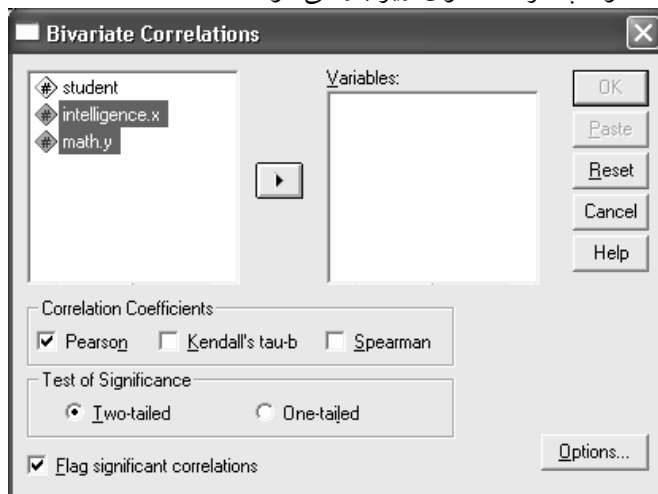
برای محاسبه ضریب همبستگی پیرسون مراحل زیر را انجام دهید:

Analyze

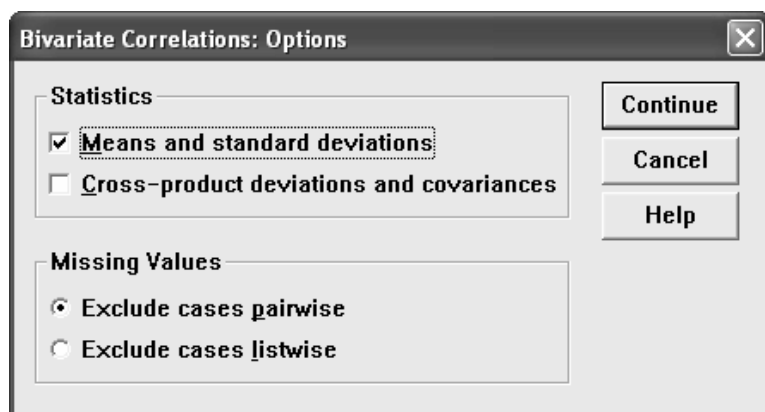
Correlate

Bivariate...

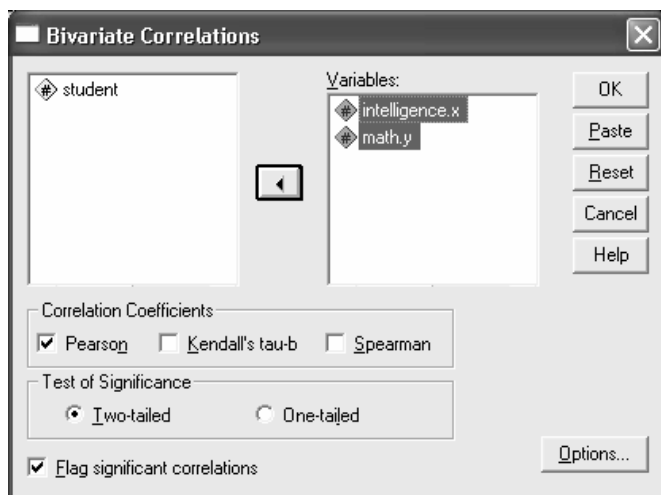
با اجرای فرمان فوق پنجره گفتگوی زیر باز می‌شود:



در تصویر فوق مجموعه متغیرهایی که ضریب همبستگی خطی آن‌ها محاسبه خواهد شد به جعبه (Variables) منتقل می‌کنیم سپس در قسمت (Correlation Coefficients) گزینه (Pearson) را انتخاب و در قسمت (Test of Significance) نوع آزمون معنی‌داری ضریب همبستگی تعیین کنید گزینه (Two-tailed) به دو سویه بودن آزمون و گزینه (One-tailed) به یک‌سویه بودن آزمون اشاره می‌کند همچنین با فعال کردن کلید (Options...) پنجره گفتگوی زیر باز می‌شود:



در پنجره فوق بخش (Statistics) گزینه (Means and standard deviations) برای نمایش میانگین‌ها و انحراف معیارها و بخش (Missing Values) نحوه برخورد با مشاهدات گمشده را نشان می‌دهد.



سپس در پنجره (Bivariate Correlations) گزینه (OK) را کلیک کنید و نتایج به صورت خروجی زیر ظاهر خواهد شد:

Correlations

Descriptive Statistics

	Mean	Std. Deviation	N
intelligence.x	116.70	9.214	10
math.y	15.70	2.983	10

Correlations

		intelligence.x	math.y
intelligence.x	Pearson Correlation	1	.744*
	Sig. (2-tailed)		.014
	N	10	10
math.y	Pearson Correlation	.744*	1
	Sig. (2-tailed)	.014	
	N	10	10

*. Correlation is significant at the 0.05 level (2-tailed).

چنانکه مشاهده می کنید میزان همبستگی متغیر هوشبهر و نمره ریاضی ۰/۷۴ است که این مقدار در سطح ۰/۰۵ معنی دار است.

ضریب همبستگی اسپیرمن

این آزمون زمانی به کار می رود که داده ها از نوع رتبه ای است و اندازه های متغیرها به صورت رتبه ای تنظیم شده است برای محاسبه ضریب همبستگی از فرمول زیر استفاده می شود. (دلاور، ۳۰۶:۱۳۸۰)

$$\rho = 1 - \frac{6(\sum d_i^2)}{n(n^2 - 1)}$$

مثال: فرض کنید، می‌خواهیم ضریب همبستگی اسپیرمن بین دو دسته نمره آزمون ریاضی و دیگری آزمون هوش محاسبه کنیم. بنابراین لازم است که قبل از محاسبه، مقیاس نمره‌ها را از فاصله‌ای به رتبه‌ای تبدیل کرد:

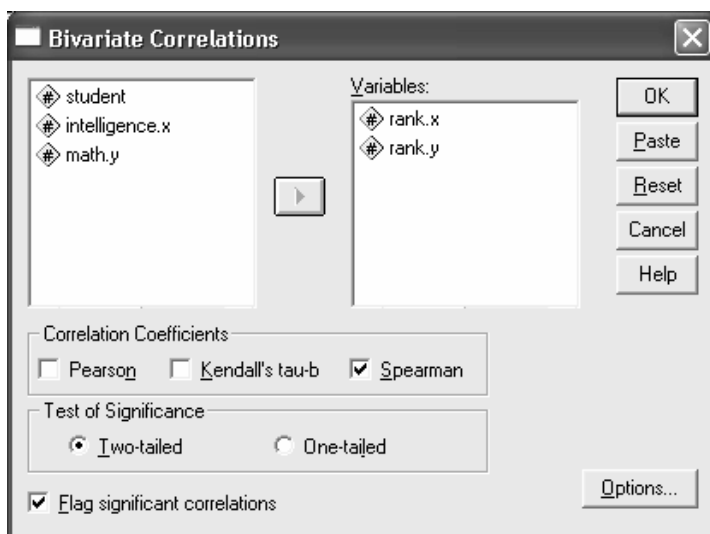
رتبه (Y)	رتبه (X)	نمره ریاضی (Y)	هوشبر (X)	دانش‌آموز
۲	۱/۵	۱۹	۱۲۸	۱
۸	۴	۱۳	۱۲۴	۲
۴	۸	۱۷	۱۱۰	۳
۵/۵	۵	۱۶	۱۱۹	۴
۹	۱۰	۱۲	۱۰۵	۵
۳	۳	۱۸	۱۲۴	۶
۵/۵	۶/۵	۱۶	۱۱۳	۷
۷	۶/۵	۱۵	۱۱۳	۸
۱۰	۹	۱۱	۱۰۳	۹
۱	۱/۵	۲۰	۱۲۸	۱۰

برای محاسبه ضریب همبستگی اسپیرمن مراحل زیر را انجام دهید:

Analyze

Correlate
Spearman

با اجرای فرمان فوق پنجره گفتگوی زیر باز می‌شود:



در تصویر فوق مجموعه متغیرهایی که ضریب همبستگی خطی آن‌ها محاسبه خواهد شد به جعبه (Variables) منتقل کنید سپس در قسمت (Correlation Coefficients) گزینه (Spearman) را انتخاب کنید و در قسمت (Test of Significance) نوع آزمون معنی‌داری ضریب همبستگی تعیین می‌شود گزینه (Two-tailed) به دو سویه بودن آزمون و گزینه (One-tailed) به یکسویه بودن آزمون اشاره می‌کند:

سپس در پنجره (Bivariate Correlations) گزینه (OK) را کلیک کنید و نتایج به صورت خروجی زیر ظاهر خواهد شد:

Nonparametric Correlations

Correlations

			rank.x	rank.y
Spearman's rho	rank.x	Correlation Coefficient	1.000	.780**
		Sig. (2-tailed)	.	.008
		N	10	10
	rank.y	Correlation Coefficient	.780**	1.000
		Sig. (2-tailed)	.008	.
		N	10	10

**. Correlation is significant at the 0.01 level (2-tailed).

چنانکه مشاهده می‌کنید، میزان همبستگی متغیر ریاضی با هوشبهر دانش‌آموزان برابر ۰٫۷۸۰ است که این مقدار در سطح ۰٫۰۱ معنی‌دار است و با ۹۹ درصد اطمینان می‌توان گفت که بین نمره ریاضی و هوشبهر دانش‌آموزان ارتباط وجود دارد.

همبستگی جزئی

اگر بخواهیم ضریب همبستگی بین دو متغیر را به شرط ثابت بودن سایر متغیرها محاسبه کنیم، ضریب همبستگی حاصل را جزئی می‌نامند. مثلاً اگر همبستگی بین متغیرهای X_1 و X_2 به شرط ثابت بودن متغیر X_3 را اندازه بگیریم این همبستگی را همبستگی جزئی می‌نامند.

مثال: فرض کنید در پژوهشی، محقق می‌خواهد همبستگی هوش دانش‌آموزان و وضعیت اقتصادی (حقوق) والدین را با پیشرفت تحصیلی درس ریاضی دانش‌آموزان محاسبه کند و نمره‌های به دست آمده به شرح زیر است:

دانش‌آموز	هوشبهر	درآمد والدین	نمره ریاضی
۱	۱۲۸	۲۷۵/۳۰۰	۱۹
۲	۱۲۴	۲۵۲/۴۰۰	۱۳
۳	۱۱۰	۳۵۶/۱۰۰	۱۷
۴	۱۱۹	۳۱۵/۰۴۲	۱۶
۵	۱۰۵	۱۹۷/۰۰۰	۱۲
۶	۱۲۴	۲۵۳/۳۰۰	۱۸
۷	۱۱۳	۴۰۰/۰۰۰	۱۶
۸	۱۱۳	۲۹۷/۵۰۰	۱۵
۹	۱۰۳	۲۱۰/۰۰۰	۱۱
۱۰	۱۲۸	۳۸۰/۰۰۰	۲۰

که پس از محاسبه نتیجه زیر به دست آمد:

Correlations

Correlations		math
intelligence	Pearson Correlation	.744*
	Sig. (2-tailed)	.014
	N	10
salary	Pearson Correlation	.644*
	Sig. (2-tailed)	.045
	N	10

*. Correlation is significant at the 0.05 level (2-tailed).

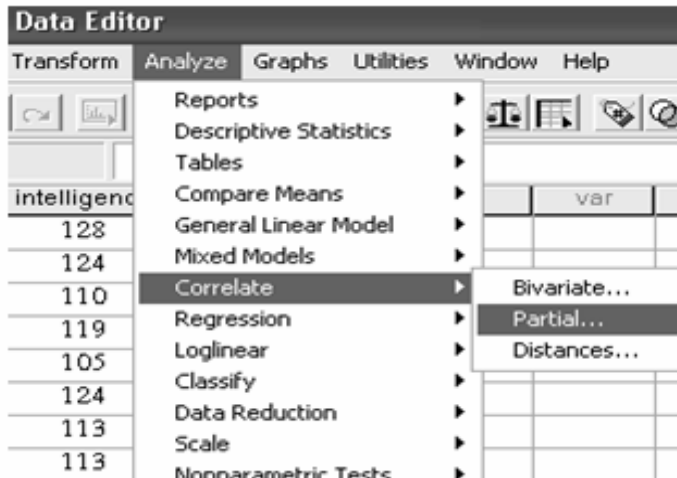
۱. میزان همبستگی متغیر میزان حقوق والدین (وضعیت اقتصادی) با پیشرفت تحصیلی درس ریاضی ۰/۶۴۴ است.
۲. میزان همبستگی متغیر هوش دانش‌آموزان با پیشرفت تحصیلی درس ریاضی ۰/۷۴۴ است.

در صورتی که پژوهشگر بخواهد متغیر هوش دانش‌آموزان را کنترل کند در این صورت پژوهشگر از روش همبستگی جزئی استفاده می‌کند برای محاسبه ضریب همبستگی جزئی مراحل زیر را انجام دهید:

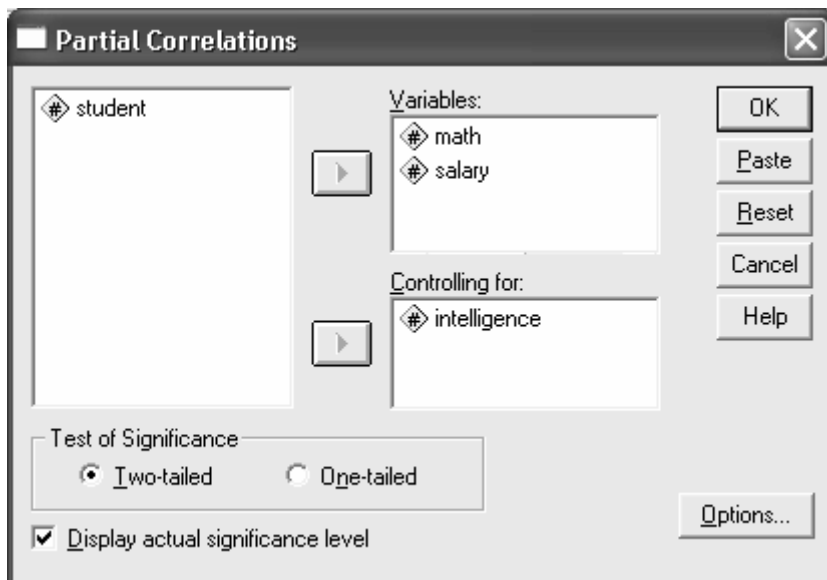
Analyze

Correlate

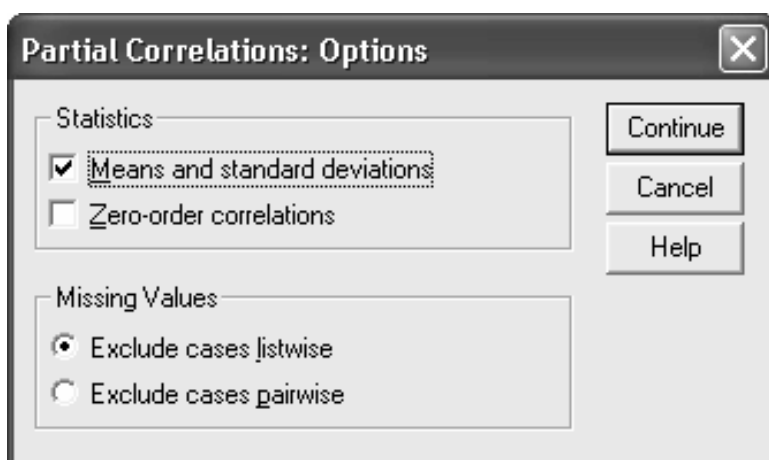
Partial...



با اجرای فرمان فوق پنجره زیر باز می‌شود:



در پنجره گفتگوی (Partial Correlations) متغیرهایی که محاسبه ضریب همبستگی جزئی آن‌ها مورد نظر است به قسمت (Variables) و متغیرهایی که در این محاسبه ثابت فرض می‌شوند به قسمت (Controlling for) منتقل کنید و در سایر گزینه‌ها مانند (Test of significance) و (Options...) مانند محاسبه ضریب همبستگی خطی عمل کنید با این تفاوت که در قسمت (Options...) در قسمت (Statistics) به جای گزینه (Cross-product deviations and covariance) گزینه مشابه (Zero-order correlations) جایگزین شده است:



با انتخاب این گزینه ضرائب همبستگی کلیه متغیرها بدون اعمال متغیرهای کنترلی (ثابت) محاسبه و آزمون می‌شوند. سپس در پنجره (Partial Correlations) گزینه (OK) را کلیک کنید و نتایج به صورت خروجی زیر ظاهر خواهد شد:

Partial Corr

Descriptive Statistics

	Mean	Std. Deviation	N
math	15.70000	2.983287	10
salary	293.6642	69.237143	10
intelligence	116.7000	9.214120	10

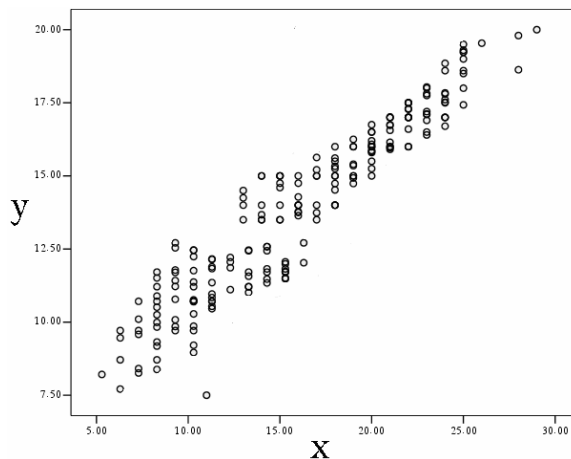
Correlations

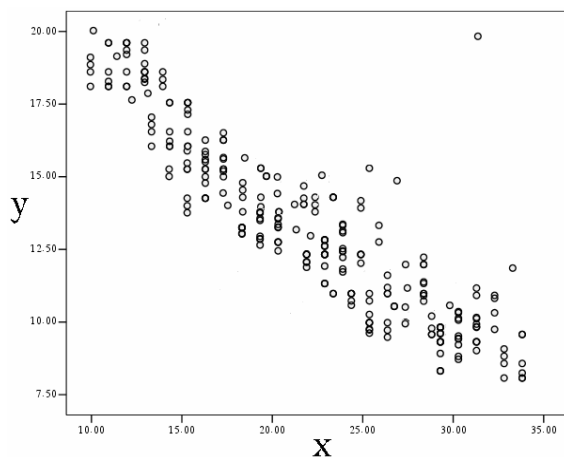
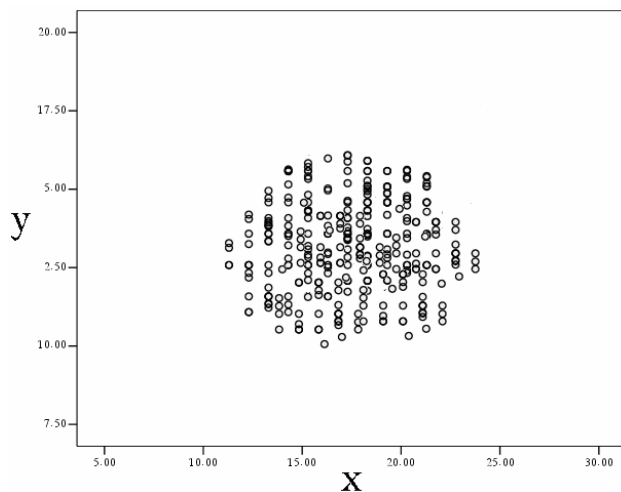
Control Variables			math	salary
intelligence	math	Correlation	1.000	.668
		Significance (2-tailed)	.	.049
		df	0	7
	salary	Correlation	.668	1.000
		Significance (2-tailed)	.049	.
		df	7	0

چنانکه مشاهده می‌کنید میزان همبستگی متغیر میزان حقوق والدین (وضعیت اقتصادی) با پیشرفت تحصیلی درس ریاضی در صورتی که متغیر هوش کنترل شود ۰٫۶۶۸ است.

خوآزمایی

۱. همبستگی چیست؟
۲. چه موقع همبستگی بین دو متغیر مثبت یا منفی است؟
۳. زیر را تفسیر کنید:
 - $r_{xy} = -۰٫۶۹$
 - $r_{xy} = -۰٫۷۱$
 - $r_{xy} = +۰٫۳۳$
۴. همبستگی‌های ترسیم شده در نمودارهای زیر را تعیین کنید:





۵. نمره‌های جدول زیر نتیجه اجرای دو آزمون روان‌شناسی عمومی و روان‌شناسی شخصیت است، برای اطلاعات مذکور نمودار پراکندگی رسم کنید:

روان‌شناسی عمومی	روان‌شناسی شخصیت
۱۴/۲۵	۱۹
۱۷	۱۸/۲۵
۱۴	۲۰
۱۵/۵	۱۴/۷۵
۱۸/۷۵	۱۹
۲۰	۱۶/۲۵
۱۷	۱۹
۱۴/۷۵	۱۳/۵

۶. نمره‌های ۱۰ دانشجو در درس روش تحقیق و آمار آنان به شرح زیر در دست است. ضریب همبستگی پیرسون را حساب کنید:

روش تحقیق	آمار
۱۵	۱۴
۲۰	۱۸
۱۹	۱۷
۱۴	۱۵
۹	۱۲
۱۱	۱۰
۱۵	۱۳
۱۰	۹
۱۸	۱۷
۹	۸

۷. برای دو دسته نمره زیر محاسبات خواسته شده را انجام دهید:

X	Y
۸	۶
۲۰	۱۴
۲۶	۱۰
۱۷	۸
۱۴	۲

- ضریب همبستگی پیرسون را حساب کنید.
- با تبدیل نمرات به رتبه ضریب همبستگی اسپیرمن را حساب کنید.
- ۸. برای نمره‌های زیر با کنترل متغیر X_r ، ضریب همبستگی X_1 را با Y محاسبه کنید:

X_1	X_r	Y
۴۵	۷۱	۷۵
۵۰	۴۰	۷۰
۵۰	۵۶	۷۰
۵۵	۵۰	۶۵
۶۰	۶۰	۶۰
۶۰	۷۰	۶۰
۶۵	۷۰	۵۵
۷۰	۵۰	۶۵
۷۰	۶۵	۵۰
۷۵	۵۱	۴۵

مراجع

۱. پاشا شریفی، حسن - نجفی زند، جعفر (۱۳۷۵). روش‌های آماری در علوم رفتاری تهران، نشر دانا.
۲. حسینی، سید یعقوب (۱۳۸۲). آمار ناپارامتریک، تهران، دانشگاه علامه طباطبائی.
۳. حافظ‌نیا، محمدرضا (۱۳۸۰). مقدمه ای بر روش تحقیق در علوم انسانی، تهران، انتشارات سمت.
۴. دلاور، علی (۱۳۸۰). روش‌های آماری در روان‌شناسی و علوم تربیتی، تهران، دانشگاه پیام‌نور.
۵. سیگل، سیدنی. آمار غیرپارامتریک برای علوم رفتاری، ترجمه یوسف کریمی (۱۳۸۰). تهران: دانشگاه علامه طباطبائی.
۶. سیف، علی اکبر (۱۳۷۶). روش‌های اندازه‌گیری و ارزشیابی آموزشی، تهران، انتشارات آگاه.
۷. شرکت آمار پردازان. راهنمای کاربران (SPSS 6.0)، تهران، چاپخانه نیل، ۱۳۷۹.
۸. شیولسون، ریچارد. ج. استدلال آماری در علوم رفتاری، ترجمه علیرضا کیامنش (۱۳۸۲)، تهران، انتشارات جهاد دانشگاهی. ۲. جلد
۹. کرلینجر، فرد. ان. مبانی پژوهش در علوم رفتاری، ترجمه حسن پاشا شریفی و جعفر نجفی زند (۱۳۷۶)، تهران، آوای نور.

10. <http://www.spss.com/corpinfo/history.2007.htm>

معرفی سایت‌های اینترنتی برای آشنایی بیشتر:

www.spss.com
<http://employees.csbsju.edu/rwiellk/psy347/spssinst.htm>
<http://www.statistica.com.au/htm/spss-instruction.htm>
<http://www.glimo.vub.ac.be/downloads/eng-spss-basic-pdf>
<http://www.ats.ucla.edu/stat/spss/modules/>

خواننده محترم

این پرسشنامه به منظور ارتقای کیفیت کتاب‌های درسی و رفع نواقص آن‌ها تهیه شده است. دقت شما در پاسخگویی به این پرسشنامه در پایان هر نیمسال ما را در تحقق این هدف یاری خواهد کرد.

نام کتاب نام مؤلف/مترجم سال انتشار
 پاسخگو: عضو علمی پیام‌نور □ عضو علمی سایر دانشگاه‌ها □ رشته تخصصی سابقه تدریس
 دانشجوی پیام‌نور □ دانشجوی سایر دانشگاه‌ها □ رشته تحصیلی ورودی سال

سوال	بله	نه	متوسط	خوب	ضعیف	بسیار ضعیف
۱. آیا از زمان تحویل و نحوه دسترسی به کتاب راضی بودید؟						
۲. آیا حجم کتاب با توجه به تعداد واحد مناسب بود؟						
۳. آیا راهنمایی‌های لازم برای مطالعه کتاب منظور شده بود؟						
۴. آیا در ترتیب مطالب کتاب سلسله مراتب شناختی (آسان به مشکل) رعایت شده بود؟						
۵. آیا تقسیم‌بندی مطالب در فصل‌ها یا بخش‌ها متناسب و بجا بود؟						
۶. آیا متن کتاب روان و ساده و جمله‌ها قابل فهم بود؟						
۷. آیا به‌روزر بودن مطالب و آمارها رعایت شده بود؟						
۸. آیا مطالب تکراری داشت؟						
۹. آیا پیوستگی مطالب با درس‌های پیش‌نیاز رعایت شده بود؟						
۱۰. آیا مثال‌ها، شکل‌ها، نمودارها، جدول‌ها و ... گویا بودند و در فهم مطلب تأثیر داشتند؟						
۱۱. مطالعه هدف‌های کلی، آموزشی/ رفتاری تا چه اندازه به درک بهتر شما کمک کرد؟						
۱۲. آیا خودآزمایی‌های کتاب به‌گونه‌ای بود که تمام مطالب درسی را شامل شود؟						
۱۳. آیا پاسخ خودآزمایی‌ها و تمرین‌ها کامل و گویا بود؟						
۱۴. چقدر با غلط‌های املائی و اشکال‌های چاپی مواجه شدید؟						
۱۵. آیا از کیفیت چاپ و صحافی کتاب راضی بودید؟						
۱۶. آیا طرح روی جلد کتاب با مطالب کتاب تناسب داشت؟						
۱۷. چنانچه دانشگاه وسایل کمک‌آموزشی از قبیل نوار، فیلم، لوح فشرده و ... در اختیارتان گذارده، آیا به درک بهتر شما کمک کرده‌اند؟						
۱۸. تا چه اندازه این کتاب شما را از حضور در کلاس بی‌نیاز کرد؟						

در مجموع کتاب را چگونه ارزیابی می‌کنید؟ عالی □ خوب □ متوسط □ ضعیف □ بسیار ضعیف □
 لطفاً چنانچه با اشکال‌های تایپی یا محتوایی و مطالب تکراری مواجه شده‌اید، فهرستی از آن‌ها را با ذکر شماره صفحه ضمیمه کنید. در صورت تمایل سایر پیشنهاد‌های خود را نیز بنویسید.

این پرسشنامه را پس از تکمیل از کتاب جدا کنید و به قسمت آموزش مرکز تحویل دهید یا مستقیماً به نشانی تهران، صندوق پستی ۴۶۹۷-۱۹۳۹۵، مدیریت تولید مواد و تجهیزات آموزشی کتاب ارسال فرمایید. آدرس وبگاه ما www.pnu.ac.ir است. با ورود به وبگاه، مسیر زیر را طی نمایید: ساختار دانشگاه/ معاونت‌ها/ فناوری اطلاعات/ مدیریت تولید مواد و تجهیزات آموزشی.

با تشکر

مدیریت تولید مواد و تجهیزات آموزشی