

به نام خداوند جان و خرد ...



راهنمای آسان تحلیل آماری با SPSS



تالیف: رامین کریمی

سرشناسه : کریمی، رامین، ۱۳۶۵ -
 عنوان و نام پدیدآور : راهنمای آسان تحلیل آماری با SPSS
 مشخصات نشر : تهران: هنگام، ۱۳۹۴ .
 مشخصات ظاهری : ۳۱۸ص: مصور، جدول، نمودار
 شابک : ۱۳۰۰۰۰ ریال : 978-600-93084-2-2
 وضعیت فهرست نویسی : فیپای مختصر
 یادداشت : فهرست نویسی کامل این اثر در نشانی
<http://opac.nlai.ir> قابل دسترسی است
 شماره کتابشناسی ملی : ۳۸۳۱۴۴۱



راهنمای آسان تحلیل آماری با SPSS

نویسنده : رامین کریمی

طراح جلد : نیوشا اشراقی

ناشر : هنگام

نوبت چاپ : اول، تابستان ۱۳۹۴

شمارگان: ۵۰۰ نسخه

امور فنی : مجتمع چاپ و نشر هنگام

شابک : ۹۷۸-۶۰۰-۹۳۰۸۴-۲-۲

قیمت : ۱۳۰۰۰ تومان

کلیه حقوق برای مولف محفوظ است

تقديم به

مادر م

آموزش های تکمیلی نرم افزار SPSS

و آموزش نرم افزارهای LISREL ، Amos ، PLS

Minitab ، Eviews و داده کاوی و ... در سایت زیر در دسترس شماست:



www.kharazmi-statistics.ir

۸	پیش گفتار
	* فصل اول: آشنایی با اصطلاحات و نرم افزارهای آماری ۱۰
۱۱	معرفی نرم افزار SPSS
۱۲	تعریف آمار
۱۲	جمعیت آماری و نمونه آماری
۱۲	جمعیت آماری
۱۳	نمونه آماری
۱۴	سازه، مفهوم و متغیر
۱۵	سطح سنجش (اندازه گیری)
۱۶	سطح سنجش اسمی
۱۶	سطح سنجش ترتیبی
۱۹	سطح سنجش فاصله ای
۲۰	سطح سنجش نسبی
۲۲	متغیرهای گسسته و پیوسته
۲۵	آمار توصیفی و استنباطی
۲۵	تفاوت آماره و پارامتر
۲۵	آماره
۲۵	پارامتر
۲۶	تکنیک های آمار استنباطی
۲۶	برآوردهای فاصله ای یا خطای استاندارد
۲۸	سطح معنی داری
۲۹	آزمون های معنی داری
۳۳	فرضیه
۳۸	انواع متغیرها
۳۸	متغیر مستقل
۲۹	متغیر وابسته
۴۰	متغیر میانجی
۲۹	متغیر تعدیل کننده
۴۳	متغیر کنترل
۴۵	آزمون های پارامتری و ناپارامتری
۴۷	داده

	※ فصل دوم: ورود داده ها و آماده سازی برای تحلیل ۵۰
۵۱	بخش اول: آشنایی با برنامه و ورود داده ها
۵۱	آشنایی با پنجره ها و منوهای برنامه SPSS
۵۶	مراحل انجام یک تحلیل آماری
۵۷	تعریف متغیرها و وارد کردن داده ها
۵۷	تعریف متغیرها
۶۳	وارد کردن داده ها
۶۵	بخش دوم: آماده سازی داده ها برای تحلیل
۶۶	داده های پرت
۷۷	داده های ناقص
۸۵	کدگذاری مجدد متغیرها
۹۰	مقیاس سازی
	※ فصل سوم: توصیف داده ها و پیش فرض های آماری ۱۰۰
۱۰۱	بخش اول: توصیف داده ها
۱۰۲	توزیع فراوانی ها
۱۰۲	فراوانی و درصد فراوانی
۱۰۴	چارک
۱۰۷	شاخص های گرایش به مرکز
۱۰۷	میانگین
۱۰۸	میانه
۱۰۸	مد
۱۱۱	شاخص های پراکندگی
۱۱۱	انحراف استاندارد
۱۱۱	واریانس
۱۱۲	ضریب تغییرات
۱۱۳	دامنه تغییرات
۱۱۶	نمودارها
۱۱۶	نمودار دایره ای
۱۱۸	نمودار ستونی (میله ای)
۱۲۰	نمودار هیستوگرام

۱۲۴	بخش دوم: پیش فرض های آماری
۱۲۵	نرمال بودن داده ها
۱۳۴	خطی بودن رابطه (نمودار پراکندگی)
۱۴۱	هم خطی چندگانه
۱۴۶	یکسانی پراکندگی
۱۴۹	تبدیل داده ها
۱۵۵	بخش سوم: تقسیم و گزینش داده ها
۱۵۵	تقسیم کردن داده ها
۱۵۷	گزینش موردها
	* فصل چهارم: سنجش اعتبار و پایایی ۱۶۲
۱۶۳	بخش اول: تحلیل عاملی اکتشافی
۱۸۱	بخش دوم: پایایی (ضریب آلفای کرونباخ)
	* فصل پنجم: تحلیل داده ها (آمار استنباطی) ۱۹۰
۱۹۱	بخش اول: ضرایب همبستگی و پیوستگی
۱۹۶	آزمون فی و وی کرامر
۲۰۰	آزمون مجذور کای استقلال
۲۰۴	آزمون همبستگی اسپیرمن
۲۰۸	آزمون همبستگی پیرسون
۲۱۳	آزمون همبستگی تفکیکی (جزئی)
۲۲۰	بخش دوم: آزمون های مقایسه ای یا تفاوتی
۲۲۲	آزمون t تک نمونه ای
۲۲۵	آزمون t مستقل
۲۲۹	آزمون t همبسته
۲۳۳	آزمون تحلیل واریانس یک راهه
۲۴۴	بخش سوم: آزمون های چندمتغیره
۲۴۴	رگرسیون چند متغیره
۲۵۴	تحلیل مسیر
۲۶۷	رگرسیون لجستیک

فهرست مطالب

۲۸۰	بخش چهارم: آزمون های ناپارامتریک
۲۸۳	آزمون تک متغیری مجذور کای
۲۸۷	آزمون ویلکاکسون
۲۹۲	آزمون فریدمن
۲۹۶	آزمون مان-ویتنی
۳۰۰	آزمون کروسکال والیس
۳۰۷	راهنمای انتخاب آزمون آماری مناسب
۳۱۱	منابع:
۳۱۴	پیوست: پرسشنامه

پیش‌گفتار

می‌گویند هستی موضوع هیچ علمی به تنهایی نیست. این گفته بدین معنی است که پرداختن به زمینه‌های یک علم، فقط گوشه‌های بسیار کوچکی از هستی را برای ما روشن می‌کند. برای آن که دامنه شناخت وسعت یابد، لازم است چند شاخه علمی به هم گره بخورد و از روش‌های یکی، در دیگری استفاده شود. ریاضیات به عنوان یک علم زیربنایی، حضور خود را در بسیاری از علوم به اثبات رسانده است. آمار نیز از یک دیدگاه به عنوان علمی مستقل و از دیدگاهی دیگر به عنوان شاخه‌ای از ریاضیات، در علوم کاربردی جایگاه ویژه خود را دارد.

نرم‌افزار SPSS یکی از پرکاربردترین برنامه‌های کامپیوتری برای محاسبات آماری است. این برنامه در بسیاری از رشته‌های دانشگاهی کاربرد دارد و بسیاری از پژوهش‌گران حوزه‌های علوم انسانی، علوم پایه، کشاورزی، تربیت بدنی و ... نیز از این برنامه به طور گسترده استفاده می‌کنند.

آموختن برنامه SPSS مفید است چرا که امکان محاسبات آماری بسیاری را فراهم می‌کند و بسیاری از رشته‌های دانشگاهی نیازمند محاسبات آماری هستند، این برنامه جزء واحدهای درسی برخی رشته‌های دانشگاهی است و در پیشرفت علمی پژوهش‌گران و دانشجویان نقش موثری دارد. برنامه SPSS نیازی به یادگیری زبان برنامه‌نویسی ندارد و از این منظر آموختن این برنامه برای تمامی مخاطبان دانشگاهی و پژوهش‌گران امکان پذیر بوده و می‌توان با تمرین و مطالعه بر آن مسلط شد. این برنامه به عنوان یک امتیاز برای کسانی که جویای کار و یا خواهان پیشرفت در کار خود هستند تلقی می‌شود.

این کتاب آموزشی دارای ویژگی‌هایی است:

- ✓ تا حد امکان از بیانی ساده و روان جهت آموزش استفاده شده است.
- ✓ به آموزش آزمون‌های مهم و پرکاربرد پرداخته است.
- ✓ شیوه آموزش تصویری است.
- ✓ شیوه گزارش نتایج توضیح داده شده است.
- ✓ تمرین‌های متناسب برای همه آزمون‌ها ارائه شده است.
- ✓ دارای شکل‌ها و نمودارهای آموزشی است که فهم مطالب را آسان‌تر می‌کند.
- ✓ آموزش‌ها مبتنی بر جدیدترین نسخه (نسخه ۲۲) این نرم‌افزار است.

کتابی که پیش‌روی شماست، ۵ فصل دارد که هر فصل از چند بخش تشکیل شده است. در فصل اول به ارائه توضیحاتی در مورد اصطلاحات پرکاربرد آماری و روش شناختی پرداخته‌ایم که در فصل‌های بعدی این اصطلاحات به کار رفته است. در فصل دوم به آشنایی با ساختار و اجزاء برنامه SPSS و نحوه وارد کردن اطلاعات پرداخته‌ایم و سپس روندهای آماده کردن داده‌ها برای انجام تحلیل‌های آماری توصیفی و استنباطی را شرح داده‌ایم. در فصل سوم به آموزش آمار توصیفی مانند فراوانی‌ها، شاخص‌های گرایش به مرکز و پراکندگی و نمودارها پرداخته‌ایم و در انتهای این فصل مروری بر دستورهای تقسیم و انتخاب داده‌ها داشته‌ایم. در فصل چهارم به بررسی اعتبار و پایایی ابزار اندازه‌گیری یا پرسشنامه پرداخته‌ایم و روش تحلیل عاملی اکتشافی و ضریب آلفای کرونباخ را آموزش داده‌ایم. در انتها در فصل پنجم، به آموزش آمارهای استنباطی پرداخته‌ایم و ماهیت و نحوه اجرای آزمون‌های همبستگی، آزمون‌های مقایسه‌ای از نوع پارامتریک و ناپارامتریک و آزمون‌های چندمتغیره (رگرسیون چندمتغیره، تحلیل مسیر و رگرسیون لجستیک) را شرح داده‌ایم.

هر فصل کتاب شامل چند بخش است که هر بخش به طور معمول شامل قسمت آموزش نظری و محض، ارائه مثال عملی، شیوه اجرای دستور در برنامه، نحوه تفسیر یافته‌ها، شیوه گزارش نتایج، مدل آموزشی و ارائه تمرین است. در انتهای هر فصل واژه نامه مربوط به آن فصل آورده شده است. توضیحات تکمیلی و سایر آموزش‌های SPSS، به همراه آموزش نرم‌افزارهای آماری دیگر در سایت www.Kharazmi-Statistics.ir ارائه شده است که علاقه‌مندان می‌توانند به این سایت مراجعه کنند.

پیشنهادهای و نقطه‌نظرات خود را در مورد این کتاب می‌توانید به ایمیل اینجانب ارسال کنید. نظرات ارزشمند شما در ویرایش‌های بعدی کتاب لحاظ خواهد شد.

ایمیل: RKarimi777@yahoo.com سایت: www.Kharazmi-Statistics.ir

موفق باشید

آشنایی با اصطلاحات و مفاهیم آماری

محتوای فصل

« تعریف آمار، جمعیت آماری و نمونه آماری، سازه، مفهوم و متغیر، سطح سنجش، آمار توصیفی و استنباطی، تفاوت آماره و پارامتر، تکنیک‌های آمار استنباطی، فرضیه و انواع آن، انواع متغیرها، آزمون‌های پارامتری و ناپارامتری، داده»

معرفی نرم‌افزار SPSS

شرکت^۱ IBM یک شرکت تولیدکننده مادر در زمینه نرم‌افزارها و فن‌آوری‌هایی است که برای تحلیل داده‌ها، گزارش‌نویسی و مدل‌سازی مورد استفاده قرار می‌گیرد. در حال حاضر این شرکت در ایلینویز شیکاگو در کبک-شور ایالات متحده آمریکا واقع شده است. برنامه آماری SPSS یکی از تولیدات این شرکت است.

برنامه آماری^۲ SPSS در سیر تکاملی خود در سال ۱۹۶۸ SPSS نام گرفت که به معنی "بسته آماری برای علوم اجتماعی" است. این نرم‌افزار پس از تولید به سرعت در دانشگاه‌های سراسر قاره آمریکا رواج یافت.

در مورد تاریخ اختراع این نرم‌افزار باید گفت که اولین زمان اختراع SPSS به سال ۱۹۶۸ برمی‌گردد، یعنی زمانی که سه مرد جوان به نام‌های نورمن اچ. نای، سی. هالای هیول و دیل اپ. بنت با زمینه حرفه‌ای متمایز یک سیستم نرم‌افزاری را با این ایده توسعه دادند که بتوانند به کمک آن‌ها، داده‌های خام را به اطلاعات لازم برای تصمیم‌گیری تبدیل کنند.

نای دانشجوی دوره دکترای جامعه‌شناسی در دانشگاه استنفورد مخاطبان هدف و مقتضیات برنامه را تعریف کرد. هیول، دانشجوی دکتری تحقیق در عملیات با گرایش تجزیه و تحلیل در دانشگاه استنفورد، ساختار فایل سیستم را طراحی کرد و بالاخره بنت، فارغ‌التحصیل دوره کارشناسی ارشد مدیریت اجرایی، کار برنامه‌نویسی این نرم‌افزار را تکمیل کرد.

SPSS نرم‌افزاری است که به منظور تحلیل داده موفق به کسب "استاندارد صنایع" شده است. این نرم‌افزار در میان پژوهش‌گران دانشگاه‌ها بیشترین میزان کاربرد را دارد، به‌ویژه پژوهش‌گرانی که در روان‌شناسی و علوم اجتماعی فعالیت می‌کنند این برنامه کاربرد وسیعی در پژوهش‌های سازمان‌های خصوصی، دولتی و شرکت‌های بزرگ خصوصی دارد.

^۱ International Business Machines Corp

^۲ Statistical Package for the Social Sciences

تعریف آمار

آمار دانشی است که با استفاده از فنون و روش‌های علمی-ریاضی به جمع‌آوری، طبقه‌بندی، توصیف، تبیین و پیش‌بینی اطلاعات کمی و کیفی و نتیجه‌گیری و تعمیم آن‌ها در جهت هدف معین می‌پردازد (ساعی، ۱۳۸۸: ۴۰). بر این اساس می‌توان گفت که آمار دو نقش اساسی در پژوهش‌های اجتماعی دارد: ۱- توصیف و تبیین ۲- تعمیم. در مرحله اول آمار به توصیف خصوصیات نمونه پرداخته و آماره‌های مختلفی مانند فراوانی، میانگین و انحراف استاندارد متغیرها را نشان می‌دهد. در همین مرحله به بررسی روابط و تأثیرات بین متغیرها هم پرداخته می‌شود. در مرحله تعمیم با استفاده از پارامترهایی مانند خطای استاندارد و سطح معنی‌داری (سطح خطا)، وضعیت تعمیم‌پذیری نتایج از نمونه به جمعیت آماری ارزیابی می‌شود.

جمعیت آماری و نمونه آماری

جمعیت آماری

جمعیت (جامعه) آماری پژوهش، موارد و یا کسانی هستند که :
 الف) قصد داریم درباره آنان اطلاعاتی برای پژوهش گردآوری کنیم و از بین آنان نمونه پژوهش را انتخاب کنیم.
 ب) ناچاریم برای کسب معرفت از جامعه به نظرات آنان (در مورد افراد) و یا ویژگی‌های آنان (در مورد اشیا و مکان‌ها و ...) توجه نماییم.
 در مجموع در مورد تعریف جمعیت آماری می‌توان گفت که جمعیت آماری عبارت است از گروهی از افراد یا اشیا (آزمودنی‌ها) که در خصوصیات یا ویژگی‌های مورد پژوهش مشترک می‌باشند و با هدف و موضوع پژوهش ارتباط دارند. جمعیت آماری را با N نشان می‌دهند.
 به عنوان مثال زمانی که قصد داریم در زمینه رضایت شغلی معلمان پژوهش کنیم، جمعیت پژوهش ما معلمان هستند و برای گردآوری اطلاعات باید به معلمان مراجعه کنیم و نمونه خود را از میان معلمان انتخاب کنیم. اما هنگامی که می‌خواهیم از تأثیر

سابقه خدمتی معلم در نحوه برخورد معلم با دانش‌آموز پژوهش کنیم این کار را با استفاده از یک پرسشنامه انجام می‌دهیم. پژوهش بر روی معلمان است ولی باید برای سنجش نحوه برخورد معلم با دانش‌آموزان، به نظرات دانش‌آموزان مراجعه کنیم. در چنین حالتی ابتدا باید معلمان را از نظر سابقه خدمت به گروه‌های مختلف تقسیم کنیم، سپس از میان معلمان در هر گروه عده‌ای را انتخاب کرده و به دانش‌آموزان آن معلم مراجعه کنیم. در این صورت معلمان و دانش‌آموزان هر دو نمونه پژوهش هستند و انتخاب دانش‌آموز تابعی از انتخاب معلم می‌باشد و جمعیت پژوهش شامل هر دو گروه معلمان و دانش‌آموزان است.

همچنین هنگامی می‌خواهیم درباره میزان محبوبیت معلمان در نزد دانش‌آموزان به پژوهش بپردازیم، با این که پژوهش درباره معلمان است ولی جمعیت و نمونه دقیق دانش‌آموزان هستند، زیرا داده‌ها نظرات دانش‌آموزان است (صفری‌شالی، ۱۳۸۶: ۸۳-۸۴).

نمونه آماری

نمونه، عضوی از جمعیت آماری است که ویژگی‌های غالب اعضاء جمعیت آماری را داراست و در واقع معرف جمعیت و یا مجموعه آزمودنی‌ها می‌باشد و نتایج حاصل از مطالعه آن، قابل تعمیم به کل جمعیت است. نمونه آماری را با π نشان می‌دهند.

تمام اعضای گروه، جمعیت خوانده می‌شود. نمونه حاصل گردآوری اطلاعات فقط درباره تعدادی از اعضای جمعیت است. نمونه می‌تواند با درجات مختلفی از دقت بازتاب جمعیتی باشد که از آن برگرفته شده است. نمونه‌ای که دقیقاً بازتاب جمعیت خود باشد نمونه معرف خوانده می‌شود.

همچنین تفاوت سرشماری و نمونه‌گیری این است که سرشماری حاصل کسب اطلاعات از همه اعضای گروه است ولی در نمونه‌گیری ما فقط درباره تعدادی از جمعیت به گردآوری اطلاعات می‌پردازیم.

به عنوان مثال اگر قصد داشته باشیم میزان رضایت شغلی معلمان استان تهران را به دست بیاوریم، در این صورت جمعیت آماری شامل تمامی معلمان استان می‌شود و از آن جا که از نظر زمانی و مالی امکان مطالعه تمامی معلمان را نداریم تعداد ۳۰۰ نفر از معلمان را به عنوان نمونه آماری انتخاب کرده و پژوهش را روی این ۳۰۰ نفر انجام می‌

دهیم. اگر فرض کنیم کل معلمان تهران ۵۰۰۰۰ نفر باشند، این ۵۰۰۰۰ نفر جمعیت آماری می‌شوند و ۳۰۰ نفری که به تصادف از بین جمعیت آماری انتخاب می‌شوند، نمونه آماری هستند.

اما چرا خصوصیات جمعیت آماری را از روی نمونه برآورد می‌کنیم و پژوهش خود را بر روی جمعیت آماری اجرا نمی‌کنیم؟ در پاسخ به این سوال باید گفت که اجرای پژوهش بر روی جمعیت آماری از نظر مالی پرهزینه بوده، زمان‌بر بوده و همچنین گاهی دسترسی به کل جمعیت آماری ممکن نبوده و نمی‌توان پژوهش را بر روی تمامی اعضای جمعیت اجرا کرد.

● سازه، مفهوم و متغیر

سازه: سازه‌ها، که از نظریه مشتق می‌شوند، مفاهیمی هستند که از بالاترین سطح انتزاع برخوردارند. هر سازه می‌تواند متشکل از چندین مفهوم باشد. مثالی که در این مورد می‌توان ارائه داد، احساس خوشبختی است. احساس خوشبختی، یک سازه است که در حالت معمولی قابل مشاهده نیست. بنابراین، برای مشاهده و سنجش آن، باید از مفاهیم تشکیل‌دهنده آن مانند احساس نشاط، احساس سرزندگی، احساس رضایت (فردی، اجتماعی، اقتصادی و سیاسی)، احساس امید و ... استفاده کرد.

مفهوم: مفهوم، تجریدی از رویدادهای قابل مشاهده است که معرف شباهت‌ها یا جنبه‌های مشترک میان آن‌هاست. در واقع مفاهیم، واژه‌های انتزاعی‌اند که برای توضیح دادن و یا معنا بخشیدن به تجربیاتمان از آن‌ها استفاده می‌کنیم. برای مثال، پیشرفت تحصیلی مفهومی است که می‌توان آن را از طریق عملکرد (نمره‌های کلاسی) دانش‌آموزان در دروس مختلف مشاهده کرد. هر مفهوم از چندین متغیر در مقیاس‌های مختلف تشکیل شده است. جهت توضیحات بیشتر به مثال سازه برمی‌گردیم که در آن، احساس رضایت فردی در زندگی به عنوان یکی از مفاهیم در سازه احساس خوشبختی معرفی شده است. در این جا برای این که بتوانیم مفهوم رضایت فردی را مورد سنجش و مشاهده قرار دهیم، باید از چندین متغیر مانند رضایت از درآمد، رضایت از شغل، رضایت از ساعات کاری و ... استفاده کنیم.

متغیر: متغیر عبارت است از ویژگی واحد مورد مشاهده. متغیر کمّیتی است که می‌تواند از واحدی به واحد دیگر یا از یک شرایط مشاهده به شرایط دیگر، مقادیر مختلفی را اختیار کند. متغیر نمادی است که اعداد یا مقادیر به آن اختصاص پیدا می‌کند. به عبارت دیگر، متغیر به ویژگی‌هایی اطلاق می‌شود که می‌توان آن‌ها را مشاهده یا اندازه‌گیری کرد و دو یا چند عدد یا نماد را جایگزین آن‌ها قرار داد. به عنوان مثال، جاده یک مفهوم است، اما طول جاده یک متغیر است. در محیط SPSS، متغیر صفتی است که مایل به انجام تحلیل‌های آماری بر روی آن هستیم. در واقع در این محیط، هر سوال پرسشنامه، یک متغیر محسوب می‌شود (حبیب‌پور و صفری‌شالی، ۱۳۸۸: ۳۱-۳۲).

● سطح سنجش (اندازه‌گیری)

هر متغیر شامل دو یا چند طبقه یا ویژگی است. مثلاً جنس متغیری است با دو طبقه زن و مرد، و محل تولد متغیری است با چند طبقه که همان محل تولد افراد است. سطح سنجش یا متغیرها به نحوه رابطه طبقات متغیر با یکدیگر مربوط است.

بر اساس نوع اندازه‌گیری، متغیرها به دو دسته کمّی و کیفی تقسیم می‌شوند. متغیر کمّی متغیری است که از نظر کمّیت (مقدار) تغییر می‌کند و می‌توان اختلاف مقادیر آن را با استفاده از عدد ثبت کرد و آن‌ها را با هم جمع کرد (جمع‌پذیر هستند). به بیان دیگر، متغیرهایی هستند که انسان توانسته است برای آن‌ها واحد و مبدأ اندازه‌گیری معین کند مانند قد، وزن و سن.

متغیر کیفی، به متغیری اطلاق می‌شود که اختلاف و تغییرات بین میزان‌های مختلف آن کیفی است و برای ثبت آن ممکن است از حروف الفبا یا کد استفاده شود. این‌گونه متغیرها را نمی‌توان جمع و تفریق کرد و برای آن‌ها مبدأ اندازه‌گیری نیز وجود ندارد (دل‌آور، ۱۳۷۸: ۴۲). متغیرهایی مثل جنس (زن یا مرد)، رنگ مو، قومیت (فارس، ترک، کرد و...)، احساس خوشبختی، بزهکاری، تعهد شغلی و انسجام اجتماعی از این دسته هستند. به عبارت دیگر متغیر کیفی متغیری است که نمی‌توان آن را با ابزار کمّی اندازه‌گیری کرد، ولی سنجش آن به صورت طبقه‌بندی صورت می‌گیرد. به عبارت دیگر، متغیر کیفی متغیری است که به صورت طبقه‌بندی قابل اندازه‌گیری است (ناییبی، ۱۳۸۸: ۱۲).

اکثر نویسندگان حیطه آمار، بر مبنای تقسیم‌بندی و تمایزی که استیونز (۱۹۴۶) بین سه سطح سنجش اسمی، ترتیبی و فاصله‌ای/نسبی قائل شده است عمل می‌کنند (برایمن و کرامر، ۲۰۰۱: ۵۶). در ادامه به تشریح هرکدام از این سه سطح سنجش (سطح اندازه‌گیری) پرداخته‌ایم.

۱) سطح سنجش اسمی

متغیر اسمی متغیری است که می‌توان بین طبقات آن تمایز قائل شده اما نمی‌توان طبقات آن را رتبه‌بندی کرد. در این سطح سنجش، افراد متعلق به طبقات مختلف با هم تفاوت دارند اما کمی کردن میزان تفاوت آن‌ها بی‌معناست. گاهی به متغیرهای اسمی، متغیرهای کیفی یا طبقه‌ای گفته می‌شود. محل تولد، جنس و وضع تاهل همه مصداق‌های متغیر اسمی‌اند.

مثال:

جنسیت: شامل دو طبقه زن و مرد.
 قومیت: شامل چند طبقه (فارس، آذری، کرد، لر و...).
 اصلیت: شامل دو طبقه بومی و غیر بومی.
 ملیت: شامل طبقات متعدد (ایرانی، عراقی، افغانی و ...).
 وضع تاهل: شامل دو طبقه متاهل و غیرمتاهل؛ و یا چندین طبقه: متاهل، ازدواج نکرده، همسر فوت شده، همسر جدا شده و همسر ترک کرده.
 نوع مسکن: شامل ملکی، اجاره‌ای، سازمانی و رایگان.
 گرایش سیاسی: شامل راست، میانه و چپ.
 سواد: شامل دو طبقه باسواد و بی‌سواد.
 نوع شغل: شامل تمام وقت، نیمه وقت و فصلی.

۲) سطح سنجش ترتیبی

متغیر ترتیبی متغیری است که رتبه‌بندی طبقات آن با معناست. بین طبقات مراتب قابل قبولی وجود دارد ولی با این وجود، کمی کردن دقیق میزان تفاوت امکان‌پذیر نیست. می‌توان از مردم پرسید تا چه حد با مطلب خاصی موافق‌اند یا مخالف. طبقات را می‌توان بر حسب شدت موافقت با مطلبی یا نگرشی رتبه‌بندی کرد. اگر از مردم بپرسیم

وضع اشتغال آن‌ها چگونه است و سه گزینه «اصلا کار نمی‌کنم»، «پاره‌وقت» و «تمام وقت» ارائه کنیم این متغیر متغیری ترتیبی خواهد بود. هر متغیری که بتوان طبقات آن را رتبه‌بندی کرد اما نتوان تفاوت طبقات آن را دقیقا به صورت عددی کمی ارائه کرد، متغیر ترتیبی است.

مثال:

طبقه اجتماعی: که عموما شامل سه طبقه است: بالا، متوسط و پایین.
 قشر اجتماعی: که شامل سه طبقه (بالا، متوسط و پایین) یا پنج طبقه (بالا، متوسط بالا، متوسط، متوسط پایین و پایین).
 نگرش‌ها: که عموما شامل سه طبقه (موافق، بی‌نظر و مخالف) و یا پنج طبقه (کاملا موافق، موافق، بی‌نظر، مخالف، کاملا مخالف).
 سلسله مراتب سازمانی: شامل چندین طبقه: رئیس، معاون، مدیر بخش، معاون بخش، کارمند معمولی.
 سطح پیشرفت صنعتی: شامل سه طبقه: کشورهای پیشرفته، در حال توسعه و توسعه نیافته.

گروه‌های درآمدی: شامل سه طبقه، پردرآمد، میان درآمد و کم درآمد.
☒ نکته: وقتی متغیرهای اندازه‌گیری شده مربوط به مقیاس‌های روانی اجتماعی، پرسش نامه-های روان شناختی یا آزمون‌های شناختی هستند، ممکن است تفاوت‌هایی در سطح اندازه‌گیری متغیر مشاهده کنیم. بسیاری از این مقیاس‌ها دارای نقاط صفر قراردادی هستند که توسط سازنده آزمون مشخص گردیده و در مقایسه با اندازه خط‌کش استاندارد اینچ یا سانتی متر، فاقد واحد مشخص اندازه‌گیری هستند. از نظر فنی، این متغیرها دارای ماهیت ترتیبی (رتبه‌ای) است، اما در عمل، پژوهش‌گران به آن‌ها در قالب مقیاس‌های فاصله‌ای یا نسبی نگاه می‌کنند و این موضوع سال‌هاست که در متون پژوهشی مورد بحث و اختلاف نظر بوده است (مونرو، ۱۳۸۹: ۱۶). در مورد سطح سنجش متغیرهای ترتیبی نظرات متفاوتی وجود دارد. تامل در برخی از این دیدگاه‌ها خالی از فایده نیست:

دیدگاه اول: گاهی برخی از پژوهش‌گران نمرات چند سوال ترتیبی را با هم جمع کرده و اقدام به ساخت مقیاس می‌کنند (مانند جمع کردن نمرات افراد در چند سوالات و ساختن مقیاس) و یا و بعد از آن نمرات مقیاس به دست آمده را در سطح سنجش فاصله‌ای در نظر گرفته و از آزمون‌های پارامتری استفاده می‌کنند. میلر معتقد است که

چون در محاسبه آزمون‌های پارامتری نمرات به دست آمده از نمونه را جمع و تقسیم و ضرب می‌کنیم در نتیجه این اعمال ریاضی را تنها برای نمراتی می‌توان به کار برد که واقعا عددی باشند و استفاده از این آزمون‌های پارامتری برای این داده‌ها، باعث تحریف داده‌ها و در نتیجه تردید درباره نتیجه‌گیری از روی آزمون می‌شود. از این‌رو، کاربرد تکنیک‌های پارامتری فقط برای نمراتی است که واقعا عددی‌اند (میلر، ۱۳۸۰).

دیدگاه دوم: متغیر ترتیبی متغیری کیفی است اما چون طبقاتش نسبت به هم بالا و پایین‌اند تا حدی به متغیرهای کمی نزدیک می‌شود. از این‌رو، در برخی مواقع مانند ترکیب کردن چند متغیر ترتیبی و ساختن مقیاسی برای یک مفهوم انتزاعی (مانند مقیاس‌های رضایت شغلی، اعتماد اجتماعی، اضطراب و ...) یا برای تحلیل‌های رگرسیونی آن‌ها را با تسامح مقیاس فاصله‌ای در نظر می‌گیرند. جز این موارد به هیچ وجه نباید با متغیر ترتیبی به مانند متغیری کمی عمل کرد. به طور مشخص نباید برای متغیر ترتیبی میانگین و میانه یا چارک حساب کرد؛ چنین اندازه‌هایی برای متغیر ترتیبی بی‌معناست (نایی، ۱۳۸۸: ۱۹).

دیدگاه سوم: مقیاس پاسخ تراکمی (مانند مقیاس لیکرت) مستلزم آن است که پاسخ دهندگان بر اساس طیف زمینه‌سازی که به وسیله لنگرگاه تعریف شده است، اندازه‌هایی را به عناصر اختصاص دهند. اعداد به طور معمول به صورت صعودی مرتب می‌شوند تا بیشتر خصیصه‌ای را که درجه بندی می‌شوند، منعکس کنند. متداول‌ترین مقیاس پاسخ تراکمی مقیاس‌های ۵ و ۷ درجه است. مقیاس تراکمی به این دلیل به این نام خوانده می‌شود که امکان جمع کردن رتبه‌ها با هم و تقسیم آن‌ها بر مقدار ثابتی (فرآیندی که به طور معمول برای تعیین میانگین استفاده می‌شود) برای تعیین نمره هر فرد در پرسشنامه وجود دارد. به بیان دیگر میانگین به دست آمده از مقیاس پاسخ تراکمی با معناست (میزر و دیگران، ۱۳۹۱: ۵۰-۴۹).

لیکرت (۱۹۳۲) اظهار می‌دارد که روش درجه‌بندی او (مقیاس لیکرت) با مقیاس درجه بندی ترستون همبستگی نزدیک به ۱ دارد که به ظاهر هر دو دارای روش فاصله‌ای هستند. گیلفورد (۱۹۴۵)، در کتابش به نام "روش‌های روان‌سنجی" تأیید کرد که مقیاس‌های پاسخ تراکمی حداقل، بیشتر ویژگی‌های مقیاس فاصله‌ای را در مقایسه با مقیاس‌های دیگر داراست. ادواردز (۱۹۵۷) گامی جلوتر می‌رود و بیان می‌کند که «اگر

به مقایسه میانگین نمره‌های نگرش دو یا چندگروه علاقه‌مندیم، این کار را می‌توانیم همانند مقیاس‌های فاصله‌ای با فواصل یکسان به وسیله مقیاس‌های رتبه‌بندی تراکمی انجام دهیم.

اگر چه ممکن است برخی پژوهش‌گران از این که تا چه اندازه نقاط روی مقیاس پاسخ تراکمی به واقع با فواصل یکسان توزیع شده‌اند، دچار تردید باشند، بیشتر پژوهش‌های منتشر شده در علوم رفتاری و علوم اجتماعی در نیم قرن گذشته، بیشتر از مقیاس‌های پاسخ تراکمی استفاده کرده‌اند، به گونه‌ای که گویی آن‌ها به مقیاس‌های فاصله‌ای شباهت دارند. پژوهشگران، نمره‌های مقیاس را با هم جمع می‌کنند و میانگین را به دست می‌آورند و این اندازه‌گیری‌ها را در تحلیل آماری که به طور معمول برای تفسیر بهتر نتایج که نیازمند اندازه‌گیری فاصله‌ای یا نسبی است، به کار می‌برند. به نظر ما، این روش در مقیاس‌های پاسخ تراکمی، قابل قبول، مناسب و کاملاً سودمند است. برای مقیاس‌های پاسخ تراکمی می‌توان میانگین را با داشتن همه ویژگی‌های مناسب آن اندازه‌گیری کرد و راه برای انجام دامنه گسترده‌ای از روش‌های آمار پارامتری مانند همبستگی پیرسون و تحلیل واریانس (آنووا)، در مورد آن باز است (میزر و دیگران، ۱۳۹۱: ۵۳-۵۴).

۲) سطح سنجش فاصله‌ای

متغیر فاصله‌ای متغیری کمی است که واحدهای اندازه‌گیری آن هم ارز و مساوی‌اند، اما فاقد صفر حقیقی است. به عبارت دیگر، واحدهای اندازه‌گیری متغیر فاصله‌ای با هم برابرند و می‌توان آن‌ها را با هم جمع کرد یا از هم تفریق کرد، اما نمی‌توان آن‌ها را بر هم تقسیم کرد و از آن‌ها نسبت گرفت (نایی، ۱۳۸۸: ۱۴). واحدهای اندازه‌گیری متغیر فاصله‌ای مساوی‌اند، اما فاقد صفر حقیقی است.

برای مثال متغیر دما متغیری فاصله‌ای است که مقادیر آن با واحد استاندارد به نام درجه سانتیگراد اندازه‌گیری می‌شود. اگر در ظهر دمای هوا ۶ درجه سانتیگراد باشد و در عصر که هوا سردتر می‌شود به ۳ درجه سانتیگراد برسد، درجه دمای عصر را از درجه دمای ظهر کم کرده، می‌گوییم عصر ۳ درجه سردتر شده است یا ظهر ۳ درجه گرم‌تر از عصر بود. اما نمی‌توانیم از این دو درجه دما نسبت بگیریم و بگوییم که دمای ظهر دو برابر عصر است یا دمای عصر یک دوم دمای ظهر است. علت این امر این است که دما از

صفر حقیقی (که حدود درجه ۲۷۳- درجه سانتیگراد) شروع نمی‌شود، بلکه از دمای صفر قراردادی شروع می‌شود که همان دمای انجماد آب است؛ یعنی دمایی که میزان دمای مناطق مسکونی عموماً حول و حوش آن است. این امر برای سهولت و کوتاه کردن مقیاس سنجش دما به عنوان مبدأ (صفر) درجه سانتیگراد انتخاب شده است.

متغیر بهره هوشی نیز متغیری فاصله‌ای است. به عنوان مثال، اگر دانش آموزی از یک آزمون هوش نمره صفر بگیرد، به این معنا نیست که او اصلاً هوش ندارد و یا کسی که نمره بهره هوشی او ۱۰۰ است نمی‌توان گفت که بهره هوشی او دقیقاً دو برابر کسی است که بهره هوشی او ۵۰ است. چرا که متغیر بهره هوشی فاقد صفر حقیقی است.

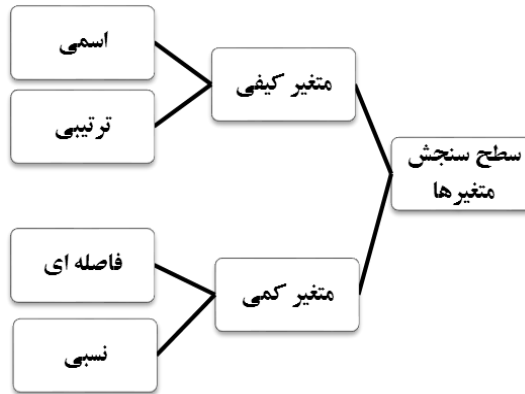
۴) سطح سنجش نسبی

چهارمین و دقیق‌ترین سطح اندازه‌گیری، مقیاس نسبی است که در آن، ویژگی‌های نمونه به صورت معنی‌داری رتبه‌بندی شده‌اند، فواصل بین رتبه‌ها مساوی‌اند و وجود نقطه صفر نه به صورت داوری و دلخواه، بلکه بر اساس ماهیت تعیین می‌شود (مونرو، ۱۳۸۹: ۱۵).

متغیر نسبی (نسبتی) متغیری کمی است که واحدهای اندازه‌گیری آن هم‌ارز و مساوی‌اند و دارای صفر حقیقی است. متغیر نسبی علاوه بر این که دارای همان خصوصیات متغیر فاصله‌ای است (واحدهای آن هم‌ارز و معادل‌اند)، چون دارای صفر حقیقی است می‌توان واحدهای آن را بر هم تقسیم یا نسبت آن‌ها را حساب کرد؛ از همین رو به آن متغیر نسبی یا نسبتی می‌گویند. برای مثال سن بر حسب سال متغیری نسبی است که واحد سنجش آن سال است. کسی که تازه به دنیا آمده صفر ساله است؛ علی ۴۰ سال دارد و سن او دو برابر حمید است که ۲۰ سال سن دارد. متغیرهای تحصیلات بر حسب سال و درآمد بر حسب ارز نیز نسبی هستند. تعداد هر مجموعه‌ای متغیری نسبی است، مانند جمعیت، تعداد خانوار، تعداد کارمندان و تعداد واحدهای مسکونی. نسبت و میزان مانند میزان مولید، میزان رشد جمعیت، میزان طلاق و ... نیز متغیرهای نسبی‌اند.

توجه: تفاوت دو مقیاس فاصله‌ای و نسبی در عمل چندان زیاد نیست و می‌توان در هنگام انجام تحلیل‌های آماری این دو مقیاس را یکسان فرض کرد. نرم افزار SPSS بین

داده‌های فاصله‌ای و نسبی تمایزی قائل نمی‌شود و از اصطلاح Scale برای توصیف هر دو دسته از متغیرهای فاصله‌ای و نسبی استفاده می‌کند. یعنی در برنامه SPSS، داده‌ها در سه دسته متغیرهای اسمی، ترتیبی و فاصله‌ای-نسبی طبقه بندی می‌شوند.



شکل ۱-۱- انواع سطوح سنجش متغیرها

چه سطحی از سنجش مناسب‌تر است؟

چون تعیین سطح سنجش بستگی به تصمیم پژوهش‌گر دارد در ارتباط با این سوال که چه سطحی از سنجش در پژوهش‌ها مناسب‌تر است باید به چند نکته توجه داشت:

- ۱- هر چه بر سطح سنجش افزوده شود دامنه روش‌های تحلیل قابل استفاده گسترده‌تر می‌گردد.
- ۲- تکنیک‌های قوی و پیچیده تحلیل، برای متغیرهای فاصله‌ای/نسبی مناسب است.
- ۳- سطوح بالاتر سنجش، اطلاعات بیشتری فراهم می‌آورد.
- ۴- پرسش‌هایی که مستلزم دقت و جزئیات زیادی‌اند ممکن است غیرقابل اعتماد باشند، چرا که مردم غالباً اطلاعات دقیق و مفصل ندارند (مثلاً ممکن است بهره هوشی خود را به طور دقیق ندانند).
- ۵- ممکن است مردم از ارائه اطلاعات دقیق اکراه داشته باشند اما از بیان کلی‌تر آن اکراه یا ابایی نداشته باشند (مثلاً مشخص کردن گروه سنی یا طبقه درآمد).

۶- متغیرهایی که در سطح فاصله‌ای/نسبی سنجیده شده‌اند را به سادگی می‌توان به سطح ترتیبی یا اسمی تقلیل داد ولی داده‌هایی را که در سطوح پایین‌تر اندازه‌گیری شده‌اند نمی‌توان به سطوح بالاتر تبدیل کرد (دواس، ۱۳۷۶: ۱۳۶).

خلاصه آن که به‌طور کلی توصیه می‌شود متغیر را در بالاترین سطح سنجش سنجید، اما ملاحظات مربوط به پایایی، میزان پاسخگویی و مقتضیات پژوهش به این معناست که اغلب اندازه‌گیری در سطوح پایین‌تر عاقلانه‌تر است.

جدول ۱-۱- راهنمای تشخیص سطح سنجش

معیارها	اسمی	ترتیبی	فاصله‌ای	نسبی
آیا طبقات مختلفی وجود دارد؟	بله	بله	بله	بله
آیا می‌توان طبقات را رتبه بندی کرد؟	خیر	بله	بله	بله
آیا می‌توان اختلاف طبقات را به صورت عددی مشخص کرد؟	خیر	خیر	بله	بله
آیا می‌توان داده‌ها را بر هم تقسیم کرده و نسبت گرفت؟	خیر	خیر	خیر	بله

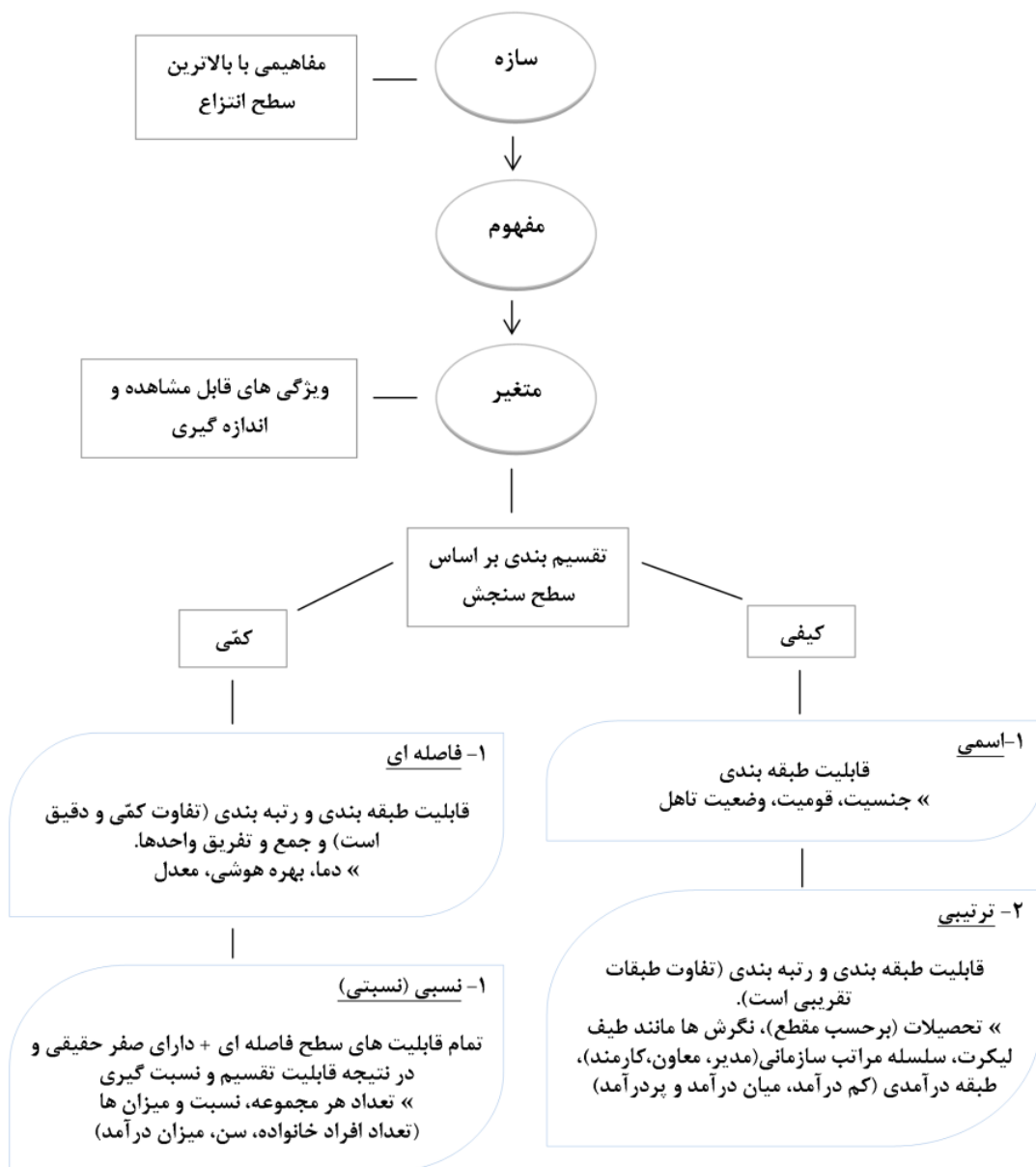
متغیرهای گسسته و پیوسته

بر مبنای یک تقسیم‌بندی دیگر، متغیرها به دو دسته متغیرهای گسسته و متغیرهای پیوسته تقسیم می‌شوند. متغیر گسسته می‌تواند اعداد یا ارزش‌هایی را که مشخص کننده یک وجه مشخص و معین از یک مقیاس هستند، به خود اختصاص دهد. به عنوان مثال، جنس یک متغیر گسسته است: یک شخص یا زن است یا مرد. اختصاص هر نوع ارزش دیگری بین این دو نوع ارزش امکان‌پذیر نیست. تعداد بازیکنان یک تیم فوتبال نیز یک متغیر گسسته است، زیرا فقط امکان داشتن ۱، ۲، ۳، ۴، ۵، ۶، ۷، ۸، ۹، ۱۰، ۱۱ بازیکن وجود دارد و نه ۷/۵ نفر بازیکن.

متغیر پیوسته، متغیری است که بین دو واحد آن هر نقطه یا ارزشی را می‌توان انتخاب کرد. در این متغیر درجات مختلف اندازه‌گیری وجود دارد و دقت وسیله اندازه‌گیری، تعداد این درجات را تعیین می‌کند. به عنوان مثال، وزن یک متغیر پیوسته است و می‌

تواند بین صفر تا بی‌نهایت باشد. وزن یک شخص می‌تواند ۵۵ یا ۵۶ کیلوگرم باشد و یا می‌تواند هر عددی بین این دو عدد باشد (مثلاً ۵۵.۶ یا ۵۵.۶۴ کیلوگرم). قد، زمان، طول یا ارتفاع پرش، درصد چاقی بدن، و سطح هموگلوبین خون متغیرهای پیوسته هستند.

البته در عمل تشخیص بین متغیر پیوسته و گسسته به صورت نظری امکان‌پذیر نیست. دلیل این امر فقدان وسایل اندازه‌گیری دقیق و مناسب است. در خیلی از متغیرهای پیوسته ما ناچاریم اعداد را به صورت کلی برای اندازه‌گیری به کار ببریم. بهره هوشی از جنبه نظری یک متغیر پیوسته است، اما، آزمونی که برای اندازه‌گیری هوش به کار برده می‌شود، به گونه‌ای است که نمره‌ها را به صورت نمره‌های گسسته نشان می‌دهد.



نکته: برنامه SPSS دو سطح سنجش فاصله ای و نسبی را یکسان در نظر می گیرد

آمار توصیفی و استنباطی

آمار به دو نوع اصلی تقسیم می‌شود: توصیفی و استنباطی. آمار توصیفی آماری است که الگوی پاسخ‌های افراد نمونه را تلخیص کرده و نشان می‌دهد. آمار توصیفی اطلاعاتی درباره مثلاً «متوسط» درآمد پاسخگویان فراهم می‌آورد یا نشان می‌دهد که سطح تحصیلات بر الگوی رأی دادن افراد نمونه مؤثر است یا خیر. اما ما عموماً در پی اطلاع از نگرش‌ها و خصوصیات مثلاً ۲۰۰۰ نفر اعضای نمونه نیستیم، بلکه در پی آن هستیم که نتایج نمونه را به جمعیت تعمیم دهیم. کار آمار استنباطی نشان دادن این نکته است که آیا الگوهای توصیف شده در نمونه، کاربردی در مورد جمعیتی که نمونه از آن انتخاب شده دارد یا خیر. با کمک آمار استنباطی می‌توانیم میزان نزدیکی آراء نمونه به آراء جمعیت را برآورد کنیم: آمار استنباطی ما را قادر به استنباط ویژگی‌های جمعیت از روی ویژگی‌های نمونه می‌کند (دواس، ۱۳۷۶: ۱۳۷-۱۳۸).

تفاوت اصلی آمار توصیفی و استنباطی این است که در آمار توصیفی هیچ‌گاه نمی‌توان نتایج به دست آمده از نمونه آماری را به کل جمعیت آماری تعمیم داد. درحالی که در آمار استنباطی و یا تحلیلی می‌توان یافته‌های به دست آمده از نمونه آماری را به کل جمعیت آماری پژوهش تعمیم داد. مفهوم کانونی آمار استنباطی، تعمیم‌پذیری است.

تفاوت آماره و پارامتر

آماره

آماره‌ها، اندازه‌گیری‌هایی هستند که ویژگی‌های یک نمونه آماری را بیان می‌کنند و از آن‌ها برای توصیف ویژگی‌های یک نمونه آماری استفاده می‌شود. به ارزش‌های عددی که داده‌های نمونه را خلاصه و توصیف می‌کنند، آماره می‌گویند.

پارامتر

پارامترها، اندازه‌گیری‌هایی هستند که ویژگی‌های یک جمعیت آماری را بیان می‌کنند. هدف عمده آمار استنباطی، استنباط پارامتر (ویژگی‌های جمعیت آماری) از آماره (ویژگی‌های نمونه آماری) است (ساعی، ۱۳۸۸: ۴۲). پارامتر را با حروف یونانی

و آماره را با حروف انگلیسی نشان می‌دهند. جدول ۱-۲ روش‌های نمایش شاخص‌های آماری در نمونه (آماره) و و جمعیت یا جامعه (پارامتر) را نشان می‌دهد.

جدول ۱-۲- مقایسه نمادهای آماری در نمونه (آماره) و در جمعیت آماری (پارامتر)

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

تکنیک‌های آمار استنباطی

آمار استنباطی شامل یک سری روش‌های آماری است که بر اساس اطلاعات نمونه، درباره ویژگی‌های جمعیت، پیشگویی‌هایی را فراهم می‌کند. تمرکز اولیه اکثر پژوهش‌ها بر روی پارامترهای جمعیت تحت مطالعه است. نمونه و آماره‌های توصیف‌کننده آن تنها تا جایی اهمیت دارند که اطلاعاتی را در مورد پارامترهای جمعیت فراهم کنند. بنابراین، یک جنبه مهم استنباط آماری، گزارش از صحت احتمالی یا درجه اطمینان از آماره نمونه‌ای است که مقدار پارامتر جمعیت را پیشگویی می‌کند (آگرستی و فینلای، ۱۹۹۷).

همان‌طور که ذکر شد از آمار استنباطی برای استنتاج الگوها و خصوصیات جمعیت از خصوصیات نمونه استفاده می‌شود. رایج‌ترین تکنیک‌های استنباطی در تحلیل یک متغیره، برآوردهای فاصله‌ای (خطای استاندارد) و در تحلیل دو یا چندمتغیره بهره‌گیری از سطح معنی‌داری (سطح خطا) است.

۱) برآوردهای فاصله‌ای یا خطای استاندارد

اگر میانگین درآمد در نمونه‌ای ۱۸۰۰۰ دلار باشد میانگین درآمد در جمعیت چقدر خواهد بود؟ چون بعید است که نمونه بازتاب کامل جمعیت باشد (خطای نمونه‌گیری)، نمی‌توان فقط با استفاده از میانگین نمونه (که برآورد نمونه خوانده می‌شود) میانگین واقعی درآمد جمعیت را (که پارامتر جمعیت خوانده می‌شود) پیدا کرد.

لازم است راهی برای برآورد میزان دقت احتمالی برآورد نمونه پیدا کرد. چنانچه نمونه ما نمونه‌ای تصادفی باشد نظریه احتمالات چنین راه‌حلی در اختیار ما قرار می‌دهد. اگر تعداد زیادی نمونه تصادفی مختلف داشته باشیم، اکثر آن‌ها تقریباً نزدیک به پارامتر جمعیت خواهند بود: میانگین اکثر آن‌ها نزدیک به میانگین واقعی جمعیت خواهد بود و فقط معدودی از آن‌ها تفاوت زیادی با میانگین جمعیت خواهند داشت. در واقع توزیع برآورد نمونه‌ها به توزیع نرمال نزدیک می‌شود.

مسئله این جاست که اگر فقط یک نمونه داشته باشیم (که دارای تعدادی عضو است) چگونه بدانیم میانگین نمونه ما چقدر به میانگین حقیقی جمعیت نزدیک است. ممکن است نمونه ما یکی از نمونه‌های پرت باشد، اما بیشتر محتمل است که یکی از نمونه‌های نزدیک به میانگین حقیقی جمعیت باشد (چون اکثر نمونه‌ها نزدیک به آن هستند). اما حتی در این صورت تا چه حد به آن نزدیک است؟ برای برآورد میزان نزدیکی میانگین نمونه به میانگین جمعیت از آماره‌ای به نام خطای استاندارد میانگین استفاده می‌کنیم که فرمول آن بدین شرح است:

$$S_m = \frac{S}{\sqrt{N}}$$

که در آن S_m خطای استاندارد میانگین، S انحراف استاندارد و N تعداد کل نمونه است.

پس از محاسبه خطای استاندارد نظریه احتمالات به ما می‌گوید که میانگین جمعیت احتمالاً در چه دامنه‌ای از میانگین نمونه قرار دارد. طبق نظریه احتمالات در ۹۵ درصد نمونه‌ها، میانگین جمعیت در محدوده ± 2 واحد خطای استاندارد از میانگین نمونه قرار دارد. به عبارت دیگر به احتمال ۹۵ درصد میانگین جمعیت در محدوده ± 2 خطای استاندارد از میانگین نمونه ما قرار دارد. بنابراین می‌توانیم برآورد کنیم که میانگین جمعیت در چه دامنه‌ای قرار دارد. به این دامنه فاصله اطمینان گفته می‌شود و میزان قطعیت قرار گرفتن میانگین جمعیت در این فاصله، سطح اطمینان خوانده می‌شود. رقمی که برای خطای استاندارد به دست می‌آوریم همواره برحسب واحد متغیر مورد بررسی است. مثلاً اگر متغیر درآمد بر حسب دلار باشد و خطای استاندارد ۱۰۰۰ باشد، بدان معناست که خطای استاندارد ۱۰۰۰ دلار است. معنای همه این‌ها به زبان ساده چیست؟ گیریم میانگین درآمد در نمونه ما ۱۸۰۰۰ دلار است و خطای استاندارد آن

۱۰۰۰ دلار. در نتیجه به احتمال ۹۵ درصد میانگین درآمد جمعیت در محدوده ۱۶۰۰۰ تا ۲۰۰۰۰ دلار است (یعنی ۱۸۰۰۰ به اضافه و منهای دو خطای استاندارد).

اندازه خطای استاندارد تابع حجم نمونه است. بنابراین برای این که بتوان میانگین جمعیت را در محدوده کوچکتری برآورد کرد؛ یعنی برای این که بتوان فاصله اطمینان را کاهش داد باید خطای استاندارد را کم کرد. بدین منظور باید بر حجم نمونه افزود: با چهار برابر شدن حجم نمونه، خطای استاندارد نصف می‌شود (دواس، ۱۳۷۶: ۱۵۳-۱۵۴).

۲) سطح معنی داری

بعد از این که مشخص شد دو متغیر با هم رابطه دارند باید دید رابطه‌ای که در نمونه مشاهده شده است به همان شدت در جمعیتی که نمونه از آن گرفته شده است نیز وجود دارد یا خیر: باید دید می‌توان نتایج حاصله از نمونه را با اطمینان تعمیم داد یا خیر. بدین منظور از آمار استنباطی (آزمون‌های معنی داری) سود می‌جوییم.

منطقی که در پس این آمارها نهفته است چیست؟ معمولاً ابتدا فرض می‌کنیم در جمعیت بین دو متغیر رابطه‌ای وجود ندارد (یعنی $I = 0$ = همبستگی). این فرضیه را فرضیه صفر می‌نامند. چنانچه مشاهده کردیم در نمونه بین دو متغیر رابطه‌ای وجود دارد (مثلاً $I = .4$) تفاوت فرض عدم رابطه در جمعیت با رابطه مشاهده شده در نمونه به دو صورت قابل تفسیر است:

۱- نمونه ما نمونه معرف نیست. چه بسا باوجود استفاده از تکنیک‌های نمونه‌گیری تصادفی، بازهم نمونه ما نمونه‌ای غیر معرف باشد. این را خطای نمونه‌گیری می‌نامند. بدین ترتیب ممکن است تفاوت بین نتیجه حاصل از نمونه ($I = .4$) و فرض عدم وجود رابطه در جمعیت ($I = 0$) این باشد که نمونه، معرف خوبی از جمعیت نیست.

۲- فرض عدم رابطه در جمعیت فرض غلطی است (فرض وجود رابطه در جمعیت صحیح است).

اگر تفسیر اول را بپذیریم و حاصل نمونه ($I = .04$) را ناشی از خطای نمونه‌گیری بدانیم، آن‌گاه می‌گوییم بعید است که رابطه بین دو متغیر در جمعیت از صفر تجاوز کند (یعنی فرض عدم رابطه را قبول می‌کنیم). از آن جا که معمولا هدف از مطالعه نمونه، استنتاج درباره جمعیت است، فرض را بر آن می‌گیریم که رابطه مشاهده شده در نمونه ($I = .04$)، پیامد خطای نمونه‌گیری است و در نتیجه آن را معادل عدم رابطه در جمعیت ($I = 0$) می‌گیریم.

اما اگر تفسیر دوم را بپذیریم و فرض اولیه عدم رابطه در جمعیت را غلط بدانیم، نتیجه حاصل از نمونه را بازتاب رابطه واقعا موجود در جمعیت تفسیر می‌کنیم و چنان عمل می‌کنیم که گویی رابطه مشاهده شده در نمونه واقعی است و محصول نمونه غیر معرف نیست. به بیان دیگر می‌توانیم وجود رابطه در جمعیت را با توجه به وجود رابطه در نمونه بپذیریم.

آزمون‌های معنی‌داری

با کمک آزمون‌های معنی‌داری می‌توان پی برد که کدام تفسیر درست است. منطق این آزمون‌ها ساده است. اگر دو متغیر در جمعیت فاقد رابطه باشند احتمال این که نمونه تصادفی ما بیان‌گر رابطه‌ای بین این دو متغیر باشد چقدر است؟ (یعنی احتمال دقیق نبودن نمونه تصادفی ما چقدر است؟). به عنوان مثال اگر صدبار نمونه‌گیری تصادفی انجام دهیم، احتمال این که یکی از آن‌ها نمونه غیر معرفی باشد، یعنی بیان‌گر رابطه‌ای باشد که واقعا در جمعیت وجود ندارد چقدر است؟ معمولا گفته می‌شود آنجا که احتمالا از هر صد نمونه بیش از پنج نمونه بیان‌گر رابطه‌ای باشند که ناشی از خطای نمونه‌گیری است، احتمال نادرست بودن نمونه بالاست. چه بسا نمونه خاص ما یکی از پنج نمونه باشد! در نتیجه باید گفت به احتمال زیاد رابطه مشاهده‌شده ناشی از خطای نمونه‌گیری است و فرض فقدان رابطه در جمعیت واقعی صحیح است.

گروهی از پژوهش‌گران محتاط‌ترند و معتقدند آنجا که بیش از یکی از صد نمونه بتواند برحسب تصادف رابطه‌ای به شدت رابطه مشاهده شده ایجاد کند، آن‌گاه احتمال خطا زیاد است. اما اگر دریابیم که صرفا شمار بسیار اندکی از نمونه‌ها ممکن است رابطه مشاهده شده را ایجاد کنند می‌توانیم قبول کنیم که رابطه مشاهده شده در نمونه ما واقعی و منعکس‌کننده رابطه در جمعیت است.

از آن جا که هرگز صدبار نمونه‌گیری نمی‌کنیم، باید برآورد کنیم که اگر صد بار نمونه‌گیری می‌کردیم چقدر احتمال داشت که نمونه ما جزء یکی از نمونه‌هایی باشد که صرفاً بر اساس شانس و تصادف بیان‌گر رابطه‌ای به شدت رابطه مشاهده شده در نمونه ما هستند.

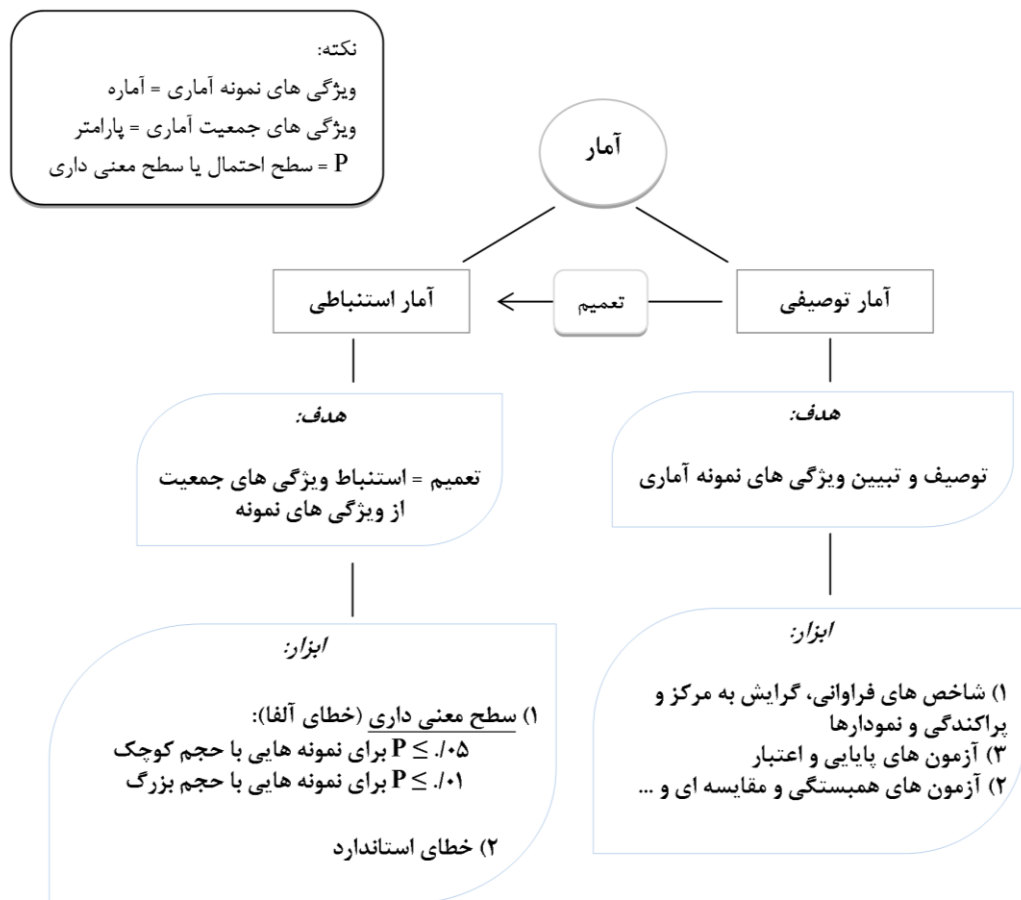
با کمک نظریه احتمالات می‌توان احتمال واقعی نبودن رابطه مشاهده شده در نمونه را (یعنی احتمال این که ناشی از خطای نمونه‌گیری باشد) برآورد کرد. (در این جا به این نظریه نمی‌پردازیم و فقط متذکر می‌شویم که فرض این نظریه این است که از نمونه‌های تصادفی استفاده می‌کنیم). آزمون معنی‌داری آماری در واقع برآورد همین احتمال است. دامنه مقادیر این آزمون‌ها از (۰) تا (۱) است و آن‌ها را سطوح معنی‌داری می‌خوانند. معنای این ارقام چیست؟ گیریم سطح معنی‌داری ۰/۵۰ است. این بدان معناست که در ۵۰ نمونه از ۱۰۰ نمونه فقط بر اثر خطای نمونه‌گیری (شانس) رابطه‌ای به قوت رابطه‌ای که ما در نمونه مشاهده کردیم دیده می‌شود. در این صورت احتیاط حکم می‌کند رابطه مشاهده شده در نمونه خود را به احتمال زیاد واقعی ندانیم و در نتیجه فرض عدم رابطه در جمعیت رد نمی‌شود (تأیید می‌شود).

اگر سطح معنی‌داری ۰/۰۵ باشد بدان معناست که فقط پنج نمونه از هر ۱۰۰ نمونه برحسب تصادف به رابطه مشاهده شده در نمونه ما منجر می‌شود. اگر سطح معنی‌داری ۰/۰۱ باشد به معنی واقعی نبودن رابطه در یک نمونه از هر ۱۰۰ نمونه و اگر ۰/۰۰۱ باشد به معنای یک نمونه در ۱۰۰۰ نمونه است. پیداست هر چه سطح معنی‌داری پایین‌تر باشد، می‌توان اطمینان بیشتری به «واقعی» بودن رابطه مشاهده شده در نمونه داشت.

در اینجا به نحوه محاسبه این آزمون‌های آماری معنی‌داری نمی‌پردازیم. فرمول این‌ها در کتاب‌های آماری وجود دارد و با برنامه‌های کامپیوتری هم به راحتی می‌توان آن‌ها را حساب کرد. اما مسأله‌ای وجود دارد که برنامه‌های کامپیوتری پاسخگوی آن نیستند. اکثر برنامه‌های کامپیوتری سطح معنی‌داری را بین ۰/۰۱ تا ۱/۰۰ محاسبه می‌کنند؛ اما در چه سطحی فرضیه عدم رابطه در جمعیت (فرضیه صفر) رد می‌شود؟ معمولاً سطح معنی‌داری را ۰/۰۵ تا ۰/۰۱ را به عنوان مبنا در نظر می‌گیرند. اما این سطوح قراردادی و اختیاری هستند.

مسأله‌ای که در کاربرد سطح معنی داری ۰/۰۵ وجود دارد، سهولت رد فرض صفر (عدم رابطه) است: چه بسا فرض عدم رابطه در جمعیت رد شود (فرض وجود رابطه تأیید شود) در حالی که واقعا رابطه‌ای وجود نداشته باشد. چنین اشتباهی را **خطای نوع اول** می‌خوانند و بیشتر در نمونه‌های بزرگ پیش می‌آید. از این رو بهتر است در نمونه‌های بزرگ سطح معنی‌داری ۰/۰۱ را به عنوان مبنا در نظر بگیریم. اما اگر همواره از سطح ۰/۰۱ استفاده کنیم ممکن است کار به **خطای نوع دوم** بکشد - یعنی سخت‌گیری بیش از حد و تصدیق فرضیه صفر در جایی که باید آن را رد کرد. احتمال چنین خطایی در نمونه‌های کوچک بیشتر است. طبق قاعده تجربی برای نمونه‌های کوچک از سطح معنی‌داری ۰/۰۵ و برای نمونه‌های بزرگ از سطح ۰/۰۱ یا کمتر استفاده می‌کنیم (دواس، ۱۳۷۶: ۱۹۰-۱۹۲).

☑ نکته: سطح معنی‌داری در برنامه SPSS به صورت Sig گزارش می‌شود که در هنگام گزارش نتایج در پایان‌نامه‌ها و مقالات باید به صورت مقدار P (P-Value) گزارش شود.



فرضیه

در روش تحقیق، فرضیه عبارت است از راه حل پیشنهادی پژوهشگر برای پاسخ به مسأله. به عبارت دیگر، ریشه یک فرضیه مناسب با انتخاب و بیان مسأله در هم آمیخته است. فرضیه ابزار نیرومندی است که پژوهشگر را قادر می‌سازد تا نظریه را به مشاهده و مشاهده را به نظریه ربط دهد. فرضیه یک قضیه شرطی یا فرضی است که تأیید یا رد آن باید بر اساس سازگاری مفاهیم آن و به استناد مدارک تجربی و دانش گذشته، آزمایش شود. فرضیه همانند نورافکن پر قدرتی است که راه را برای پژوهشگر روشن می‌کند (دلاور، ۱۳۹۰: ۶۱).

فرضیه راهنمای پژوهش است چرا که وظیفه پژوهشگر را مشخص می‌کند که چه کاری را انجام دهد و به دنبال چه چیزی بگردد. اگر مسأله پژوهش وظیفه پژوهشگر را بیان می‌کند، فرضیه چگونگی انجام آن را روشن می‌سازد. مسأله مشخص می‌کند که پژوهشگر چه کار باید انجام دهد اما فرضیه مشخص می‌کند که این کار چگونه انجام گیرد. به عبارتی دیگر مسأله جا را نشان می‌دهد و فرضیه راه را.

فرضیه‌ها قوی‌ترین ابزار پژوهش هستند. از این رو اگر بخواهیم فرضیه‌ها را بررسی کنیم لازم است که برای اندازه‌گیری مفاهیم موجود در فرضیه‌ها از مقیاس‌های متناسب (اسمی، رتبه‌ای و فاصله‌ای یا نسبی) استفاده کرد. شیوه بیان و نگارش هر فرضیه نوع آزمون آماری متناسب با آن را مشخص می‌کند. در نتیجه قبل از بیان فرضیه بهتر است با انواع فرضیه آشنا شویم.

انواع فرضیه: فرضیه‌ها انواع مختلفی دارند و به شیوه‌های مختلف می‌توان فرضیه‌ها را طبقه‌بندی کرد. دو نوع از تقسیم‌بندی فرضیه‌ها در ادامه بیان شده است.

۱) فرضیه مقایسه‌ای، فرضیه رابطه‌ای و فرضیه علی

۱- **فرضیه مقایسه‌ای:** در این فرضیه‌ها ما به دنبال بررسی و مقایسه تفاوت اثر دو یا چند متغیر بر یک یا چند متغیر دیگر هستیم. در این حالت فرضیه به صورتی بیان می‌شود که تفاوت‌ها حدس زده می‌شوند و مقایسه‌ای انجام می‌گیرد. برای مثال می‌توان فرض نمود که: «میانگین نمرات دانشجویانی که با

روش "الف" آموزش داده می‌شوند از میانگین نمرات دانشجویانی که با روش "ب" آموزش می‌بینند بیشتر است». بدین ترتیب بین نمرات دانشجویان روش الف و ب مقایسه انجام می‌گیرد. یا می‌توان فرض کرد که «میزان رضایت شغلی در زنان و مردان شاغل متفاوت است» عمدتاً در پژوهش‌هایی که دارای گروه تجربی (گروه مورد آزمایش) و گروه شاهد (کنترل) می‌باشند از این نوع فرضیات استفاده می‌کنند.

آزمون فرضیه‌های تفاوتی، با توجه به نوع مقیاس داده‌ها، هم از طریق آزمون‌های پارامتری (مانند آزمون‌های t و آزمون تحلیل واریانس) و هم از طریق آزمون‌های ناپارامتری (مثل ویلکاکسون، مان ویتنی، کروسکال والیس و ...) انجام می‌گیرد.

۲- **فرضیه رابطه‌ای:** در این حالت پژوهش‌گر در پی مطالعه و بررسی میزان همبستگی و جهت آن بین دو یا چند متغیر است. پژوهش‌گر قصد دارد که صرفاً درجه و جهت همبستگی متغیرهای مورد مطالعه را کشف کند و نه رابطه علت و معلولی بین آن‌ها را. برای مثال می‌توان فرضیه‌های زیر را عنوان کرد «بین میزان تحصیلات و گرایش افراد به خودکشی رابطه وجود دارد» یا «بین میزان ورزش و میزان افسردگی زنان رابطه وجود دارد». در پژوهش‌ها با این گونه فرضیه‌ها، نتایج به دست آمده از آزمون‌های آماری می‌تواند از مقدار α تا ۱ (با جهت مثبت و منفی) باشد.

در این حالت برای سنجش رابطه معمولاً از آزمون‌های پیوستگی یا همبستگی مانند مجذور کای استقلال، پیرسون و اسپیرمن و کندانال استفاده می‌شود.

۳- **فرضیه علی:** هدف در فرضیه‌های علی، کشف و تعیین رابطه علت و معلولی دو یا چند متغیر است. در این جا هدف صرفاً تعیین ارتباط و همبستگی دو یا چند متغیر نیست، بلکه می‌خواهیم عمیق‌تر و ریشه‌ای‌تر با آن برخورد کرده و بگوییم که متغیری علت به وجود آمدن متغیر دیگر است. برای مثال می‌توان فرض کرد که: «هوش دانشجویان ممتاز علت پیشرفت تحصیلی آنان است» یا این که «پایگاه اقتصادی- اجتماعی افراد بر میزان اعتماد اجتماعی آن‌ها تأثیر دارد» و یا «مصرف داروی X موجب کاهش فشار خون می‌شود».

برای سنجش فرضیه‌های علی با توجه به نوع داده‌ها (پارامتری و ناپارامتری) از آزمون‌های متناسب استفاده می‌شود. معمولاً برای داده‌های ناپارامتری از آزمون مجذور کای استقلال استفاده می‌شود و برای داده‌های پارامتری از آزمون‌های رگرسیون (دو یا چند متغیره)، تکنیک تحلیل مسیر و تکنیک مدل‌یابی معادلات ساختاری استفاده می‌شود.

۲) فرضیه پژوهشی و فرضیه آماری

در تقسیم‌بندی دیگر، فرضیه‌ها به دو دسته فرضیه پژوهشی و آماری تقسیم می‌شوند:

۱. **فرضیه آماری:** فرضیه آماری جمله یا عبارتی است که پیرامون ویژگی‌های جامعه بیان می‌شود و امکان دارد درست نباشد، ولی پژوهش‌گر صرفاً به خاطر برقرار کردن یک شرایط قابل آزمایش، آن را مطرح می‌کند. به عبارتی، فرضیه آماری یک بیان کمی درباره پارامتر جامعه است که با استفاده از نمادهای آماری نوشته می‌شود و نقش آن، هدایت پژوهش‌گر در انتخاب آزمون آماری است. فرضیه آماری به دو صورت فرضیه صفر و فرضیه خلاف بیان می‌شود.

الف) فرضیه صفر: فرضیه صفر که با علامت H_0 (اچ صفر) نشان داده می‌شود، فرضی است که پژوهش‌گر مایل به رد کردن آن می‌باشد. بنابراین، فرض صفر بیان‌گر عدم ارتباط بین متغیرها، عدم تفاوت یک متغیر در بین یک یا چند گروه و عدم تأثیر یک متغیر بر متغیر دیگر می‌باشد. در پژوهش بیان فرضیه صفر ضروری نیست.

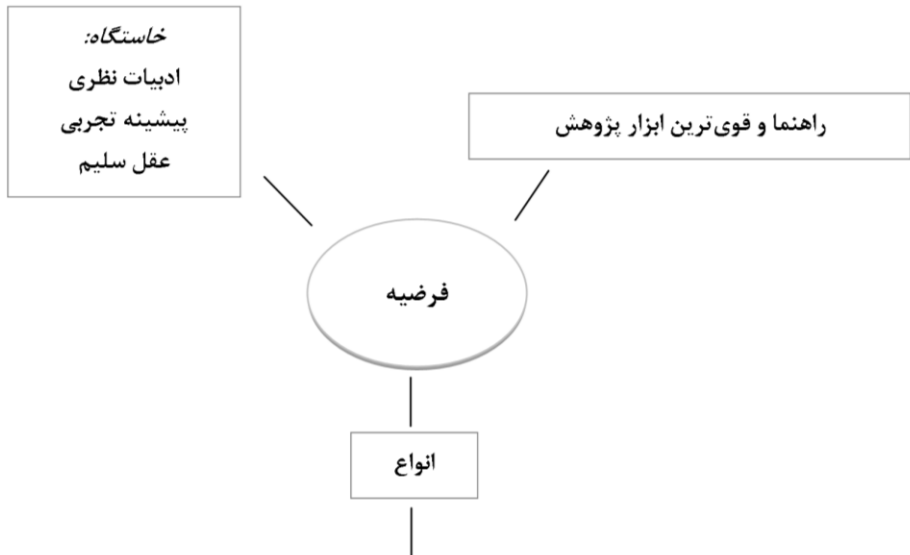
ب) فرض خلاف یا بدیل: این فرض با علامت H_1 (اچ یک) یا H_A نشان داده می‌شود. فرض خلاف غالباً منطبق با فرضیه پژوهشی است. به بیان دیگر، این فرضیه بیان‌کننده انتظار پژوهش‌گر درباره نتایج پژوهش است.

۲. **فرضیه پژوهشی:** فرضیه‌ای که در پژوهش بیان و گزارش می‌شود فرضیه پژوهشی است اما فرضیه‌ای که در پژوهش مورد آزمون قرار می‌گیرد فرضیه آماری است. به عنوان مثال فرضیه‌هایی مانند: «بین میزان دین‌داری و گرایش به خودکشی رابطه منفی وجود دارد.» یا فرضیه «گرایش به فرزندآوری در بین روستائینان بیشتر از شهرنشینان است»؛ فرضیه پژوهشی هستند و در متن پژوهش گزارش می‌شوند.

در جدول ۳-۱ به مقایسه شیوه بیان فرضیه‌های آماری (فرضیه صفر و خلاف) و فرضیه پژوهشی پرداخته‌ایم:

جدول ۳-۱- مقایسه شیوه ارائه فرضیه‌های آماری و فرضیه پژوهشی

فرضیه پژوهشی	فرض صفر	فرض خلاف
بین بهره هوشی و معدل دانش آموزان رابطه وجود دارد	$r = 0$	$r \neq 0$
بین بهره هوشی و معدل دانش آموزان رابطه مثبت وجود دارد	$r \leq 0$	$r > 0$
میزان بهره هوشی دانش آموزان دختر و پسر متفاوت است	$\mu_A = \mu_B$	$\mu_A \neq \mu_B$



تقسیم بندی بر اساس ماهیت:

۱. فرضیه مقایسه ای: میزان دینداری در مردان و زنان متفاوت است (یا زنان از مردان دیندارترند).

آزمون های آماری: آزمون های t، آزمون تحلیل واریانس، مان ویتنی و....

۲. فرضیه رابطه ای: بین میزان تحصیلات و میزان دینداری افراد رابطه وجود دارد.

آزمون های آماری: آزمون های همبستگی پیرسون، اسپیرمن، کندال، کای اسکوتر استقلال، فی و وی کرامر

۳. فرضیه علی: دین داری منجر به کاهش ارتکاب جرم می شود (دین داری بر میزان ارتکاب جرم موثر است)

آزمون های آماری: رگرسیون، تحلیل مسیر، مدل یابی معادلات ساختاری

تقسیم بندی بر اساس شیوه ارائه:

۱. فرضیه پژوهشی: در پژوهش بیان و گزارش می شود

۲. فرضیه آماری: بیان فرضیه با نماد های آماری. هدایت در انتخاب آزمون های آماری. بیان فرضیه آماری در پژوهش

ضروری نیست.

الف) فرضیه صفر یا H_0 : فرضیه ای که در عمل آزمون می شود. پژوهش مایل به رد کردن فرضیه صفر است

الف) فرضیه خلاف یا بدیل یا H_1 : منطبق با فرضیه پژوهشی. انتظار پژوهشگر از نتایج پژوهش را می رساند.

انواع متغیرها

متغیر، 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. به بیان دیگر متغیر یک نماد است که می‌توان عدد یا ارزش را جایگزین کرد. به عنوان مثال سن افراد، میزان تحصیلات، وزن و قد افراد، رضایت شغلی، استرس، اعتماد اجتماعی، تعهد سازمانی و ... متغیر محسوب می‌شوند. متغیرها بر اساس نقشی که در پژوهش دارند به انواع مختلفی تقسیم می‌شوند که مهم‌ترین آن‌ها عبارتند از: متغیر مستقل، وابسته، میانجی، تعدیل‌کننده و متغیر کنترل.

(۱) متغیر مستقل

متغیر مستقل، متغیری است که بر روی متغیر دیگر تأثیر می‌گذارد. بنابراین متغیر مستقل، متغیر علت می‌باشد و به طور طبیعی در مطالعه رابطه متغیرها، متغیر مستقل از قبل وجود دارد که با تغییرات خود موجب تغییر در متغیری می‌شود که وابسته نام دارد. پس:

الف- متغیر مستقل، متغیری است که از قبل وجود دارد (قبل از متغیر وابسته) مثلاً در مطالعه نقش جنسیت در میزان علاقه به فوتبال، جنسیت که از قبل وجود دارد متغیر مستقل نام دارد. یا در مطالعه تأثیر روش تدریس خاص در پیشرفت تحصیلی همواره روش تدریس خاص اعمال می‌شود، پس اثر آن در پیشرفت مشاهده می‌شود و این روش تدریس متغیر مستقل است.

ب- متغیر مستقل، در فرضیه نقش علت را دارد، بنابراین نقش اثرگذار را دارد.

پ- متغیر مستقل در آزمایش‌ها، متغیری است که در کنترل پژوهش‌گر است، می‌تواند آن را وارد فرآیند آزمایش کند یا از آزمایش خارج کند و یا مقدار آن را تغییر دهد.

ت- متغیر مستقل، متغیری است که پژوهش‌گر از قبل با آن و تغییرات آن آشنایی دارد.

ث- متغیر مستقل، نقش معین‌کننده دارد.

ج- متغیر مستقل، در فرآیند محرک و پاسخ، نقش محرک را دارد.

چ- در آخر این که معمولاً علامت آن در معادلات (به ویژه معادلات رگرسیونی) به صورت X است. در رگرسیون متغیر مستقل را متغیر پیش‌بین هم می‌نامند.

۲) متغیر وابسته

معمولاً در پژوهش‌ها، مطالعه پژوهش‌گر بر روی متغیر وابسته متمرکز است. به عبارت دیگر در عمل و در مرحله برخورد با مسأله، پژوهش‌گر مشاهده می‌کند که تغییری تغییر می‌کند، اما نمی‌داند علت این تغییر یا علت اصلی این تغییر چیست و نیز قادر نیست این تغییرات را پیش‌بینی کند. بنابراین متغیر وابسته تغییری است که تغییرات آن وابسته به متغیر (متغیرهای) مستقل می‌باشد و حالت معلولی دارد و پژوهش‌گر سعی می‌کند ابتدا در مطالعه مقدماتی با محور قراردادن متغیر وابسته، متغیر یا متغیرهای مستقل را شناسایی نموده، سپس با دست‌کاری متغیر مستقل، تغییرات متغیر وابسته را بررسی نماید. به طور خلاصه:

الف- معمولاً در هر پژوهشی، پژوهش‌گر با یک متغیر وابسته و حداقل یک متغیر مستقل سروکار دارد. احتمال این که ما چند متغیر مستقل داشته باشیم زیاد است ولی معمولاً متغیر وابسته (نهایی) یکی است. اگر در پژوهشی چندین متغیر وابسته وجود داشته باشد، پژوهش‌گر عملاً با چندین پژوهش مستقل سروکار خواهد داشت.

ب- متغیر وابسته در فرضیه نقش معلول را دارد بنابراین اثرپذیر است و زمانی تغییر می‌کند که متغیر مستقل در فرآیند آزمایش (یا به طور خودبه‌خودی در پژوهش‌های زمینه‌یابی یا پیمایشی) نقش خود را ایفا کرده باشد. البته در همه پژوهش‌ها نقش علت و معلولی با این قاطعیت بیان نمی‌شود، بلکه جریان پژوهش و شیوه مطالعه طوری است که پژوهش‌گران می‌کوشند به جای رابطه علی از وجود رابطه صحبت کنند و حتی زمانی که مشخصه‌های آماری به گونه‌ای انتخاب می‌شوند که این ارتباط حالت تداعی پیدا کند، با این همه اگر ارتباط و تداعی را با علامت پیکان (فلش) نمایش دهیم در اغلب موارد این پیکان یک‌سویه و از متغیر مستقل به متغیر وابسته است.

پ- متغیر وابسته در آزمایش‌ها، تغییری است که تغییرات آن در کنترل پژوهش‌گر نیست، بلکه در گروه تغییرات متغیر مستقل است، پژوهش‌گر نمی‌تواند به طور مستقیم متغیر وابسته را حذف کند، یا دست‌کاری کند و یا جهت آن را تغییر دهد.

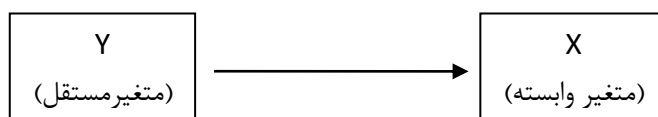
ت- پژوهش‌گر از آغاز با تغییرات متغیر وابسته آشنا نیست، بلکه شرایط پژوهش را از آن روی فراهم کرده است که این تغییرات را مطالعه کند.

ث- متغیر وابسته در پژوهش‌ها نقش پیش‌بینی‌شونده را دارد.

ج- متغیر وابسته در فرآیند محرک و پاسخ، نقش پاسخ یا برون‌داد را برعهده دارد.

چ- معمولاً متغیر وابسته در عنوان پژوهش می‌آید. مثلاً در عنوان پژوهش «بررسی میزان اعتماد اجتماعی در بین مردم شهر قم و عوامل مؤثر بر آن»، اعتماد اجتماعی متغیر وابسته است.

ح- معمولاً علامت آن در معادلات (به ویژه معادلات رگرسیونی) بصورت Y است. در رگرسیون متغیر وابسته را متغیر ملاک هم می‌نامند.

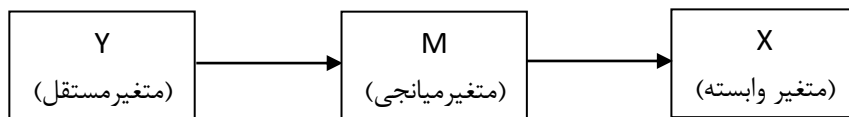


شکل ۱-۲- نحوه ارتباط متغیرهای مستقل و وابسته

۳) متغیر میانجی (واسط)

میانجی‌گری یا اثر غیرمستقیم زمانی رخ می‌دهد که اثر یک متغیر مستقل بر متغیر وابسته از طریق متغیر میانجی منتقل شود. متغیر میانجی متغیری است که بین دو متغیر دیگر قرار می‌گیرد و موجب ارتباط غیرمستقیم آن‌ها با یکدیگر می‌شود. متغیر میانجی از متغیر پیشین خود (متغیر مستقل) اثر پذیرفته و بر متغیر پسین خود (متغیر وابسته) اثر می‌گذارد. فرض کنیم که میزان تحصیلات کارمندان بر میزان حقوق دریافتی آنان مؤثر است. در این حالت اثر میزان تحصیلات بر میزان حقوق مستقیم و بدون واسطه است. اما می‌توانیم فرض کنیم که تأثیر میزان تحصیلات بر میزان درآمد، هم به صورت مستقیم و هم به صورت غیرمستقیم از طریق جایگاه شغلی است. با این استدلال که افزایش میزان تحصیلات با بهبود و ارتقاء جایگاه شغلی کارمندان همراه است که این ارتقاء جایگاه شغلی خود موجب افزایش حقوق کارمندان می‌شود. در این حالت متغیر جایگاه شغلی یک متغیر میانجی است. تعداد متغیرهای میانجی می‌تواند یک متغیر یا بیشتر باشد.

متغیر میانجی زمانی مورد توجه قرار می‌گیرد که بین متغیر مستقل و وابسته رابطه بسیار قوی و بالا باشد (بارون و کنی^۳، ۱۹۸۶). متغیر میانجی را با علامت M نشان می‌دهند.



شکل ۱-۳- نقش متغیر میانجی در پژوهش

۴) متغیر تعدیل‌کننده:

متغیر تعدیل‌کننده، متغیر مستقلی است که نقش ثانویه دارد و پژوهش‌گر مایل است اثر آن را در فرآیند آزمون فرضیه در کنار متغیر مستقل مطالعه کند. متغیر تعدیل‌کننده، متغیری است که بر رابطه متغیر مستقل و متغیر وابسته تأثیر اقتصادی دارد. یعنی، حضور متغیر سوم (متغیر تعدیل‌کننده)، رابط مورد انتظار اصلی بین متغیرهای مستقل و وابسته، را تغییر می‌دهد. متغیر تعدیل‌گر متغیری است که جهت و شدت ارتباط بین متغیر مستقل و وابسته را تحت تأثیر قرار می‌دهد. متغیر تعدیل‌کننده زمانی مورد توجه قرار می‌گیرد که بین رابطه بین متغیر مستقل و وابسته از حد انتظار پایین‌تر بوده و یا ناسازگار (جهت رابطه عکس انتظار ما باشد) است (بارون و کنی، ۱۹۸۶). در صورت‌بندی فرضیه از متغیر تعدیل‌کننده همانند متغیرهای مستقل و وابسته نام برده می‌شود. متغیر تعدیل‌کننده را معمولاً با علامت Z نشان می‌دهند.

نقش این متغیر را در مثال‌های زیر بهتر می‌توان درک کرد. فرض کنید پژوهش‌گری در زمینه رضایت شغلی افراد مطالعه کرده است. او فرضیه‌ای تدوین می‌کند که: «رضایت شغلی افراد با میزان درآمد آنان رابطه دارد». مسلماً برای هر پژوهش‌گری این سوال مطرح می‌شود که آیا میزان رضایت شغلی افراد شاغل که ناشی از افزایش درآمد ماهیانه است در میان شاغلان زن و مرد تفاوتی ندارد؟ یعنی درست است که می‌گوییم (به عنوان فرضیه) با افزایش درآمد، رضایت شغلی هم بیشتر می‌شود، اما آیا تأثیر این اقدام در شاغلان زن و مرد یکسان است؟ ممکن است اهمیت میزان درآمد در بین مردان (به دلیل نقشی که در تامین مخارج خانواده دارند) بیشتر از زنان باشد و افزایش یکسان درآمد یک گروه از شاغلان زن و مرد، موجب افزایش یکسانی در رضایت شغلی آنان

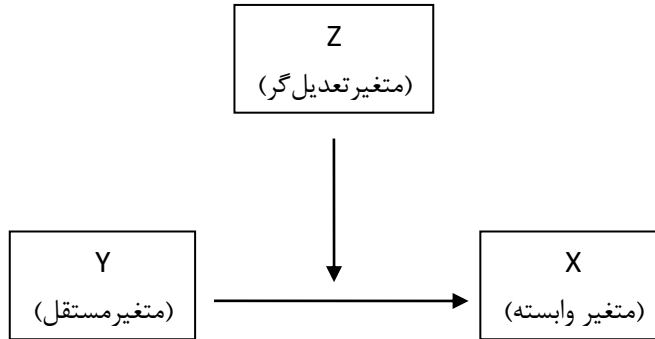
³ Baron and Kenny

نشود، چرا که ممکن است در زنان عوامل مهم‌تری (سالم بودن محیط کاری، ساعات کاری، میزان مرخصی و ...) وجود داشته باشند که بیشتر از درآمد، بر رضایت شغلی موثر باشند. در این جا جنس افراد متغیر تعدیل‌کننده است که بر شدت رابطه بین درآمد و رضایت شغلی اثر می‌گذارند.

مثال دیگری در این زمینه می‌تواند در فهم بهتر متغیر تعدیل‌کننده به ما کمک کند. یکی از تئوری‌های جدید مدعی است که تنوع نیروی کار (مشمول بر مبانی اخلاقی، نژادها و ملیت‌های مختلف) در اثربخشی سازمانی نقش ایفا خواهد کرد زیرا هر گروه، مهارت‌ها و تخصص‌های فنی خاصی را با خود به محیط کاری می‌آورد. این هم‌افزایی می‌تواند پژوهش‌گر شود، فقط اگر مدیران بدانند چگونه استعدادهای خاص گروه کاری متنوع را به کار ببندند. در این مثال، اثربخشی سازمانی متغیر وابسته است که به وسیله نیروی کاری متنوع به طور مثبت تحت تأثیر قرار می‌گیرد، که در واقع متغیر مستقل است. درعین حال، برای استفاده از نیروهای بالقوه، مدیران باید بدانند چگونه استعدادهای گروه‌های مختلف را برای اثربخش ساختن امور، تشویق و هماهنگ سازند. اگر ندانند، هم‌افزایی پژوهش‌گر نخواهد شد. به عبارت دیگر، بهره‌برداری اثربخش از استعدادها، دیدگاه‌های مختلف و توانایی‌های حل مسأله برای افزایش اثربخشی سازمانی متکی و وابسته به مهارت مدیران در عمل کردن به عنوان، یک عامل تسریع‌کننده است. این مهارت فنی مدیریتی به صورت متغیر تعدیل‌کننده در می‌آید.

مثال دیگری که می‌توان زد این است که می‌دانیم در مشاغل دولتی بین میزان تحصیلات و میزان حقوق دریافتی کارمندان رابطه مثبت وجود دارد، بدین صورت که با افزایش میزان تحصیلات افراد، حقوق دریافتی آنان هم افزایش می‌یابد. اما وجود روند مثبت و افزایشی بین تحصیلات و حقوق در مشاغل دولتی، در بین همه مشاغل دولتی یکسان نیست. ممکن است که افزایش میزان تحصیلات در بین معلمان در مقایسه با سایر کارمندان ادارات دولتی تفاوت داشته باشد و افزایش تحصیلات معلمان موجب شود تا حقوق کمتری در مقایسه با کارمندان ادارات دولتی بگیرند. همچنین در ارگان‌ها و سازمان‌های مختلف هم این افزایش حقوق متفاوت است، مثلاً ممکن است در ارگان‌های مرتبط با وزارت بهداشت یا وزارت نفت، افزایش میزان تحصیلات منجر شود که حقوق این دسته از کارمندان افزایش بیشتری (به عنوان مثال در مقایسه با کارمندان وزارت

آموزش و پرورش) داشته باشد. متغیر نوع شغل یا نوع ارگان دولتی کارمندان، نقش تعدیل‌کنندگی را بین تحصیلات و حقوق کارمندان ایفا می‌کند و موجب می‌شود شدت رابطه بین تحصیلات و حقوق با توجه به نوع سازمان در حال خدمت، تغییر کند و تعدیل شود.



شکل ۱-۴- نقش متغیر تعدیل‌گر در پژوهش

۵) متغیر کنترل:

برخی متغیرها وجود دارند که لازم است پژوهش‌گر آن‌ها را بشناسد و مدنظر قرار دهد و درباره آن‌ها تدابیری اتخاذ کند، اما در بیان فرضیه نامی از آن‌ها برده نمی‌شود، بلکه پژوهش‌گر جدا از عبارت یا جمله فرضیه به ذکر آن‌ها می‌پردازد، این متغیرها را کنترل می‌نامند. منظور از متغیرهای کنترل شده متغیرهایی هستند که در جریان پژوهش مورد کنترل قرار می‌گیرند و به تفسیر دیگر، تأثیرات آن‌ها کنترل یا خنثی می‌شود. از این متغیرها عمدتاً در پژوهش‌هایی که دو گروه آزمایش و شاهد (کنترل) داریم، استفاده می‌کنیم.

مثال: می‌خواهیم تأثیر استفاده از وسایل کمک آموزشی چند رسانه‌ای (مانند رایانه و ویدئو پرژکتور) را بر یادگیری دانش‌آموزان بسنجیم. در این مثال، وسایل کمک آموزشی متغیر مستقل و یادگیری متغیر وابسته است. بدین منظور دو گروه از دانش‌آموزان را انتخاب می‌کنیم که از نظر ویژگی‌هایی مثل سن، جنس و بهره‌هوشی با یکدیگر برابر یا شبیه باشند. همچنین معلم هر دو کلاس را یک نفر انتخاب می‌کنیم. برای آموزش درس شیمی برای یک ترم در یک کلاس از وسایل کمک آموزشی چند رسانه‌ای استفاده می‌کنیم و در کلاس دیگر از وسایل کمک آموزشی استفاده نمی‌کنیم. سپس از دانش

آموزان هر دو کلاس امتحان می‌گیریم که میانگین امتحان درس شیمی در کلاسی که در آن از رسانه‌های تصویری استفاده شده است ۱۹ و در کلاسی که در آن از وسایل کمک آموزشی چند رسانه‌ای استفاده نشده بود برابر با ۱۵ به دست آمده است.

نتیجه‌ای که می‌توانیم بگیریم آن است که به احتمال زیاد تفاوت نمره موجود بین دو گروه به استفاده از وسایل کمک آموزشی چند رسانه‌ای برمی‌گردد. وقتی می‌توانیم این جمله را بگوییم که صرفاً متغیر مستقل در یک جا وارد شده باشد و در جای دیگر وارد نشده باشد و به تعبیر دیگر، همه متغیرها غیر از متغیر مستقل کنترل شده باشند. در این‌جا، سن، جنس، بهره هوشی و فرد آموزش‌دهنده (معلم) متغیرهای کنترل می‌باشند.

۱. مستقل (X)	— نقش علت را دارد = روی متغیرهای دیگر اثر می‌گذارد
۲. وابسته (Y)	— نقش معلول را دارد = از متغیرهای دیگر اثر می‌پذیرد (پیش‌بینی شونده)
۳. میانجی (M)	— بین دو متغیر قرار می‌گیرد و هم نقش متغیر مستقل و هم وابسته را دارد اثر متغیر مستقل بر وابسته از طریق متغیر میانجی منتقل می‌شود
۴. تعدیل‌گر (Z)	— متغیر مستقلی است که نقش ثانویه دارد جهت و شدت رابطه متغیر مستقل و وابسته به میزان متغیر تعدیل‌گر بستگی دارد
۵- کنترل	— اثراتش در تحلیل حذف (کنترل) می‌شود عمدتاً در تحقیقات آزمایشی با گروه کنترل و آزمایش

آزمون‌های پارامتری و ناپارامتری

آزمون‌های آماری را بر اساس ویژگی‌های داده‌ها یا متغیرها به دو دسته تقسیم می‌کنند: آزمون‌های پارامتری و آزمون‌های ناپارامتری.

در جریان رشد روش‌های آماری مدرن، اولین تکنیک‌های استنباطی که پدید آمدند، تکنیک‌هایی بودند که برای جمعیتی که نمرات از آن برگرفته می‌شود مفروضات فراوانی قائل بودند. از آن جا که مقادیر جمعیت «پارامتر» اند، این تکنیک‌های آماری را پارامتری می‌خوانند. این دسته از تکنیک‌ها به نتیجه‌گیری‌هایی می‌انجامد که با قید و شرط همراه‌اند؛ برای مثال، «اگر مفروضات شکل جمعیت (جمعیت‌ها) معتبر باشد، آنگاه می‌توانیم نتیجه‌گیری کنیم که...»

در سال‌های اخیر تعداد زیادی تکنیک استنباطی ابداع شده است که مفروضات سختی را درباره پارامترها پیش نمی‌کشند. این تکنیک‌های جدیدتر، ناپارامتری یا توزیع آزاد هستند به طوری که نتیجه‌گیری‌ها در آن مستلزم هیچ قید و شرطی نیست.

آزمون‌های پارامتری

زمانی که داده‌های ما ویژگی‌های زیر را دارا باشند می‌توانیم از آزمون‌های پارامتری استفاده کنیم:

- ۱- مشاهده‌های ما باید از جمعیت‌های آماری که دارای توزیع نرمال هستند به دست آمده باشند.
- ۲- این جمعیت‌های آماری باید دارای واریانس برابر باشند (فرض همگن بودن واریانس‌ها).
- ۳- متغیرها باید حداقل با مقیاس فاصله‌ای اندازه‌گیری شده باشند.
- ۴- مشاهده‌های ما باید مستقل از یکدیگر باشند. یعنی انتخاب هر موردی از میان جمعیت، برای قراردادن آن مورد در نمونه، نباید مشکل و خللی در شانس انتخاب شدن هریک از موارد دیگر در نمونه ایجاد کند (سیگل، ۱۳۷۲: ۲۴-۲۵).

تمام شرایط بالا عناصری از مدل آمار پارامتری هستند. به استثنای فرض همگن بودن واریانس‌ها، این شرایط را معمولاً در جریان یک تحلیل آماری آزمون نمی‌کنند. بلکه این

شرایط «پیش‌فرض‌هایی» هستند که پذیرفته می‌شوند، و درستی یا نادرستی آن‌ها، معنی‌دار بودن عبارات نتیجه‌گیری ما را در کاربرد آزمون‌های پارامتری معین می‌کنند.

آزمون‌های پارامتریک چند مزیت دارند. در صورت مساوی بودن سایر شرایط، این آزمون‌ها نسبت به آزمون‌های ناپارامتریک توان و انعطاف بیشتری دارند. آزمون‌های پارامتریک نه تنها به پژوهش‌گر اجازه می‌دهند تا تاثیر متغیرهای مستقل متعددی را بر روی متغیر وابسته مطالعه کند، بلکه امکان مطالعه تاثیر متقابل آن‌ها را هم فراهم می‌کنند. نمونه‌های کوچک و انحراف در داده‌ها موجب می‌شود که از آزمون‌های ناپارامتریک استفاده کنیم (مونرو، ۱۳۸۹: ۱۳۴).

وقتی دلایلی موجود باشد که نشان بدهد اطلاعات جمع‌آوری شده ما واجد شرایط ذکر شده هستند در آن صورت یقیناً باید از آزمون‌های پارامتری از قبیل آزمون t یا آزمون F برای تجزیه و تحلیل داده‌ها استفاده شود. چنین انتخابی از آن جهت مطلوب است که آزمون‌های پارامتری در صورتی که مفروضات آن‌ها رعایت شود، توانمندترین آزمون‌ها برای رد کردن فرض صفر وقتی که باید آن را رد کنیم می‌باشند (سیگل، ۱۳۷۲: ۲۵).

آزمون‌های ناپارامتری

اما در صورتی که اطلاعات جمع‌آوری شده شرایط و پیش‌فرض‌های ذکر شده را نداشته باشند، مثلاً جامعه آماری ما توزیع نرمال نداشته باشد یا توزیع آن مشخص نباشد و یا زمانی که مقیاس متغیرها فاصله‌ای نباشد و ترتیبی و اسمی باشد چه باید کرد؟ در صورتی که داده‌ها یکی از مفروضات ذکر شده را نداشته باشند باید از آزمون‌های ناپارامتری استفاده کرد. گرچه شواهد تجربی نشان می‌دهد که انحراف جزئی از مفروضات آزمون‌های پارامتری تأثیرات چندانی شدید و عمیقی بر روی نتایج به دست آمده نمی‌گذارد. با این حال هنوز روی این مسأله که "انحراف جزئی" چقدر است اتفاق نظر عمومی و کلی به دست نیامده است (سیگل، ۱۳۷۲). در جدول ۱-۴ برخی از مهم‌ترین آزمون‌های پارامتری و ناپارامتری را می‌توان مشاهده کرد.

جدول ۱-۴- مقایسه انواع آزمون‌های پارامتریک و ناپارامتریک

نوع آزمون	سنجش رابطه	مقایسه میانگین یا میانه
پارامتری	همبستگی پیرسون - رگرسیون خطی	t تک نمونه‌ای - t گروه‌های مستقل t جفتی یا همبسته - تحلیل واریانس - تحلیل اندازه‌های مکرر
ناپارامتری	همبستگی کندال و اسپیرمن ضرایب فی و وی کرامر - مجذور کای استقلال	مجذور کای نیکویی برازش - مان ویتنی - ویلکاکسون کروسکال والیس - فریدمن

داده

داده یا مشاهده آماری، به اندازه متغیری خاص در موردی خاص گفته می‌شود. برای مثال، وقتی که در پیمایشی سن فردی (موردی خاص) را می‌پرسیم و او پاسخ می‌دهد ۳۰ سال دارد، «۳۰ ساله بودن سن آن پاسخگو» مشاهده آماری یا داده است. یا وقتی که میزان رشد جمعیت در افغانستان طی یک دهه گذشته را حساب می‌کنیم و معلوم می‌شود که میزان رشد سالانه جمعیتش ۳ درصد است، این «۳ درصد بودن رشد سالانه جمعیت افغانستان» یک داده یا یک مشاهده آماری است.

واژه نامه فصل اول

Alternative hypothesis	فرضیه خلاف	Parametric tests	آزمون های پارامتری
Null hypothesis	فرضیه صفر	Non-Parametric tests	آزمون های ناپارامتری
Variable	متغیر	Statistics	آمار
Nominal Variable	متغیر اسمی	Inferential statistics	آمار استنباطی
Continuous Variable	متغیر پیوسته	Descriptive Statistics	آمار توصیفی
Moderator Variable	متغیر تعدیل گر	Statistic	آماره
Interval Variable	متغیر فاصله ای	Measure	اندازه یا مقدار
Quantitative Variable	متغیر کمی	Parameter	پارامتر
Control Variable	متغیر کنترل	Ordinal Variable	متغیر ترتیبی (رتبه ای)
Qualitative Variable	متغیر کیفی	Statistical population	جمعیت آماری
Discrete Variable	متغیر گسسته	Standard Error	خطای استاندارد
Independent Variable	متغیر مستقل	Standard Error of Mean	خطای استاندارد از میانگین
Mediator Variable	متغیر میانجی	Data	داده
Ratio Variable	متغیر نسبی	Construct	سازه
Dependent Variable	متغیر وابسته	Level of measurement	سطح سنجش
Concept	مفهوم	Significance level	سطح معنی داری
Sample	نمونه	Confidence interval	فاصله اطمینان
Sampling	نمونه گیری	Hypothesis	فرضیه

ورود داده‌ها و آماده‌سازی برای تحلیل

محتوای فصل

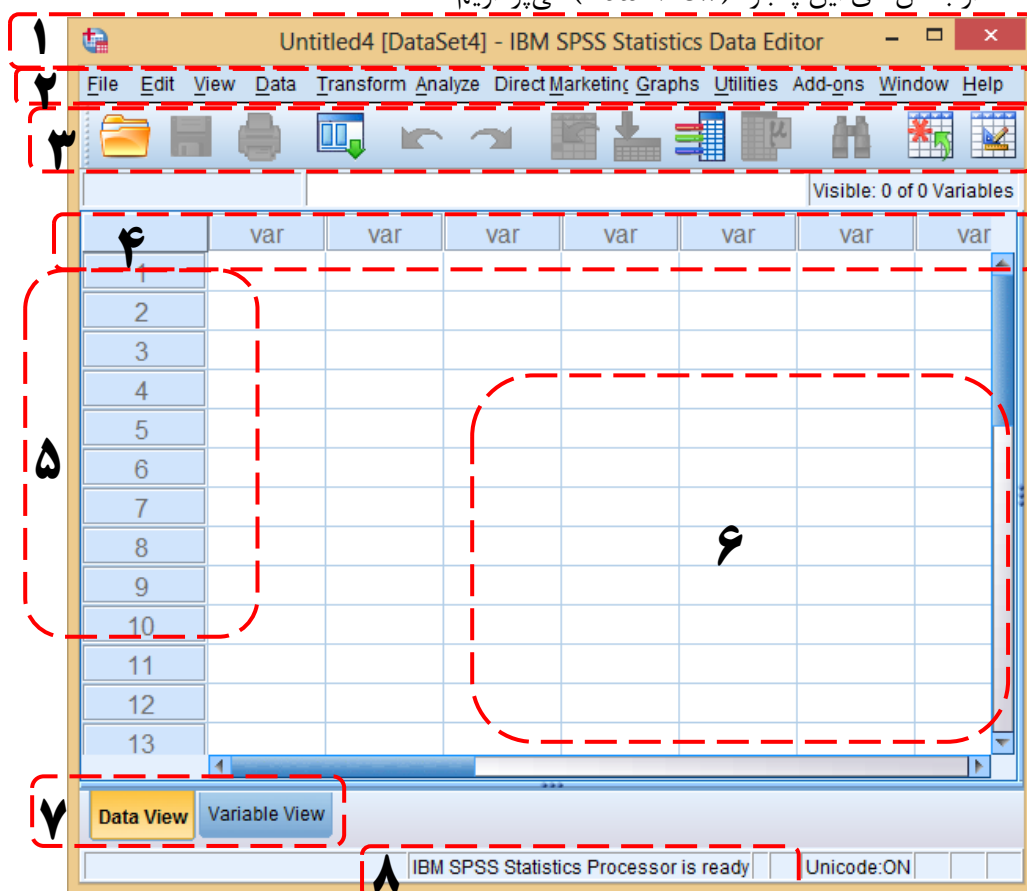
« بخش اول: آشنایی با برنامه و ورود داده: آشنایی با پنجره‌ها و منوهای برنامه ، مراحل انجام تحلیل آماری، تعریف متغیرها و وارد کردن داده‌ها.
بخش دوم: آماده سازی داده‌ها برای تحلیل: داده‌های پرت، داده های ناقص ، کدگذاری مجدد متغیرها، مقیاس سازی»

بخش اول:

آشنایی با برنامه و ورود داده ها

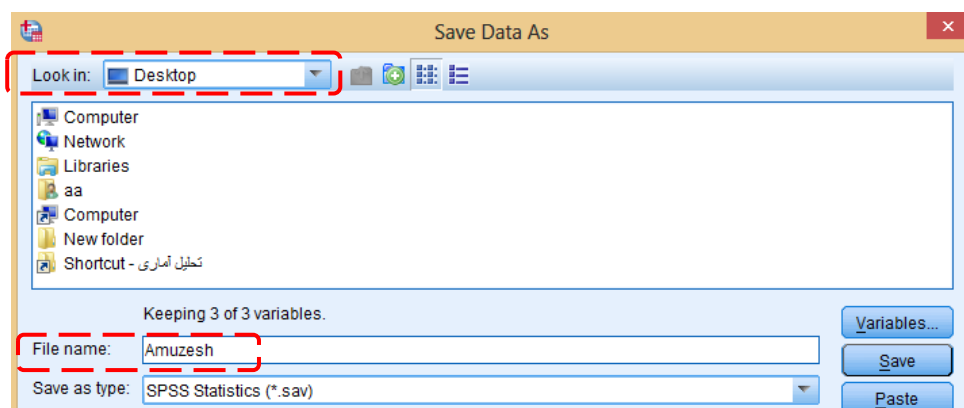
● آشنایی با پنجره ها و منوهای برنامه SPSS

در این بخش با اجزاء و دستورات برنامه SPSS آشنا می شویم. شکل زیر پنجره اصلی برنامه را تشکیل می دهد. با باز کردن برنامه، نخست با این پنجره مواجه می شویم. این پنجره دارای اجزاء و دستورات متنوعی است. در ادامه به طور مختصر به توضیح هر کدام از بخش های این پنجره (Data View) می پردازیم.



۱. **سربرگ یا عنوان فایل:** مانند تمامی برنامه‌های مشابه باید برای هر صفحه از SPSS که باز می‌کنیم و قصد ورود و تحلیل داده‌ها را در آن داریم، نامی انتخاب کنیم. برای انتخاب نام و ذخیره اطلاعات ابتدا باید اطلاعاتی را وارد برنامه کنیم. بعد از تعریف داده‌ها در پنجره متغیرها (Variable View)، می‌توانیم جهت ذخیره فایل داده‌ها از منوی دستورات به شیوه زیر عمل کنیم:

مسیر File > Save <--- را دنبال می‌کنیم:



و در کادر Look in محلی را که قصد ذخیره اطلاعات در آن را داریم مشخص می‌کنیم و در قسمت File name نامی برای فایل خود انتخاب می‌کنیم (نام انتخابی را نیز ترجیحاً انگلیسی می‌گذاریم) و سپس گزینه Save را می‌زنیم. بعد از این کار مشاهده می‌شود که به جای نام Untitled در سربرگ، نامی را که ما انتخاب کرده‌ایم جایگزین می‌شود. در مثال کتاب، نام Amuzesh در سربرگ مشاهده می‌گردد.

۳. **نوار دستورات:** برنامه SPSS جهت اجرای دستورات مختلف، منو یا نواری به نام نوار دستورات دارد. در این منو هر دستوری را که می‌خواهیم به برنامه بدهیم در دسترس است. در جدول ۱-۲ دستورات اصلی به همراه کاربردهای عمده آنان توضیح داده شده است.

جدول ۲-۱- مروری بر دستورات اصلی برنامه SPSS

نام دستور	کاربردهای اصلی
File	بازکردن فایل های قدیمی - ایجاد فایل های جدید - ذخیره داده ها - پرینت
Edit	کپی و انتقال داده ها - افزودن ستون یا ردیف جدید در فایل داده ها - یافتن اطلاعات - تغییر تنظیمات اولیه (Options) و...
View	نمایش یا پنهان سازی نوار ابزار یا خطوط زمینه در صفحه - تغییر نوع و اندازه فونت نوشتاری...
Data	مرتب سازی متغیرها (Sort) - تقسیم بندی داده ها - انتخاب پاسخگویان - وزن بخشی ...
Transform	ترکیب متغیرها و محاسبه مقادیر جدید - کدگذاری مجدد داده ها - جایگذاری مقادیر ناقص
Analyze	تحلیل داده ها یا انجام تمامی آمارهای توصیفی و استنباطی
Graphs	ترسیم انواع نمودارها

۳. **نوار ابزار:** شامل برخی ابزارهای پرکاربرد است که پژوهشگر می تواند جهت افزایش سرعت در کار خود از آن ها استفاده کند. لازم به ذکر است که تمامی ابزارهایی که در نوار ابزار وجود دارد در نوار دستورات نیز وجود دارد. ولی به جهت پرکاربرد بودن آن ها، در قالب یک نوار جداگانه در برنامه آورده شده است.

۴. **ستون نام متغیرها:** زمانی که در پنجره متغیرها (Variable View) برای هر کدام از متغیرها نامی انتخاب می کنیم، نام انتخاب شده در این ستون ها نمایش داده می شود. لازم به ذکر است که در پنجره داده ها (Data View) هر ستون به یک متغیر اختصاص دارد.

۵. **ردیف موارد یا پاسخگویان:** در این بخش اطلاعات مربوط به هر مورد یا پاسخگو یا هر پرسشنامه وارد می شود. هر ردیف اختصاص به یک مورد یا پاسخگو دارد.

۶. **خانه داده ها:** در این بخش اطلاعات و داده های مربوط به پاسخگویان وارد می شود. هر خانه یا سلول اختصاص به یک سوال از یک مورد یا پاسخگو دارد. یعنی در هر خانه تنها اطلاعات مربوط به یک سوال از پرسشنامه یک پاسخگو وارد می شود (یک داده).

۷. پنجره ویرایش‌گر داده‌ها: برنامه SPSS از دو پنجره تشکیل شده است که شامل پنجره ویرایش‌گر داده (Data Editor) و پنجره خروجی یا برون داد (Output یا SPSS Viewer) است.

پنجره ویرایش‌گر داده‌ها، اصلی‌ترین پنجره SPSS است. خود این پنجره از دو بخش تشکیل شده است که هر کدام محتوا و کارکرد خاصی دارند. این دو پنجره شامل پنجره داده‌ها (Data View) و پنجره متغیرها (Variable View) می‌شود.

- ۱) پنجره متغیرها (Variable View): جهت نام‌گذاری هر کدام از سوالات و تعریف خصوصیات هر کدام از متغیرها به کار می‌رود. در این پنجره ما باید ویژگی‌های هر متغیر را برحسب ۱۱ معیار و در ۱۱ ستون مشخص کنیم.
- ۲) پنجره داده‌ها (Data View): مهم‌ترین پنجره بوده و در آن اطلاعات و داده‌ها وارد می‌شوند. زمانی که پژوهش‌گر نام و ویژگی‌های متغیرها را در پنجره متغیرها تعیین کرد، بعد از آن تمام دستورات و تحلیل‌ها با پنجره داده‌هاست.

۸. پیغام نصب برنامه: در این کادر دو عبارت ممکن است نوشته شده باشد:

SPSS Processor is ready

و یا این عبارت

SPSS Processor is Unavailable

زمانی که پیغام اول نوشته شده باشد به این معناست که برنامه به طور کامل و صحیح نصب شده است و تمامی تحلیل‌های آماری را می‌توان انجام داد. زمانی که پیغام دوم نوشته شده باشد بدین معناست که نصب برنامه به طور ناقص و اشتباه صورت گرفته در نتیجه برنامه قادر به اجرای برخی و یا هیچ کدام از دستورات نیست. اگر چنین پیغامی در برنامه شما وجود دارد باید برنامه را حذف کرده و مجدداً و به شیوه صحیح اقدام به نصب برنامه کنید. اگر با نصب مجدد برنامه، باز هم با همین پیغام مواجه شدید می‌توانید از نسخه‌های بالاتر و یا پایین‌تر استفاده کنید و یا همان نسخه را از شرکت دیگری خریداری کنید.

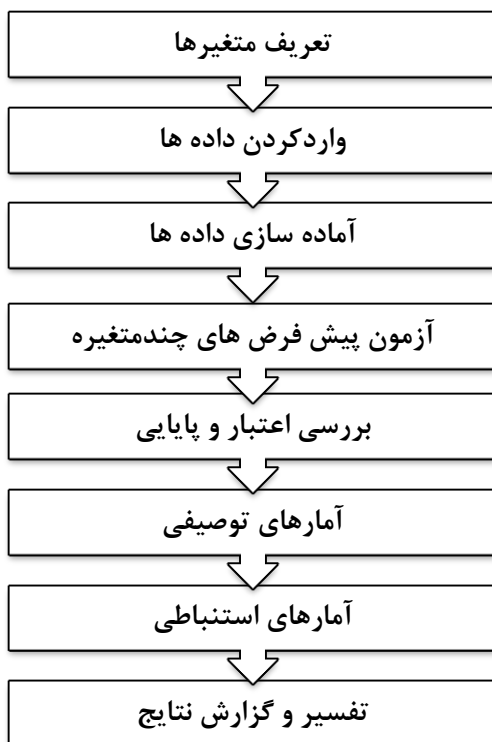
منوی ابزار شامل ابزارهایی است که معادل دستورات هستند. ابزارها میان بر دستورات هستند و امکان دسترسی به رایج ترین و پرکاربردترین دستورات را فراهم می کنند. در جدول ۲-۲ نام و عملکرد برخی از مهم ترین ابزارها ارائه شده است.

جدول ۲-۲- مروری بر ابزارهای اصلی برنامه SPSS

ابزار میان بر	نام دستور	کاربرد
	Open	برای پیدا کردن و باز کردن فایل های داده یا خروجی به کار می رود.
	Save	برای ذخیره فایل داده یا خروجی به کار می رود.
	Recall recently used dialogs	برای دسترسی سریع به آخرین دستورات اجرا شده به کار می رود.
	Undo a user action	آخرین دستور اجرا شده را لغو می کند و داده و خروجی را به حالت ماقبل دستور آخر برمی گرداند.
	Variables	می توانیم خصوصیات هر کدام از متغیرها را با اجرای این دستور مشاهده کنیم.
	Find	برای یافتن عدد و داده ای در داخل فایل داده ها به کار می رود.
	Insert Cases	برای افزودن سطری جدید در پنجره داده ها به کار می رود
	Insert Variable	برای افزودن ستونی جدید برای تعریف متغیری جدید به کار می رود. هم در پنجره متغیرها و هم در پنجره داده ها کاربرد دارد.
	Split File	برای تقسیم داده ها به کار می رود. می توانیم داده ها را بر حسب متغیری خاص (مثلا زن یا مرد بودن) تفکیک کرده گروه ها را به طور جداگانه بررسی کنیم
	Select Cases	برای انتخاب موردهای (یا پرسشنامه های) خاص صورت می گیرد. با این دستور می توانیم تحلیل ها را فقط بر روی بخشی از موردهای گزینش شده اجرا کنیم.
	Value Labels	برای تبدیل اعداد به برجسب (نام) آن ها به کاری می رود. با انتخاب این گزینه، چنانچه در پنجره متغیرها، متغیری را کدگذاری کرده باشیم (ستون Value) به جای کد داده شده، نام هر طبقه ظاهر می شود.

● مراحل انجام یک تحلیل آماری

دیاگرام زیر مراحل انجام یک تحلیل آماری را نشان می‌دهد. بعد از این که داده‌ها از طریق پرسشنامه یا به طرق دیگر جمع‌آوری شد و پرسشنامه‌های ناقص حذف گردید، جهت ورود داده‌ها (اطلاعات) و انجام آزمون‌های مورد نیاز بر روی آن‌ها مراحل زیر باید طی گردد. لازم به ذکر است که برخی مراحل ذکر شده جنبه تکمیلی داشته و ممکن است برخی از مراحل (مانند بررسی اعتبار و پایایی) در همه پژوهش‌ها نیاز نباشد اما سایر مراحل تقریباً در همه تحلیل‌های آماری مشترک است.



شکل ۱-۲- مراحل انجام تحلیل آماری

تعریف متغیرها و وارد کردن داده‌ها

برای انجام محاسبات کامپیوتری با برنامه SPSS، ابتدا باید داده‌های گردآوری شده را وارد برنامه کرده، ویژگی‌های آن‌ها را برای برنامه تعریف کنیم و در مرحله بعد به پردازش و اصلاحاتی مانند برطرف کردن خطاهایی که طی گردآوری داده‌ها و انتقال آن‌ها به برنامه پیش می‌آید پردازیم تا برای محاسبات آماری آماده شود. در این بخش نحوه تعریف متغیرها و سپس نحوه ورود داده‌ها در برنامه SPSS آموزش داده می‌شود.

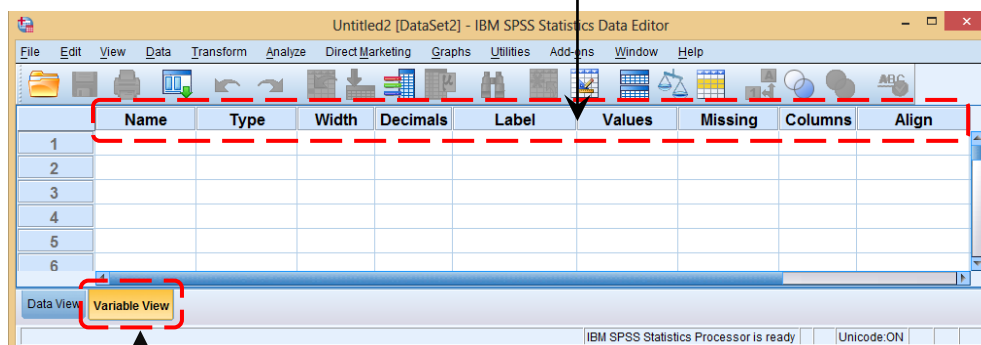
۱) تعریف متغیرها

به طور معمول (در پژوهش‌های مبتنی بر پرسشنامه) منظور از متغیر، سوالات پرسشنامه است. سوالاتی مثل جنس، سن، تحصیلات، معدل و غیره همگی یک متغیر محسوب می‌شوند و باید هر کدام از آن‌ها را به عنوان یک متغیر در برنامه تعریف کرد. دقت شود که هر سوال پرسشنامه، یک متغیر محسوب می‌شود.

در پرسشنامه این کتاب، ما با ۳۱ سوال (۱۳ سوال مربوط به خصوصیات استاد، ۵ سوال مربوط به تمایل به آموختن برنامه SPSS و ...) مواجهیم و باید ۳۱ متغیر در برنامه تعریف کنیم. در ادامه به نحوه تعریف متغیرها می‌پردازیم و متغیرهای جنس افراد (زن یا مرد بودن) و سن پاسخگویان را به عنوان تمرین و مثال در برنامه تعریف می‌کنیم.

جهت تعریف متغیرها باید وارد پنجره (Variable View) شویم. در این پنجره همان طور که مشخص است ۱۱ ستون (در نسخه‌های پایین‌تر ۱۰ ستون) وجود دارد که هر ستون مربوط به تعریف و تعیین یک خصوصیت متغیر است.

ستون‌های مشخصات متغیر: در هر کدام از این ستون‌ها باید یک خصوصیت هر متغیر را برای برنامه تعریف کنیم.



پنجره متغیرها: زمانی که این پنجره فعال باشد رنگ آن تغییر می‌کند و نارنجی می‌شود.

☑ نکته: گاهی تنها تعریف نام متغیر ضرورت دارد و ضرورتی به تعریف سایر خصوصیات متغیر نیست. اما چنانچه متغیر ما در سطح اسمی یا ترتیبی باشد لازم است مقادیر کد یا ارزش‌های متغیر هم تعریف گردد.

☑ نکته: همیشه قبل از تعریف کردن متغیرها و ورود داده‌ها، پرسشنامه‌های حاوی اطلاعات را شماره‌گذاری کرده و در برنامه اولین تغییری که تعریف می‌کنیم متغیر شماره پاسخگو یا شماره پرسشنامه باشد.

در جدول ۲-۳ به توضیح کاربرد هر کدام از اجزای پنجره تعریف متغیرها پرداخته ایم.

جدول ۲-۳- کاربرد اجزاء مختلف پنجره تعریف متغیر

نام ستون	نام خصوصیت	کاربرد و نحوه تعریف
Name	نام	جهت انتخاب نام هر متغیر استفاده می گردد. نام متغیر را تا حد امکان کوتاه انتخاب می کنیم.
Type	نوع متغیر (داده)	جهت تعیین ماهیت داده استفاده می شود مانند عددی یا حرفی بودن داده. پیش فرض برنامه Numeric (عددی) بوده و چون داده ها باید در برنامه به صورت عددی وارد شوند نیازی به تغییر نوع متغیر نیست.
Width	عرض (تعداد کاراکترها)	جهت تعیین تعداد ارقام متغیر استفاده می شود. پیش فرض برنامه برای داده های عددی ۸ کاراکتر است که دو رقم آن برای اعشار در نظر گرفته می شود. اگر داده هایی که وارد می کنیم در این عرض نگنجد، برنامه عدد را گرد می کند.
Decimals	تعداد اعشار	جهت تعیین تعداد ارقام اعشاری به کار می رود و در متغیرهای کمی کاربرد دارد. چنانچه متغیر ما فاقد اعشار باشد می توانیم تعداد اعشار را برابر صفر قرار دهیم. پیش فرض برنامه قرار دادن ۲ رقم اعشار برای هر متغیر است.
Label	برچسب متغیر	ارائه عبارت یا توضیحاتی در مورد سوال است که به ما در یادآوری نوع سوال کمک می کند. اگر سوال طولانی باشد می توان در ستون نام، نامی کوتاه برای آن انتخاب کرد و در بخش برچسب، کل عبارت را نوشت.
Values	کد یا ارزش	ارزش ها یا کدهای هر متغیر در این قسمت وارد می شود. معمولاً فقط برای متغیرهای اسمی و ترتیبی استفاده می شود. در واقع به هر کدام از طبقات متغیرها یک کد (عدد) اختصاص می دهیم.
Missing	مقادیر نامشخص	ارزش یا عددی را که برای مقادیر نامشخص (ناقص) در نظر گرفته ایم، در این بخش وارد می کنیم.
Columns	طول ستون	می توانیم عرض ستون هر متغیر را کوچکتر یا بزرگتر کنیم.
Align	چیدمان	به نحوه چینش مقادیر اختصاص دارد. می توانیم داده ها را در خانه ها به صورت راست چین، وسط چین و چپ چین انتخاب کنیم.
Measure	سطح سنجش	مقیاس و سطح سنجش متغیرها را تعیین می کنیم. انتخاب سطح سنجش هیچ گونه تاثیری در نتایج ندارد.

مثال

تعریف متغیر جنس (زن و مرد)

اجرا:

جهت فعال کردن هر خانه مربوط به تعریف متغیرها کافی است در داخل خانه‌ها کلیک کرده و روی کادر یا علامت ظاهر شده مجدداً کلیک کنیم.

۱) نام گذاری:

در کادر NAME نامی برای متغیر خود انتخاب می‌کنیم. در نسخه‌های جدید برنامه می‌توانیم اسامی را فارسی بگذاریم ولی در نسخه‌های قدیمی‌تر ممکن است با مشکلاتی در فارسی نوشتن نام‌ها مواجه شویم. همچنین بدین دلیل که قراردادن فاصله بین کلمات در کادر نام امکان‌پذیر نیست باید یا نام‌ها را تک کلمه‌ای انتخاب کرد و یا بین دو کلمه به جای قراردادن فاصله، یک علامت مانند نقطه قرار داد. موارد زیر را باید در نگارش نام متغیرها رعایت کرد:

نام متغیر نباید بیش از ۸ کاراکتر (نه حرف) باشد.

نام متغیر را نمی‌توان با یک عدد یا علامت (مانند !، ؟، * یا @) شروع کرد، اما می‌تواند شامل حروف، اعداد یا یکی از کاراکترهای @، _ یا نقطه باشد. اگر از کاراکترهای غیرمجاز استفاده کنیم، هنگام خارج شدن از ستون نام، برنامه پیام خطا خواهد داد و نام‌گذاری انجام نخواهد شد.

بین کاراکترهای نام متغیر، نباید فاصله وجود داشته باشد.

در نام‌گذاری، حروف کوچک و بزرگ با یکدیگر فرقی ندارند.

نام متغیر باید منحصر به فرد باشد (دو متغیر نمی‌توانند نام یکسان داشته باشند).

۲) تعیین اعشار:

در کادر Decimals می‌توانیم تعداد اعشار متغیرها را تغییر بدهیم. به طور پیش‌فرض تعداد ۲ اعشار برای متغیرها در نظر گرفته شده است. در مورد متغیر جنس چون تنها با دو عدد ۱ و ۲ مواجه هستیم، نیازی به اعشار نداریم و تعداد اعشار را صفر قرار می‌دهیم.

۳) کدگذاری:

در کادر Value کد یا عددی را که برای هر طبقه در نظر گرفته ایم می نویسیم. متغیر جنس از دو طبقه زن و مرد تشکیل شده است که ما به صورت قراردادی به زنان کد ۱ و به مردان کد ۲ را اختصاص می دهیم.

جنس افراد متغیری اسمی است و شامل دو طبقه مرد و زن می شود. در برنامه SPSS مقدار هر متغیر را با عدد نشان داده و تعریف می کنیم. در نتیجه ما باید به جای این که در برنامه، در ستون مربوط به جنس افراد بنویسیم زن یا مرد؛ زن یا مرد بودن افراد را باید با عدد مشخص کنیم (کدگذاری). رایج است که برای تعریف متغیرهای اسمی و ترتیبی از اعداد تک رقمی استفاده می کنند و به اولین طبقه کد ۱ ارزش را می دهند و به طبقه دیگر ۲ و به همین شکل الی آخر. در مثال مربوط به جنس افراد نیز چون ما با دو طبقه زن یا مرد روبرو هستیم در نتیجه باید به یک طبقه کد ۱ و به طبقه دیگر کد ۲ بدهیم. در متغیرهای اسمی (مثل قومیت، وضعیت تاهل و دین) تفاوتی وجود ندارد که به کدام متغیر چه ارزش یا کدی را اختصاص بدهیم. در نتیجه ما به طور اختیاری به کسی که زن باشد در برنامه کد ۱ می دهیم و به کسی که مرد باشد در برنامه کد ۲ می دهیم.

چنانچه متغیر ما ترتیبی (میزان تحصیلات یا میزان درآمد شامل سه طبقه کم درآمد، میان درآمد و پردرآمد) باشد باید ترتیب طبقات لحاظ گردد و کدگذاری متناسب با رتبه و ترتیب طبقات باشد. مثلاً اگر تحصیلات افراد را به صورت ترتیبی سنجیده و طبقات شامل سیکل، دیپلم، لیسانس و فوق لیسانس باشد، در این صورت باید ترتیب کدهای متناسب با میزان تحصیلات افراد باشد و به سیکل کد ۱، به دیپلم کد ۲، به لیسانس کد ۳ و به فوق لیسانس کد ۴ اختصاص بدهیم. در مورد متغیرهایی که در قالب طیف لیکرت سنجیده شده اند نیز این نکته باید رعایت شود. مثلاً در یک طیف سه گزینه ای؛ به مخالفم کد ۱، به بی نظر کد ۲ و به موافقم کد ۳ اختصاص می دهیم.

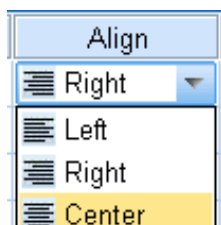
☑ نکته: کدگذاری طبقات متغیرهای اسمی تنها جهت سهولت ورود اطلاعات به کامپیوتر است و کدگذاری به هیچ وجه به معنای کمی شدن طبقات نیست و به این معنا نیست که زن معادل عدد ۱ است و مرد معادل عدد ۲ است. همچنین برای متغیرهای فاصله ای/نسبی (مانند سن، معدل یا بهره هوشی)، کدگذاری صورت نمی گیرد و نمره افراد در این متغیرها به طور مستقیم وارد برنامه می شود.

۴) تعریف برچسب:

در کادر Label برچسب هر متغیر را می‌نویسیم. متغیر جنس نیازی به توضیح بیشتر ندارد و برخورداری از نام کفایت می‌کند. قواعدی که در نام‌گذاری متغیرها صادق است، در مورد تعریف کردن برچسب متغیرها در ستون Label صدق نمی‌کند. برای متغیر جنس پاسخگویان ضرورتی به تعریف سایر مشخصات (تعیین نوع متغیر، عرض ستون، تعریف برچسب، تعداد اعداد و تعریف داده‌های ناقص) و تغییر پیش‌فرض‌های آنان نیست و تعریف این مشخصات اختیاری است.

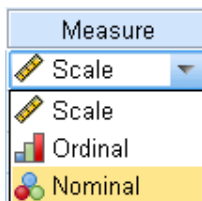
۵) چیدمان:

در کادر Align می‌توانیم طریقه چینش داده را در درون هر خانه تغییر بدهیم. برای نمونه می‌توانیم با کلیک در درون هر خانه، طریقه چینش متغیرها را وسط چین (Center) انتخاب کنیم.



۶) سطح سنجش

در کادر Measure می‌توانیم سطح سنجش متغیر را انتخاب کنیم جنس افراد متغیری اسمی است. در نتیجه در کادر مربوطه، گزینه Nominal (اسمی) را انتخاب می‌کنیم (Ordinal به معنای رتبه‌ای و Scale معادل فاصله‌ای/نسبی است).



شکل زیر نحوه تعریف متغیر جنس را نشان می‌دهد.

Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure
جنس	Numeric	8	0		{1, زن}...	None	8	Center	Nominal

۲) وارد کردن داده‌ها:

بعد از تعریف متغیرها در پنجره متغیرها (Variable View)، وارد پنجره داده‌ها (Data View) می‌شویم. در این پنجره باید اطلاعات مربوط به پاسخگویان یا اطلاعات پرسشنامه‌ها را وارد نماییم.

در هنگام وارد کردن داده‌ها در پنجره داده‌ها، به صورت ردیفی یا سطری پیش می‌رویم و در هر مرحله تمامی اطلاعات مربوط به یک پاسخگو را وارد برنامه می‌کنیم. به طور معمول، هر ردیف نشان دهنده یک پاسخگو و هر ستون نشان دهنده یک متغیر است.

اجرا:

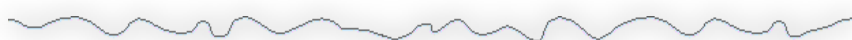
ابتدا تمامی پرسشنامه‌ها را از عدد ۱ شروع به شماره‌گذاری می‌کنیم. شماره‌گذاری پرسشنامه‌ها در مرحله بازبینی و اصلاح به ما کمک می‌کند تا اشتباهات را یافته و تصحیح نماییم. ضمناً از دوباره‌کاری‌های احتمالی جلوگیری کرده و به فرآیند وارد کردن اطلاعات نظم می‌بخشد.

☑ نکته: بعد از این‌که همه متغیرها را در پنجره متغیرها تعریف کردیم به پنجره داده‌ها می‌رویم. همچنین در هر مرحله اطلاعات یک پاسخگو یا پرسشنامه را وارد می‌نماییم و بعد از این‌که اطلاعات هر پاسخگو به‌طور کامل وارد شد، به سراغ پاسخگوی بعدی می‌رویم.

اطلاعات مربوط به یک پاسخگو		اطلاعات مربوط به متغیر جنس				
	شماره پاسخگو	جنس	تعداد اعضا	تحصیلات پدر	قومیت	سن
1	1	1	7	1	2	22
2	2	1	5	2	5	26
3	3	1	4	2	4	29

شکل بالا مربوط به پنجره داده‌ها بوده و در آن اطلاعات پاسخگویان وارد برنامه شده است. مطابق شکل، هر سطر نشان دهنده یک پاسخگو است و هر ستون نیز اختصاص به یک متغیر دارد. مثلاً در مورد سطر مشخص شده باید گفت که ستون اول (موردها) نشان‌دهنده شماره پاسخگو یا شماره پرسشنامه است و از ۱ شروع شده و به ترتیب بالا می‌رود (۲، ۳، ۴ و ...). ستون دوم (جنس) نشان‌دهنده متغیر جنس پاسخگو است که

برابر با ۱ است و نشان می‌دهد که پاسخگوی اول زن است. ستون سوم (تعداد اعضا) نشان دهنده تعداد اعضای خانواده است و تعداد اعضای خانواده فرد اول برابر با ۴ نفر است. ستون چهارم (تحصیلات) نشان‌دهنده میزان تحصیلات است و نشان می‌دهد که پاسخگوی اول میزان تحصیلاتش برابر با فوق‌دیپلم (کد ۳) است و سایر متغیرها نیز به همین ترتیب. ستون مشخص شده نیز جنس هر پاسخگو را نشان می‌دهد، که در کادر مشخص شده می‌توان گفت که سه نفر اول زن هستند (کد ۱ مربوط به جنس زن و کد ۲ مربوط به جنس مرد است).



تعریف

۱- سن ده دانشجو در ادامه آمده است. متغیری به نام سن در برنامه تعریف کرده و اطلاعات زیر را وارد کنید.

۲۳-۳۰-۲۶-۲۲-۲۹-۲۹-۲۵-۲۴-۲۵

۲- در پژوهشی ۱۵ آزمودنی شرکت داشته‌اند که اطلاعات مربوط به وزن آن‌ها آمده است. اطلاعات برحسب استان محل سکونت آمده است. دو متغیر جدید به نام وزن و استان تعریف کرده و اطلاعات زیر را وارد کنید. به استان کرمان کد ۱، شیراز کد ۲ و اصفهان کد ۳ بدهید.

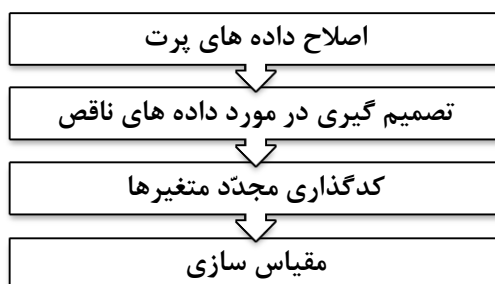
وزن	استان محل سکونت
۵۷-۶۲-۶۸-۶۲-۷۵	کرمان
۸۸-۷۹-۹۰-۷۷-۷۵	شیراز
۵۴-۸۰-۶۴-۷۴-۷۷	اصفهان

بخش دوم:

آماده سازی داده ها برای تحلیل

قبل از آن که بر روی داده ها تحلیل نهایی انجام شود، باید داده ها بررسی شده و اصلاحات و تغییراتی بر روی آن ها انجام شود تا برای تحلیل نهایی آماده شوند. برخی از دلایل آماده سازی داده ها عبارتند از:

- ۱- ممکن است در ورود داده ها اشتباهاتی صورت گرفته باشد. مثلاً ممکن است به جای عدد ۲ به اشتباه عدد ۲۲ وارد شده باشد و یا این که برخی داده ها سهواً وارد برنامه نشده باشند.
 - ۲- وجود داده های ناقص امری محتمل است و همواره باید در مورد داده های ناقص چاره ای اندیشید. چرا که داده های نامشخص بر نتایج تحلیل تأثیرگذار است.
 - ۳- قبل از انجام تحلیل باید بر روی داده ها کدگذاری مجدد انجام شود. ممکن است، نیاز داشته باشیم تغییری فاصله ای را به متغیری ترتیبی یا اسمی تبدیل کنیم (مثلاً افراد را بر حسب نمره شان در یک درس در سه گروه ضعیف، متوسط و قوی قرار دهیم). و یا این که جهت کدگذاری برخی متغیرها یا سوالات را معکوس کنیم تا همه متغیرهای یک مقیاس، هم جهت شوند.
 - ۴- در بسیاری از پژوهش های علوم انسانی و رفتاری ما با مقیاس سازی مواجهیم و باید با جمع کردن و ترکیب متغیرها، اقدام به ساختن متغیری جدید و یا مقیاسی برای یک مفهوم انتزاعی نماییم.
- با توجه به موارد بالا می توان مراحل آماده سازی داده ها را شامل چهار مرحله دانست. باید دانست که ممکن است در یک پژوهش به همه مراحل نیازی نیست.



شکل ۲-۲- مراحل آماده سازی داده ها

داده‌های پرت

همیشه باید داده‌هایی (اطلاعاتی) که وارد برنامه‌هایی مانند اکسل یا SPSS می‌کنیم را بررسی و بازبینی کنیم. همواره احتمال دارد که در داده‌ها با مقادیر غیرعادی مواجه شویم. موارد غیرعادی می‌تواند شامل مقادیر تعریف نشده و مقادیر پرت (دور افتاده) باشد. همواره قبل از انجام هرگونه تحلیل آماری بر روی داده‌ها، باید چاره‌ای در مورد مقادیر پرت بیندیشیم.

افرادی که اندازه‌های انتهایی یا غیرمعمول در یک متغیر واحد (تک متغیری) یا در ترکیبی از متغیرها (چندمتغیری) دارند، دور افتاده یا پرت نامیده می‌شوند. داده‌های پرت اغلب سه یا بیش از سه واحد انحراف معیار (± 3 SD) از میانگین مربوط به خودشان فاصله دارند که از مشکلات احتمالی در ابزار اندازه‌گیری، شیوه ثبت یا ضبط پاسخ‌ها یا عضویت شرکت‌کنندگان در جامعه‌ای که فرض می‌شود از آن نمونه‌گیری شده است، ناشی می‌شود. حضور داده‌های پرت می‌تواند نتایج تحلیل را به گونه‌ای نامطلوب تحت تأثیر قرار دهد (تحریف کند). به همین دلیل بیشتر متخصصان پیشنهاد می‌کنند که اندازه‌های پرت قبل از تحلیل داده‌ها باید حذف شوند (میزر و دیگران، ۱۳۹۱: ۲۳۶).

انواع داده‌های پرت

داده‌های پرت را می‌توان در دو دسته داده‌های پرت تک متغیری و داده‌های پرت چندمتغیری تقسیم کرد:

۱) داده‌های پرت تک متغیری

داده‌های پرت تک متغیری مربوط به یک متغیر می‌شوند. به عنوان مثال وقتی که در یک پژوهش دانشجویی در زمینه میزان رضایت مردم از عملکرد شهرداری تهران؛ ما در متغیر سن افراد با عدد ۱۵۰ روبرو می‌شویم؛ به احتمال زیاد با داده پرت مواجه شده ایم. چرا که می‌دانیم احتمال وجود فردی با چنین سن و سالی بسیار بعید است! و یا وقتی که در متغیر درآمد، شخصی درآمد ماهانه خود را از یک کار تمام وقت ۲۵ هزار تومان اعلام می‌کند و یا وقتی که در پاسخ سوالی که از فرد می‌پرسیم تا چه اندازه به آینده امیدوار است و او باید میزان رضایت خود را از عدد ۱ (به معنای خیلی کم) تا عدد

۵ (به معنی خیلی زیاد) اعلام کند، در فایل داده‌ها با عدد ۶ روبرو می‌شویم (به دلیل اشتباه در ورود داده)، همگی نشان از وجود داده‌های پرت تک متغیری دارد که نخست باید آن‌ها را شناسایی کرد و سپس در مورد آن‌ها چاره‌ای اندیشید.

البته زمانی که با متغیرهای کیفی (اسمی و ترتیبی) سروکار داریم گاهی با مقادیری در داده‌ها روبرو می‌شویم که داده پرت محسوب نمی‌شوند اما مقادیری هستند که به اشتباه وارد شده‌اند و باید حذف شوند. مثلاً در متغیر جنس، اگر ما زنان را با کد ۱ و مردان را با کد ۲ تعریف کرده باشیم و در این حال با عدد ۱.۵ در داده‌ها مواجه شویم؛ با داده پرت مواجه نیستیم اما با داده‌های اشتباه مواجه شده‌ایم (به دلیل اشتباه پاسخگو در پاسخ به سوال یا اشتباه در ورود داده) و باید آن‌ها را شناسایی کرده و حذف یا اصلاح نماییم.

شناسایی داده‌های پرت تک متغیری

برای شناسایی داده‌های پرت تک متغیری باید از جدول فراوانی و نمودار جعبه‌ای استفاده کرد. از جدول فراوانی برای شناسایی داده‌های پرت در متغیرهای اسمی و ترتیبی استفاده می‌کنیم و از نمودار جعبه‌ای برای شناسایی داده‌های پرت در متغیرهای فاصله‌ای/نسبی. البته از جدول فراوانی هم می‌توان برای شناسایی داده‌های پرت در متغیرهای فاصله‌ای/نسبی استفاده کرد ولی نمودار جعبه‌ای برتری دارد و آسان‌تر است.

الف) جداول فراوانی

از جدول فراوانی برای کشف مقادیر پرت تک متغیری در متغیرهای اسمی و ترتیبی استفاده می‌کنیم. متغیرهایی مثل جنس، وضعیت تاهل، قومیت، تحصیلات و درآمد (هر دو به صورت چندگزینه‌ای و ترتیبی سنجیده شده باشند، مثلاً تحصیلات در قالب سوالات دیپلم، فوق دیپلم، لیسانس و... سنجیده شده باشد) و یا تمام سوالاتی که در قالب طیف لیکرت سنجیده شده باشند. یعنی سوالاتی که پاسخ‌های آنان معمولاً ۳ تا ۷ گزینه دارد و پاسخ‌هایی مثل کاملاً موافقم تا کاملاً مخالفم، اصلاً تا همیشه و خیلی کم تا خیلی زیاد را در بر می‌گیرد. همچنین اگر متغیری فاصله‌ای/نسبی داشته باشیم که تعداد طبقات آن محدود (مثلاً حدود ۱۰ طبقه) باشد، می‌توانیم از جدول فراوانی استفاده کنیم.

مثال

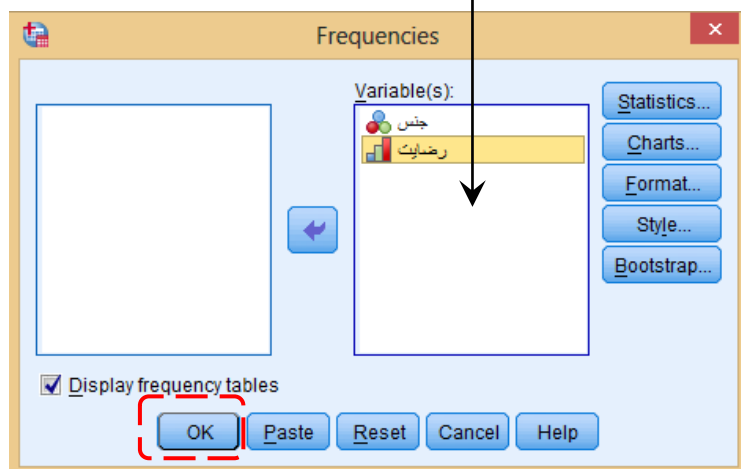
در یک پژوهش (فرضی) از دانشجویان دختر و پسر دانشگاه شهید بهشتی خواسته شد تا میزان رضایت خودشان از عملکرد ریاست دانشگاه را اعلام کنند. بر این اساس از دانشجویان تعدادی سوال پرسیده شد که دو سوال آن عبارت بود از جنس دانشجویان و میزان رضایتشان از عملکرد ریاست دانشگاه. جنس دانشجویان شامل دو جنس (دختر کد ۱، و پسر کد ۲) و میزان رضایت در طیف لیکرت ۵ گزینه‌ای (خیلی کم کد ۱، کم کد ۲، متوسط کد ۳، زیاد کد ۴ و خیلی زیاد کد ۵) سنجیده شد. همان‌طور که مشاهده می‌شود ما هنگام ورود اطلاعات مربوط به جنس افراد به دانشجویان دختر کد ۱ یا عدد ۱ و به دانشجویان پسر کد ۲ داده‌ایم و در فایل داده‌ها و خروجی (برون‌داد) مربوط به آن، تنها باید عدد ۱ و عدد ۲ مشاهده کنیم. در مورد متغیر میزان رضایت هم تنها باید اعداد ۱، ۲، ۳، ۴ و ۵ را مشاهده کنیم و نباید اعداد دیگری را (مثلاً ۶، ۱.۵، ۲۰) مشاهده کنیم.

اجرا:

دستور فراوانی را اجرا می‌کنیم:

Analyze ---> Descriptive Statistics ---> Frequencies

متغیرهای جنس و رضایت را وارد کادر Variables (متغیرها) می‌کنیم و گزینه OK را می‌زنیم



نتایج:

نتایج جدول فراوانی دو متغیر جنس و میزان رضایت در ادامه ارائه شده است. در جدول فراوانی جنس افراد مقادیر پرت مشاهده نمی شود، چرا که تنها دو کد یا طبقه ۱ و ۲ (دختر و پسر) وجود دارند. توجه شود که داده های گمشده (Missing) جزء داده های پرت به حساب نمی آیند. ما در فایل داده ها مقادیر گمشده را با عدد ۹ نشان داده ایم و در فایل خروجی اعداد گمشده با عدد ۹ ظاهر شده اند. به غیر از اعداد ۱ و ۲ و مقادیر گمشده، عدد دیگری در فایل خروجی جنس دانشجویان دیده نمی شود و بدین معناست که در متغیر جنس دانشجویان داده پرت وجود ندارد.

اما در متغیر میزان رضایت ما با اعدادی غیر از ۱، ۲، ۳، ۴ و ۵ مواجه ایم و این اعداد مقادیر گمشده هم نیستند و نشان می دهد که دو مقدار پرت در داده ها وجود دارد (۱.۳ و ۲۲) که باید در فایل داده ها شناسایی و حذف شود. چون پاسخگویان تنها می توانستند یکی از اعداد ۱، ۲، ۳، ۴ و ۵ را انتخاب کنند در نتیجه اعداد دیگری که وجود دارند (۱.۳ و ۲۲) مقادیر پرت حساب می شوند و باید از تحلیل حذف شوند.

لازم به ذکر است که عدد ۱.۳ داده پرت به حساب نمی آید و یک داده غیرعادی و تعریف نشده است. در این جا به جهت آسان تر شدن آموزش، داده های غیرعادی و تعریف نشده در ارتباط با متغیرهای اسمی و ترتیبی را داده پرت به حساب آورده ایم.

جنس

	Frequency فراوانی	Percent درصد فراوانی	Valid Percent درصد معتبر	Cumulative Percent درصد تجمعی
Valid مرد	133	66.5	66.8	66.8
زن	66	33.0	33.2	100.0
Total	199	99.5	100.0	
Missing 9	1	.5		
Total	200	100.0		

میزان رضایت

	Frequency	Percent	Valid Percent	Cumulative Percent
1	94	47.0	47.0	47.0
1.3	1	.5	.5	47.5
2	48	24.0	24.0	71.5
3	37	18.5	18.5	90.0
4	12	6.0	6.0	96.0
5	7	3.5	3.5	99.5
22	1	.5	.5	100.0
Total	200	100.0	100.0	

ب) نمودار جعبه‌ای

اگر متغیرهایی که سنجیدیم از نوع متغیرهای فاصله‌ای/نسبی باشند هم می‌توان از جداول فراوانی استفاده کرد و هم از نمودار جعبه‌ای. البته پیشنهاد می‌شود از نمودار جعبه‌ای استفاده شود زیرا در نمودار جعبه‌ای داده‌های پرت در داخل خود نمودار مشخص می‌شود و شماره موردی (پاسخگویی) که دارای داده پرت در فایل داده‌هاست در نمودار مشخص می‌شود. همچنین مبنای انتخاب داده پرت در نمودار جعبه‌ای، داشتن فاصله‌ای به اندازه حداقل ± 3 واحد انحراف استاندارد با میانگین است که به صورت خودکار توسط برنامه محاسبه می‌شود و در نمودار جعبه‌ای نشان داده می‌شود.

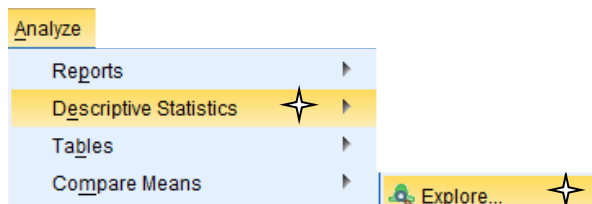
مثال

در پرسشنامه رضایت دانشجویان علاوه بر سوالات جنس و میزان رضایت، سن و معدل کارشناسی دانشجویان هم پرسیده شد. سن و معدل دانشجویان در سطح سنجش فاصله‌ای/نسبی سنجیده شد. دو نمودار جعبه‌ای سن و معدل دانشجویان در ادامه آورده شده است.

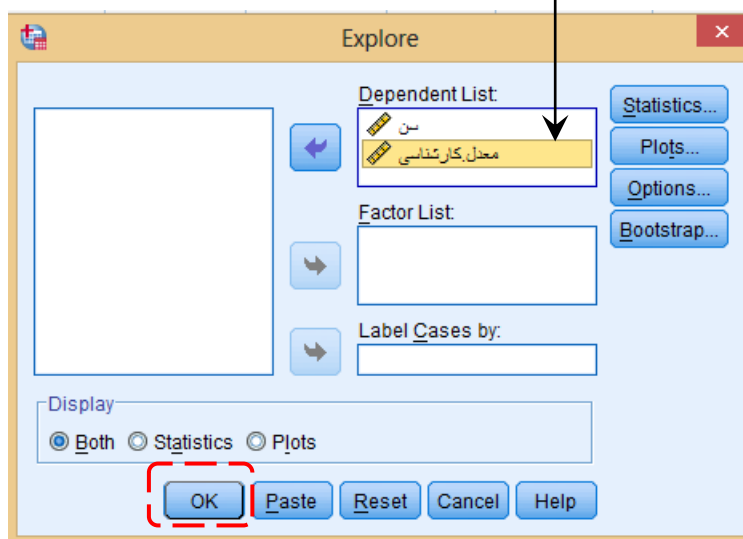
اجرا:

دستور اجرای نمودار جعبه‌ای در دستور Explore است. مراحل زیر را دنبال می‌کنیم:

Analyze ---> Descriptive Statistics ---> Explore



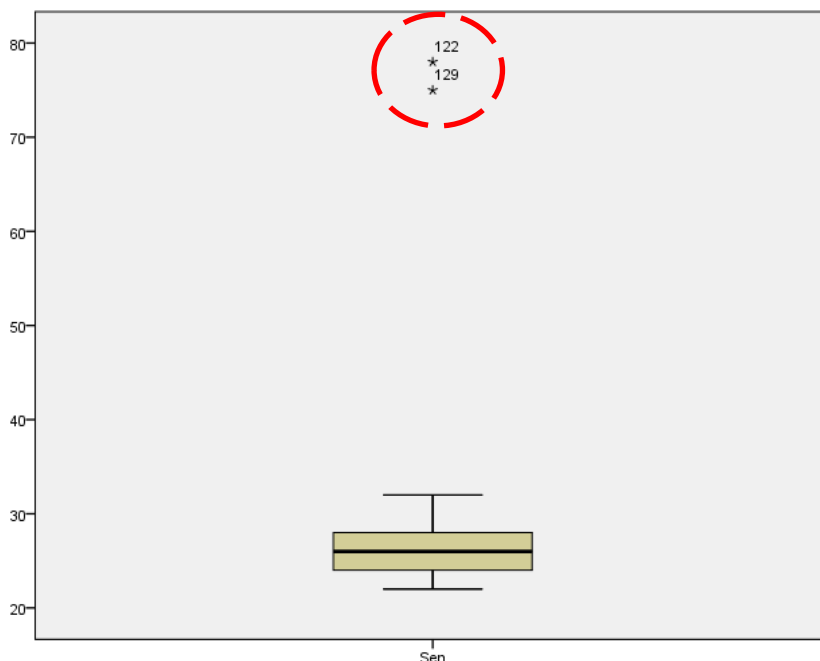
متغیرهای مدنظر را وارد کادر Dependent List می‌کنیم.
در انتها گزینه Ok را انتخاب می‌کنیم.



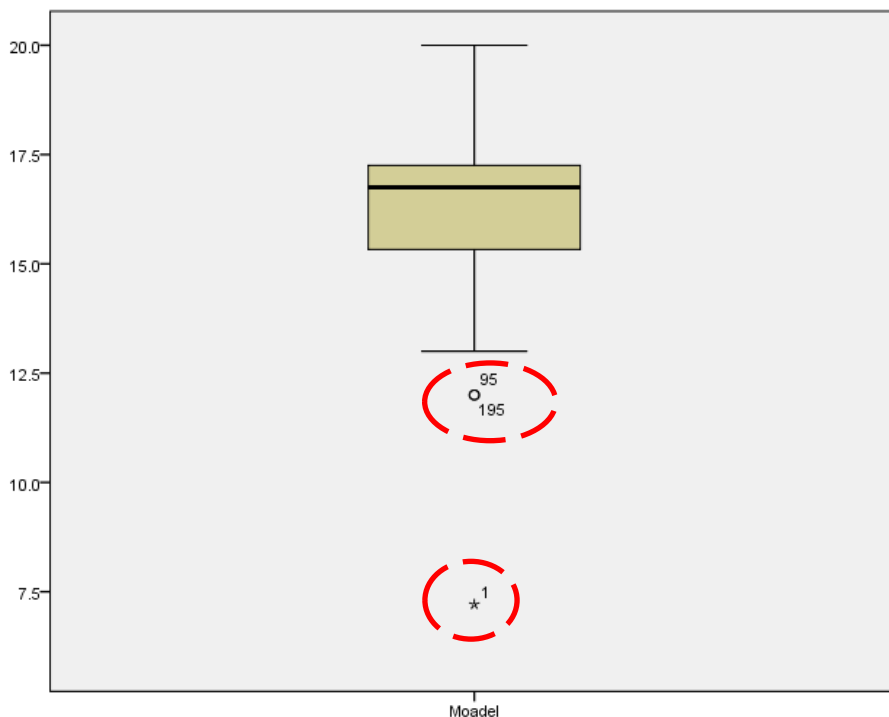
نتایج:

نمودار سن افراد نشان می‌دهد دو مورد دارای مقادیر پرت هستند. شماره این دو مورد در نمودار با علامت ستاره یا دایره کوچک مشخص شده است. افراد شماره ۱۲۲ و ۱۲۹ دارای مقادیر پرت هستند و به احتمال زیاد باید از داده‌ها حذف شوند. در فایل داده‌ها، افراد (پرسشنامه‌های) شماره ۱۲۲ و ۱۲۹ را پیدا کرده و سن آنان را نگاه می‌کنیم. سن این افراد به ترتیب ۷۸ و ۷۵ سال است. با توجه به این که احتمال وجود دانشجویانی با چنین سنی بسیار بعید است سن این افراد باید از تحلیل حذف شود.

در متغیر معدل سه داده پرت وجود دارد که شامل موردهای ۹۵، ۱۹۵ و ۱ می‌شود. معدل این افراد به ترتیب ۱۲، ۱۱.۲۵ و ۵.۲۱ است. تصمیم در مورد حذف این داده‌ها با پژوهش‌گر است. معدل این افراد حداقل ± 3 واحد انحراف استاندارد با معدل کل دانشجویان فاصله دارد. معدل کل دانشجویان برابر با ۱۶.۵۰ و انحراف استاندارد ۱.۷۵ است. به نظر می‌رسد با این که معدل‌های ۱۲ و ۱۱.۲۵، معدل پایینی است و احتمال این که دانشجویی چنین معدل‌هایی داشته باشد اندک است؛ اما همیشه چنین دانشجویانی وجود داشته‌اند و فرض وجود چنین معدل‌هایی محتمل است. بنابراین موردهای ۹۵ و ۱۹۵ باقی مانده و حذف نمی‌شوند. ولی مورد ۱ که دارای معدل ۵.۲۱ است حذف می‌شود، چون دارای معدل خیلی پایینی است. به احتمال زیاد این فرد معدل خود را به عمد اشتباه اعلام کرده است و یا این که در هنگام ورود داده‌ها از پرسشنامه به برنامه اشتباهی صورت گرفته است.



شکل ۲-۳- نمودار جعبه‌ای متغیر سن



شکل ۲-۴- نمودار جعبه‌ای معدل

۲) داده‌های پرت چند متغیری

بعد از بررسی داده‌ها برای شناسایی داده‌های پرت تک متغیری، سنجش داده‌های پرت چند متغیری انجام می‌شود. یک روش عینی برای ارزیابی وجود داده‌های پرت چندمتغیری محاسبه فاصله مهالانوبیس هر فرد است. آماره فاصله مهالانوبیس یعنی D^2 ، «فاصله» چندمتغیری بین هر فرد و میانگین چندمتغیری گروه را (که کانون نامیده می‌شود) اندازه‌گیری می‌کند.

هر فرد با استفاده از توزیع مجذورکای با سطح آلفای دقیق ۰.۰۱/ ارزیابی می‌شود. افرادی را که به این آستانه معنی‌داری می‌رسند می‌توان به عنوان موارد پرت چندمتغیری تلقی کرد و به احتمال باید از نمونه حذف شود (میزر و دیگران، ۱۳۹۱: ۱۰۶). چنانچه در سطح آلفای ۰.۰۱/ (سطح معنی‌داری کمتر از ۰.۰۱/) مقدار مجذورکای به دست آمده

معنی دار باشد، نشان می‌دهد مورد یا فرد مورد نظر دارای داده پرت چندمتغیری (در آن تعداد متغیرها) است.

مثال

در پژوهشی که بر روی کارمندان یک اداره دولتی انجام شد میزان رضایت شغلی، استرس شغلی و تعهد شغلی آنان سنجیده شد. سوالات در قالب طیف لیکرت طرح شد. می‌خواهیم وجود داده پرت چندمتغیری را در متغیرهای فوق بسنجیم.

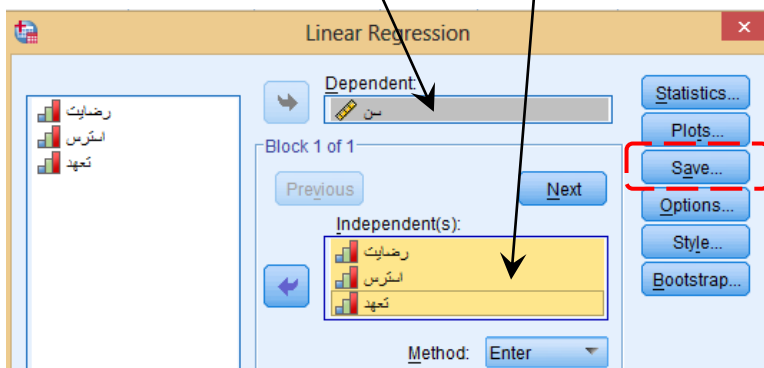
☑ نکته: دستور شناسایی داده‌های پرت چندمتغیری در فرمان رگرسیون خطی است.

اجرا:

دستور زیر را اجرا می‌کنیم:

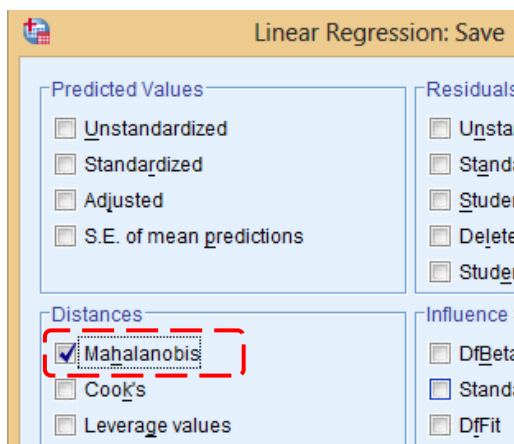
Analyze ---> Regression ---> Linear

سه متغیر رضایت شغلی، استرس شغلی و تعهد شغلی را وارد کادر Independent می‌کنیم. یک متغیر دیگر را وارد کادر Dependent می‌کنیم (سن افراد). گزینه Save را انتخاب می‌کنیم.



☑ نکته: اهمیتی ندارد که چه متغیری را به عنوان متغیر وابسته وارد کادر Dependent کنیم. تعریف یک متغیر وابسته تنها جهت اجرای رگرسیون است و غیر از سه متغیر رضایت شغلی، استرس شغلی و تعهد شغلی می‌توانیم هر متغیر دیگری را وارد کادر Dependent کنیم. در کادر Dependent، متغیری غیر از متغیرهایی که می‌خواهیم پرت بودنشان را بسنجیم، قرار می‌دهیم.

در کادر Distance گزینه Mahalanobis را فعال می کنیم.
در انتها گزینه Continue و Ok را انتخاب می کنیم.



نتیجه:

حاصل این دستورات، خروجی رگرسیون و یک متغیر جدید در فایل داده ها است. متغیر جدید « Mah_1 » نام دارد که در فایل داده ها تشکیل می شود و آخرین متغیر است.

برای ارزیابی داده های پرت چندمتغیره باید مقادیر به دست آمده برای فاصله مهالانوبیس را با توزیع مجذور کای (کای اسکور) مقایسه می کنیم (جدول توزیع مجذور کای در انتهای برخی کتب آماری یا در سایت های آماری موجود است^۴). برای این مقایسه ابتدا درجه آزادی فاصله مهالانوبیس را به دست می آوریم که از تفریق تعداد متغیرهای مستقل (وارد شده در کادر Independent رگرسیون) منهای عدد یک به دست می آید. در این مثال ما سه متغیر مستقل داریم و در نتیجه درجه آزادی برابر با ۲ است (۳ متغیر مستقل داشتیم که از عدد ۱ کم می شود و عدد به دست آمده درجه آزادی نام دارد که برابر با عدد ۲ است). مقدار مجذور کای متناظر با درجه آزادی ۲، عدد ۱۳.۸۲ است و هر مورد یا پاسخگویی که فاصله مهالانوبیس آن از عدد مذکور بیشتر باشد داده پرت محسوب می شود. برای یافتن اعداد بزرگتر از ۱۳.۸۲ در فایل داده ها و در متغیر Mah_1 بهتر است بر روی نام متغیر در پنجره Data view کلیک راست کنیم و گزینه Sort Descending را انتخاب کنیم تا داده ها از زیاد به کم مرتب شوند. حال

^۴ به عنوان مثال می توانید به سایت www.Kharazmi-Statistics.ir رجوع کنید.

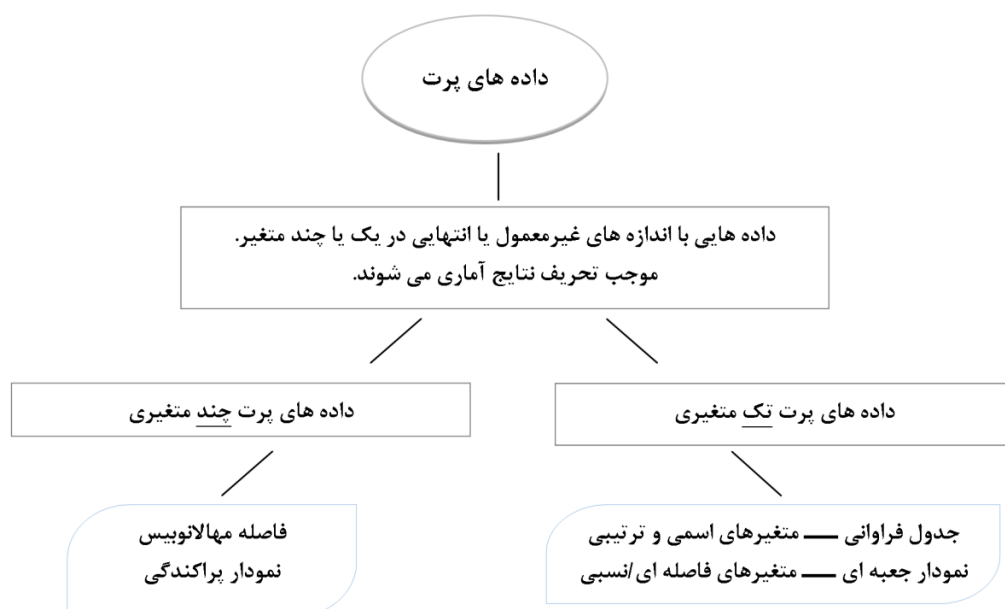
کافیست اعداد بزرگتر از ۱۳.۸۲ را پیدا کنیم و سپس مورد یا پاسخگوی مورد نظر را از فایل داده‌ها حذف کنیم.

مقدار مجذورکای برای درجه‌های آزادی ۱ تا ۸ در جدول بعد گزارش شده است. مقادیر مجذورکای ارائه شده، برای سطح آلفای ۰.۰۱ کاربرد دارد.

درجه آزادی	۸	۷	۶	۵	۴	۳	۲	۱
مقدار مجذور کای	۲۶.۱۲	۲۴.۳۲	۲۲.۴۶	۲۰.۵۲	۱۸.۴۷	۱۶.۲۷	۱۳.۸۲	۱۰.۸۳

نتایج بدست آمده نشان می‌دهد که بالاترین عدد بدست آمده برابر با ۱۰.۴۴ است که کمتر از مقدار ۱۳.۸۲ است و نشان می‌دهد داده یا موارد پرت چندمتغیری در داده‌های ما وجود ندارد.

MAH_1	var
10.44309	
10.07063	
9.02496	
9.02496	
8.70621	
8.53228	



داده های ناقص

هنگامی با مقادیر ناقص مواجه شده‌ایم که در یک یا چند متغیر، مقادیر معتبر برای تحلیل وجود نداشته باشد. تحلیل کردن داده‌های ناقص امری ضروری است چون ممکن است در تصمیم‌پذیری نتایج اثر منفی بگذارد (هیر و دیگران، ۲۰۱۰: ۴۹). داده‌های ناقص (که به داده‌های گم‌شده یا بدون پاسخ هم معروف‌اند) و روش برخورد با آن از مباحث مهم در آماده‌سازی داده‌هاست.

داده‌های ناقص سوالاتی هستند که بی پاسخ مانده‌اند و مقدار آن‌ها مشخص نیست. پاسخگوها همیشه به همه سوالات پاسخ نمی‌دهند، چون سهواً متوجه آن سوال نشده اند یا در مورد آن اطلاعی ندارند و یا پاسخ به آن را برای خود دردسرساز می‌دانند. مثلاً ممکن است برخی افراد از اعلام سن خود یا میزان درآمد خود اجتناب کنند. ممکن است پاسخگویان اطلاع دقیقی از برخی سوالات نداشته باشند، مثلاً احتمال این که برخی افراد میزان بهره هوشی خود را ندانند زیاد است. گاه نیز ممکن است پاسخگویان به دلایلی مانند عجله در پاسخ به سوالات، سهواً برخی از سوالات را بدون پاسخ بگذارند. در نتیجه وجود بی‌پاسخی در داده‌ها امری محتمل و رایج است و باید به شیوه مناسب با آن برخورد نمود.

با داده‌های ناقص چه باید کرد؟

در نخستین گام باید دانست که فرآیند گردآوری اطلاعات از پاسخگویان باید به شیوه‌ای انجام گیرد که احتمال وجود داده‌های ناقص را به حداقل ممکن برساند. با این وجود چنانچه در پژوهش خود با مسأله داده‌های ناقص مواجه شدیم می‌توانیم روش‌های زیر را به کار ببریم (جهت بحث کامل‌تر در مورد روش‌های برخورد با داده‌های ناقص می‌توانید به کتاب پیمایش در تحقیقات اجتماعی، دی. ای. دواس، ترجمه هوشنگ نایی، فصل ۱۶ مراجعه نمایید).

در برخورد با داده‌های ناقص می‌توان سه روش یا راه حل را در پیش گرفت که در ادامه به توضیح هر کدام از این سه روش پرداخته‌ایم:

(۱) چشم‌پوشی و بی‌اعتنایی به داده ناقص

۲) حذف متغیر یا پاسخگویی که دارای داده ناقص است

۳) جایگذاری داده ناقص

۱) چشم‌پوشی از مساله داده‌های ناقص

چنانچه نسبت تعداد داده‌های ناقص در مورد متغیر مدنظر بسیار ناچیز باشد می‌توان ساده‌ترین راه را انتخاب کرد و وجود داده‌های ناقص را نادیده گرفت و به ادامه تحلیل پرداخت. به عنوان یک قاعده کلی اگر داده‌های از دست رفته ۵ درصد یا کمتر باشد می‌توان آن را نادیده گرفت (تاباچینک و فیدل، ۲۰۰۱). مثلاً چنانچه ۳ نفر از ۱۰۰ پاسخگوی میزان تحصیلات خود را اعلام نکرده‌اند (و متغیر میزان تحصیلات از متغیرهای مهم پژوهش باشد) می‌توان آن‌ها را نادیده گرفت و با حجم نمونه ۹۷ نفر به ادامه تحلیل پرداخت.

۲) حذف مورد یا متغیر

الف) حذف مورد:

در این روش، چنانچه تعداد داده‌های ناقص کم باشد و سوالی که بی‌پاسخی در آن وجود دارد سوال مهمی در پژوهش به حساب بیاید، می‌توان اقدام به حذف آن فرد (یا پرسشنامه) کرد. در مثال کتاب، اگر یک یا چند نفر از افراد به سوالات مربوط به معدل کارشناسی و یا تعداد ساعات مطالعه (سوالات مهم پژوهش) پاسخ نداده باشند می‌توانیم این افراد یا پرسشنامه‌ها را به طور کلی از تحلیل حذف کنیم. تنها باید دقت کرد که حذف این افراد به کاهش چشمگیر حجم نمونه منتهی نشود و حذف موردها کمتر از ۱۵ درصد نمونه کل شود. در برنامه SPSS از دو روش برای حذف مورد استفاده می‌شود: روش حذف فهرستی^۵ (لیستی) و روش حذف زوجی^۶.

حذف فهرستی

در این روش همه افرادی که دارای داده‌های ناقص هستند از تحلیل‌های آماری معینی حذف می‌شوند. در این روش افرادی که در هریک از متغیرها، داده از دست رفته داشته باشند، حذف می‌کنیم. در این روش، یک اندازه از دست رفته‌ی تک، فقط در یک متغیر

^۵ Listwise

^۶ Pairwise

واحد سبب می شود که در مورد داده های آن فرد تحلیل های آماری انجام نشود. این روش یک عمل استاندارد در بیشتر نرم افزارهای کامپیوتری مانند SPSS است.

از مزایای حذف فهرستی این است که این روش را می توان در مورد گسترده ای از روش های چندمتغیری (برای مثال رگرسیون چندمتغیری، تحلیل عاملی و مدل یابی معادلات ساختاری) به کار برد و این روش به طور معمول به هیچ فرمان یا محاسبات اضافی نیاز ندارد. یک نگرانی آشکار در مورد این رویکرد، از دست دادن افرادی است که به سختی و با صرف هزینه (به لحاظ وقت یا منابع دیگر) درباره آنان اطلاعات جمع آوری می شود. نگرانی دیگر، کاهش حجم نمونه است که ممکن است در نمونه های نسبتاً بزرگ را که برای بیشتر روش های چندمتغیری لازم است، به نمونه ای با حجم کوچک تبدیل کند (میزر و دیگران، ۱۳۹۱: ۹۸).

حذف زوجی

در این روش، افرادی که داده های ناقص در یک یا دو متغیر خاص دارند، در صورتی که در سایر متغیرها اندازه های معتبر داشته باشند، همچنان در تحلیل باقی می ماند. بنابراین لازم نیست هیچ فردی به طور کامل از تحلیل ها حذف شود. این رویکرد خلاصه شاخص های آماری (برای مثال، میانگین ها، انحراف استانداردها، همبستگی ها) را برای همه موارد موجود که اندازه های معتبر دارند، محاسبه می کند. پیشنهاد می شود وقتی رگرسیون چندمتغیری، تحلیل عاملی یا مدل یابی معادلات ساختاری را انجام می دهیم از حذف زوجی استفاده نکنید (از روش حذف حذف فهرستی استفاده کنیم).

☑ نکته: لازم به ذکر است که برنامه SPSS به طور خودکار از روش های حذف فهرستی و زوجی استفاده می کند. بدین صورت که هنگام استفاده از آزمون های آماری و متناسب با نوع آزمون، اقدام به حذف داده های ناقص می کند. به عنوان مثال برنامه SPSS به طور پیش فرض از روش حذف زوجی برای همبستگی ها (پیرسون، اسپیرمن و...) و از روش حذف فهرستی برای آزمون های چندمتغیره (رگرسیون چندمتغیره و تحلیل عاملی) استفاده می کند و نیازی به اجرای دستور حذف فهرستی یا زوجی برای این آزمون ها نیست.

ب) حذف متغیر:

گاه برخی سوالات با بی پاسخی چشمگیری مواجه می شوند. چنانچه این سوال از سوالات مهم پژوهش نباشد (به عنوان نمونه متغیر درآمد در مثال کتاب) می توان این سوال و ستون مربوط به داده های آن در برنامه را از تحلیل نهایی حذف کرد. اما اگر

سوالی که نسبت بی‌پاسخی زیادی دارد از سوالات مهم باشد (به عنوان مثال ساعات مطالعه)، نمی‌توان سوال را حذف کرد. در این مواقع باید با روش‌های مناسب اقدام به جایگذاری داده‌های ناقص نمود و چنانچه نتوان داده‌های ناقص را به شیوه دقیقی تخمین زد و جایگذاری کرد باید دوباره پژوهش را اجرا کرده و داده‌های لازم و مناسب را جمع‌آوری کرد.

۳) جایگذاری داده ناقص:

در روش جایگذاری، داده ناقص با مقدار یا مقداری جایگزین می‌شود. این مقادیر به روش‌های مختلفی به دست می‌آیند. در ادامه سه روش رایج جایگذاری داده‌های ناقص را توضیح داده شده است.

الف) جایگذاری با میانگین متغیر

در این روش، میانگین متغیری که دارای داده ناقص است را به دست می‌آوریم و سپس این مقدار میانگین کل را جایگزین مقادیر داده‌های ناقص آن متغیر می‌کنیم. در این روش یک عدد که همان میانگین کل متغیر است، جایگزین تمامی داده‌ها ناقص می‌شود. این روش ساده‌ترین روش در بین روش‌های جایگذاری است اما روش جایگذاری با میانگین خرده گروه‌ها بر این روش برتری دارد.

ب) جایگذاری با میانگین خرده گروه‌ها

در این روش نمونه را برحسب متغیرهای زمینه‌ای و طبقه‌ای (مانند قومیت، جنس، تحصیلات) که همبستگی خوبی با متغیر دارای داده ناقص دارند به گروه‌هایی تقسیم می‌کنیم. سپس میانگین هر گروه را به جای مقادیر ناقص همان گروه می‌گذاریم. فرض کنیم که ۱۰ نفر (شامل ۳ دختر و ۷ پسر) از ۱۰۰ نفر به سوال مربوط به تعداد ساعات مطالعه پاسخ نداده‌اند. با استدلال منطقی و نیز با بررسی داده‌های خود به این نتیجه رسیدیم که دختران و پسران میزان مطالعه متفاوتی دارند و دختران به طور چشمگیری بیشتر از پسران در هفته مطالعه می‌کنند.

در این حالت، ابتدا با توجه به داده‌های موجود و اطلاعات مربوط به ۹۰ نفر باقی مانده، میانگین ساعات مطالعه دختران و پسران را به طور جداگانه محاسبه می‌کنیم. سپس مقادیر به دست آمده از میانگین ساعات مطالعه هر گروه را با توجه به جنس پاسخگو

جایگزین می‌نماییم. مثلاً اگر تعداد ساعات مطالعه دختران به طور میانگین ۵ ساعت و پسران ۳ ساعت باشد باید عدد ۵ را در گروه دخترانی که داده ناقص دارند و عدد ۳ را نیز در گروه پسرنی که داده ناقص دارند قرار بدهیم.

مثال

در پژوهشی در بین دانشجویان مقطع کارشناسی یکی از دانشگاه‌های کشور، از آنان خواسته شد که معدل کارشناسی خود را بیان کنند. تعداد ۱۰۰ نفر از دانشجویان مقطع کارشناسی به طور تصادفی انتخاب شدند که پاسخ آنان به متغیر معدل کارشناسی وارد برنامه شد. جدول فراوانی داده‌ها نشان داد که ۲ مورد داده ناقص در معدل کارشناسی وجود دارد و دو نفر به سوال مربوط به معدل پاسخ نداده‌اند. در اینجا قصد داریم ۲ مورد داده ناقص را با روش‌های مختلف جایگذاری کنیم.

۲ نفر از ۱۰۰ نفر به سوال مربوط به معدل کارشناسی پاسخ نداده‌اند، و ما با دو داده ناقص یا بدون پاسخ مواجهیم. چون درصد بی‌پاسخی کمتر از ۵ درصد کل است (۲ درصد است)، در نتیجه می‌توانیم از داده‌های بدون پاسخ چشم‌پوشی کرده و با باقی داده‌ها که ۹۸ مورد می‌شود به تحلیل‌ها ادامه دهیم. در اینجا چون هدف، آموزش جایگزینی داده‌های ناقص با مقادیر معتبر است اقدام به جایگزینی دو مورد فوق با روش‌های ذکر شده می‌کنیم. در اینجا با دو روش جایگزینی با میانگین متغیر و با میانگین خرده گروه‌ها اقدام به جایگزینی داده‌های بدون پاسخ می‌کنیم.

روش اول: جایگزینی با میانگین متغیر

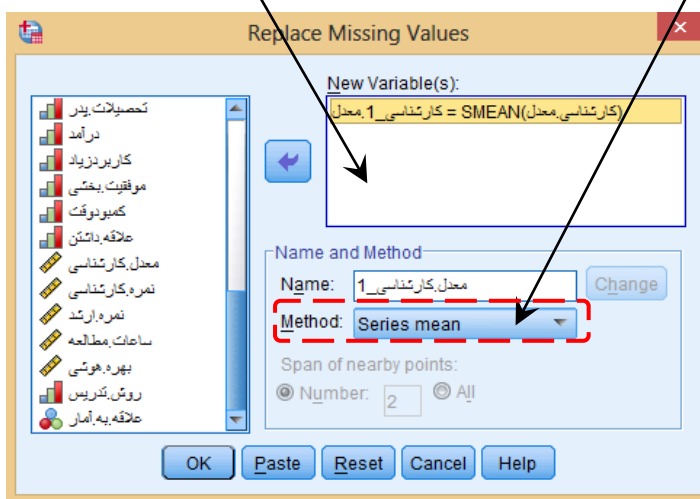
در این روش میانگین کل معدل کارشناسی ۹۸ نفری که به این سوال پاسخ داده‌اند محاسبه شده و جایگزین داده‌های ناقص می‌شود. در این روش از دستور جایگذاری مقادیر ناقص استفاده می‌کنیم.

اجرا:

دستور زیر را اجرا می‌کنیم:

Transform ---> Replace Missing Values

متغیر معدل کارشناسی را از کادر سمت چپ وارد این کادر می‌کنیم. متد یا روش جایگزینی، روش میانگین سری (Series Mean) است که در آن میانگین داده‌های معتبر محاسبه شده و جایگزین داده‌های بدون پاسخ می‌شود. با انتخاب گزینه OK جایگزینی میانگین به طور خودکار انجام می‌شود و متغیری جدید با همان نام و با پسوند ۱_ در انتهای فایل داده‌ها ایجاد می‌شود.



نتیجه:

با اجرای دستور جایگزینی مقادیر ناقص، متغیری جدید در انتهای فایل داده‌ها تشکیل می‌شود که در آن داده ناقصی وجود ندارد و مقادیر ناقص با میانگین کل معتبر جایگزین شده‌اند. میانگین جایگزین شده برابر با ۱۶.۴۷ است.

(ب) جایگزینی با میانگین خرده‌گروه‌ها:

در این روش باید متغیرهای طبقه‌ای را که با متغیر معدل کارشناسی دارای رابطه هستند را بر اساس مبانی نظری، پژوهش‌های پیشین و یا استدلال منطقی بباییم. همچنین می‌توانیم اقدام به مقایسه معدل کارشناسی دانشجویان برحسب طبقات هر کدام از متغیرهای جنس، تحصیلات، قومیت و غیره نمود. فرض می‌کنیم که معدل کارشناسی دانشجویان دختر و پسر متفاوت است و مطابق استدلال نظری و یافته‌های قبلی، دختران معدل بالاتری در مقایسه با پسران کسب می‌کنند. بر این اساس، میانگین

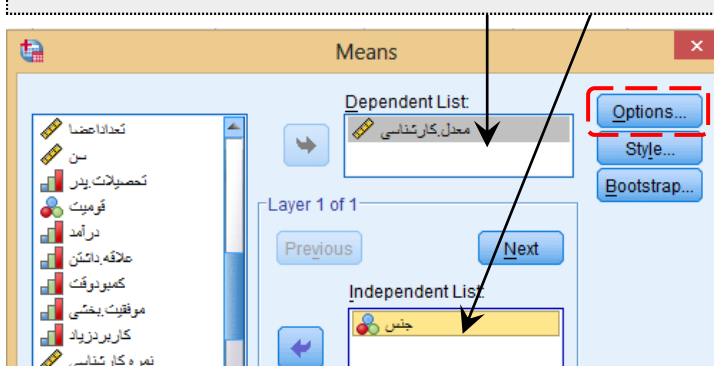
معدل کارشناسی دختران و پسران را به طور مجزا محاسبه کرده و چنانچه داده ناقص مربوط به دانشجوی دختر باشد از میانگین معدل کارشناسی دانشجویان دختر استفاده می کنیم و چنانچه داده ناقص مربوط به دانشجوی پسر باشد از میانگین معدل کارشناسی دانشجویان پسر استفاده می کنیم و میانگین ها را جایگزین می کنیم. همچنین می توانیم میانگین ها را برحسب دو یا چند متغیر به دست بیاوریم. مثلا می توانیم معدل مقطع کارشناسی را برحسب جنس و تحصیلات به دست آورد که برآورد دقیق تری از معدل کارشناسی خواهد بود.

اجرا:

دستور زیر را اجرا می کنیم. از این دستور برای محاسبه جداگانه میانگین معدل کارشناسی دانشجویان دختر و پسر استفاده می کنیم.

Analyze ---> Compare Means ---> Means

متغیر معدل کارشناسی را از کادر سمت چپ وارد کادر متغیر وابسته می کنیم. متغیر جنس که قصد داریم بر اساس آن میانگین معدل کارشناسی را جداگانه بدست بیاوریم وارد کادر متغیر مستقل می کنیم. در انتها گزینه OK را انتخاب می کنیم.



نتایج:

جدول زیر خروجی آزمون محاسبه میانگین به تفکیک گروه هاست. میانگین معدل کارشناسی دختران و پسران در جدول گزارش شده است. نتایج نشان می دهد بین معدل کارشناسی دختران و پسران تفاوت قابل توجهی وجود دارد. میانگین معدل دانشجویان پسر برابر با ۱۶.۱۸ و دانشجویان دختر برابر با ۱۶.۷۶ است.

با توجه به یافته‌های فوق، ما در هنگام جایگزینی داده‌های ناقص، میانگین‌ها را متناسب با جنس دانشجویان جایگزین می‌کنیم. بدین صورت که ابتدا دانشجویی که معدل کارشناسی خود را اعلام نکرده پیدا می‌کنیم و جنس او را شناسایی می‌کنیم. اگر پسر باشد عدد ۱۶.۱۸ را در جای خالی قرار می‌دهیم و اگر دانشجوی مورد نظر دختر باشد معدل ۱۶.۷۶ را به جای داده ناقص قرار می‌دهیم. در روش جایگزینی با میانگین متغیر، عدد ۱۶.۴۷ جایگزین داده ناقص دانشجویان دختر و پسر شده بود در حالی که با توجه به جنس دانشجویان، میانگین پسران کمتر از میانگین کل و میانگین دختران بالاتر از میانگین کل است و اگر جایگزینی معدل کارشناسی با در نظر گرفتن جنس دانشجویان انجام شود دقت جایگزینی افزایش می‌یابد.

Report

معدل کارشناسی

	Mean	N	Std. Deviation
Jens	میانگین	تعداد یا فراوانی	انحراف استاندارد
دختران	16.7610	49	1.847
پسران	16.1827	49	1.320
Total	16.4718	98	1.623

کدگذاری مجدد متغیرها

کدگذاری مجدد، به معنای تغییر کدهایی است که در ابتدا به متغیرها و طبقات داده‌ایم و در موارد متعددی نیازمند این هستیم که کدهایی که ابتدا اختصاص داده‌ایم را تغییر دهیم. کدگذاری مجدد یک از پرکاربردترین مراحل آماده‌سازی داده‌ها محسوب می‌شود. برای مقاصد زیر متغیرها را کدگذاری مجدد می‌کنیم:

۱- تبدیل مقادیر فاصله‌ای به مقادیر ترتیبی: برای مثال می‌توانیم معدل افراد

را از حالت فاصله‌ای/نسبی خارج کرده و در سه دسته معدل پایین (زیر ۱۴)، معدل متوسط (۱۴ تا ۱۷) و معدل بالا (بالای ۱۷) قرار دهیم. یا می‌توانیم افراد را برحسب درآمدشان در سه طبقه کم‌درآمد، میان‌درآمد و پردرآمد قرار دهیم.

۲- تقلیل و کاهش تعداد طبقات: برخی از مصادیق کاربرد این شیوه عبارت

است از: الف) بخواهیم افراد را بر حسب تحصیلات به جای پنج طبقه در دو طبقه قرار دهیم: دیپلم و پایین‌تر و طبقه دارای تحصیلات دانشگاهی. ب) بخواهیم طبقات گویه‌های طیف لیکرت را از پنج طبقه به سه طبقه تقلیل بدهیم. در این حالت، پنج طبقه کاملاً مخالفم، مخالفم، بی‌نظر، موافقم و کاملاً موافقم را در سه طبقه مخالفم، بی‌نظر و موافقم خلاصه می‌کنیم. ج) زمانی که پاسخ‌های ارائه شده به طبقه‌ای خیلی اندک است بهتر است آن طبقه را با طبقه دیگری ادغام کنیم. فراوانی بسیار پایین ممکن است جدول را گمراه کننده و نتایج و تحلیل‌ها را دچار تحریف کند.

۳- یکسان‌سازی جهت سوالات: زمانی که در سنجش یک مفهوم یا سازه با

سوالات منفی و مثبت به صورت همزمان مواجه هستیم باید سوالات مثبت و منفی را هم‌جهت کنیم. فرض کنیم برای سنجش اعتماد بین فردی سه سوال در قالب طیف لیکرت سه گزینه‌ای (مخالفم، تاحدودی و موافقم) طرح کرده‌ایم که شامل دو سوال مثبت است که با این جمله‌بندی مطرح شده است: «به اعضای خانواده‌ام اعتماد دارد» و «به بستگانمان اعتماد دارد» و یک سوال منفی است که با این جمله‌بندی طرح شده است: «به همسایگانمان کاملاً بی‌اعتمادم». در هنگام ساختن مقیاس باید شیوه نمره‌دهی سوالات مثبت و منفی را یکسان و هم‌جهت کرد تا نمره کل هر فرد در مقیاس تفسیرپذیر باشد.

در سوالات مثبت اعتماد (اعتماد به اعضای خانواده و بستگان)، انتخاب گزینه موافقم به معنای اعتماد داشتن است اما در سوال اعتماد به همسایگان که یک سوال منفی است، انتخاب گزینه موافقم به معنای بی‌اعتمادی است که در این جا باید جهت نمره‌دهی این سوال منفی را با سوالات مثبت یکسان کرد و کدگذاری مجدد کرد.

مثال

در مثال کتاب، معدل کارشناسی در سطح فاصله‌ای سنجیده شده است که ما قصد داریم پاسخگویان را برحسب معدل کارشناسی‌شان در سه طبقه قرار دهیم:

معدل پایین: شامل معدل کمتر از ۱۴

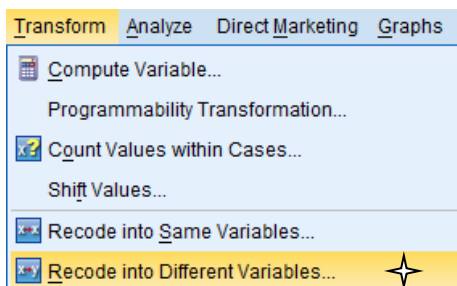
معدل متوسط: شامل معدل ۱۴ تا ۱۷

معدل بالا: شامل معدل بالاتر از ۱۷

اجرا:

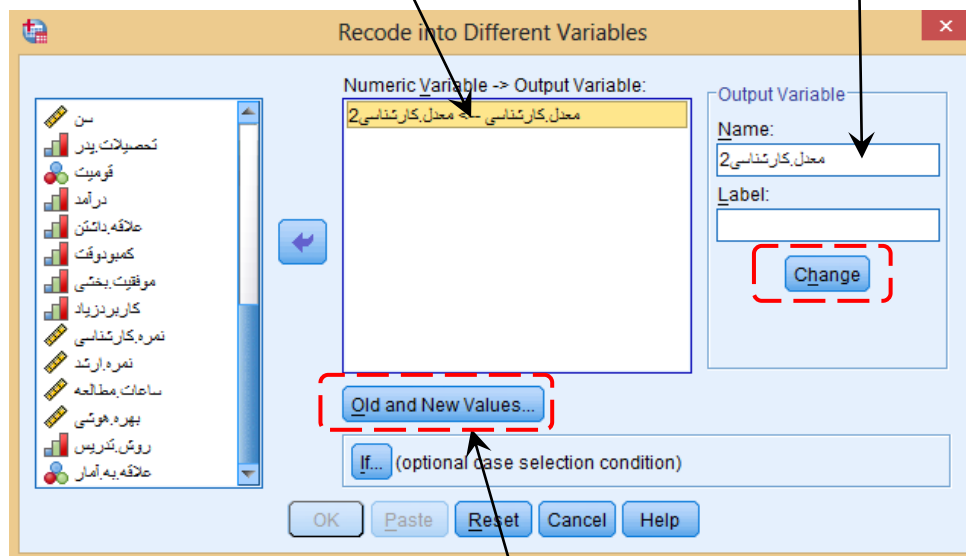
دستور زیر را اجرا می‌کنیم:

Transform ---> Recode into Different Variables



(۱) ابتدا متغیر موردنظر را از کادر سمت چپ به این کادر منتقل می کنیم.

(۲) در کادر Name نام جدیدی برای متغیر جدید می سازیم. گزینه Change را انتخاب می کنیم تا نام جدید اعمال شود.



گزینه کدهای قدیم و جدید (Old and New Values) را انتخاب می کنیم.

در پنجره کدها یا مقادیر جدید و قدیم (Old and New Values) مطابق شکل تغییرات را اعمال می کنیم. این پنجره برای آسانی آموزش به دو بخش تقسیم شده است.

در پنجره مقادیر جدید و قدیم برحسب دامنه نمراتی که تعیین کرده ایم کدهای تغییر یافته هر طبقه را وارد می کنیم:

۱۰ تا ۱۳.۹۹ = کد ۱

۱۴ تا ۱۷ = کد ۲

۱۷.۰۱ تا ۲۰ = کد ۳

به عنوان مثال برای طبقه آخر یا کد ۳ (و در کادر کدهای قدیم یا Old Value) به ترتیب زیر عمل می کنیم:

گزینه مقدار یا کد (Value) هنگامی به کار می رود که بخواهیم یک عدد را با یک عدد دیگر جایگزین کنیم، مثلاً بخواهیم کد ۵ را با کد ۱ جایگزین کنیم. اگر بخواهیم مانند مثال کتاب، دامنه ای از نمرات را با یک عدد جایگزین کنیم، از دستور دامنه نمرات (Range) استفاده می کنیم.

در کادر کدها یا مقادیر قبلی (Old Value) گزینه Range (دامنه نمرات) را انتخاب می کنیم و در کادرهای ایجاد شده عدد ۱۷.۰۱ تا (through) عدد ۲۰ را می نویسیم. در واقع دامنه نمرات مدنظر هر طبقه را مشخص می کنیم.

بعد از انجام دستورات فوق، وارد کادر کد جدید (New Value) می شویم.

در کادر Value کد جدید را که عدد ۳ است می نویسیم. سپس بر روی گزینه Add کلیک می کنیم تا طبقه بندی جدید اعمال شود. این مراحل را برای دو طبقه دیگر هم انجام می دهیم. بعد از تعیین تمامی کدهای جدید، گزینه های Continue و سپس OK را انتخاب می کنیم.

☑ نکته: هیچ کدام از طبقات نباید با هم همپوشانی داشته باشند و دامنه نمرات هر کدام از طبقات باید از یکدیگر متمایز باشد. مثلاً نمی‌توانیم عدد ۱۷ را در دو طبقه قرار بدهیم. باید در یک طبقه که طبقه دوم است مقدار ۱۷ و در طبقه دیگر مقدار بزرگتر از ۱۷ (یعنی عدد ۱۷.۰۱) قرار داده شود.

نتایج:

با اجرای دستور بالا متغیری جدید در پنجره داده‌ها تشکیل می‌شود که نمرات و کدهای آن تغییر کرده است و مطابق کدها و طبقاتی است که ما مشخص کرده‌ایم.

مقایسه دو روش کدگذاری مجدد: برای کدگذاری مجدد متغیرها دو روش در برنامه SPSS وجود دارد. که دستور این دو در کنار یکدیگر بوده و در جدول زیر به مقایسه این دو روش پرداخته شده است.

شکل ۲-۴- جدول مقایسه روش‌های کدگذاری

در این روش متغیر جدیدی ساخته نمی‌شود و تغییرات روی همان متغیر قبلی انجام می‌شود. در نتیجه در این روش داده‌های اولیه از بین رفته و کدهای جدید، جایگزین کدهای قبلی می‌شود و داده‌های جدید تشکیل می‌شود.	روش کدگذاری روی متغیر موجود Recode into Same Variables
ستونی جدید و متغیری جدید در پنجره داده‌ها تشکیل می‌شود. در این روش هم متغیر اولیه باقی مانده و هم متغیری جدید در برنامه تشکیل می‌شود. به جهت حفظ داده‌های اولیه، توصیه می‌گردد که همواره از این روش استفاده کنیم.	روش کدگذاری با ساخت متغیر جدید Recode into Different Variables

مقیاس سازی

مقیاس‌ها در علوم انسانی اهمیت و کاربرد فراوانی دارند. پایه و اساس مقیاس‌سازی جزئی از زندگی روزمره است. وقتی برای اولین بار با شخصی روبرو می‌شویم سعی می‌کنیم تصویری از او بسازیم: تصویری از درستی، هوش و ذکاوت و قابل اعتماد بودن او و غیره. این تصورات به ندرت بر مبنای فقط یک مورد اطلاع در مورد این شخص است، بلکه تصویری است ترکیبی مبتنی بر چند سرنخ و نشانه.

مقیاس‌هایی که در پژوهش‌های علوم انسانی به کار می‌روند فقط شکل صورت‌بندی شده و منظم‌تر همین عمل روزمره است. مقیاس، اندازه ترکیبی یک مفهوم است: اندازه‌ای مرکب از داده‌های برآمده از چند سوال و معرف. در ایجاد مقیاس فقط اطلاعات ارائه شده در چند متغیر نسبتاً معین را به متغیر تازه و انتزاعی‌تر تبدیل می‌کنیم. مثلاً در سنجش دین‌دار بودن افراد به جای این که از آن‌ها بپرسیم آیا دین‌دار هستید یا نه، چندین سوال (که نشان دهنده دین‌دار بودن است) از افراد می‌پرسیم. سپس نمرات این سوالات را ترکیب کرده و به صورت مقیاس دین‌داری درمی‌آوریم.

هر مقیاس حاوی پاسخ‌هایی به چند سوال است. پاسخگویان بسته به پاسخ‌شان نمره‌ای در سوالات می‌گیرند. سپس نمرات همه سوال‌های فرد با هم جمع می‌شود و نمره کل فرد به دست می‌آید (نمره مقیاس). این نمره مقیاس نشان‌دهنده جایگاه و موقعیت پاسخگو در آن مفهوم یا مطلب انتزاعی است.

در مثال کتاب، مقیاسی به نام شایستگی علمی اساتید داریم که با تعدادی سوال سنجیده شده است. از آن جا که شایستگی علمی اساتید مفهومی گسترده و چندبُعدی بوده و موارد و مصداق‌های متعددی را شامل می‌شود با تعدادی سوال آن را می‌سنجیم. برخی از سوالات آن در جدول ۲-۵ آمده است. در این مقیاس، هر فرد به سوالات مربوط پاسخ داده و در هر مورد دور عددی که مدنظرش است خط می‌کشد و در هر مورد یک نمره کسب می‌کند. در انتها نمره کل فرد که مجموع نمرات فرد در تمامی سوالات است محاسبه شده و به عنوان نمره فرد در مقیاس شایستگی استاد منظور می‌شود.

به عنوان نمونه، فردی به پنج سوال مطرح شده در ارتباط با شایستگی اساتید پاسخ داده است و مجموع نمرات فرد برابر با ۲۰ شده است و نشان دهنده نمره فرد در مقیاس شایستگی علمی اساتید است. چنانچه مقدار بیست را بر تعداد سوالات که پنج است تقسیم کنیم میانگین نمره فرد در مقیاس شایستگی علمی اساتید به دست می آید که برابر با ۴ است (توجه داریم که میانگین نمرات فرد در مقیاس از ۰ تا ۱۰ است). با به دست آوردن نمره تمامی پاسخگویان می توانیم نمره شایستگی علمی افراد (نمره‌ای که هر فرد به استاد مربوطه داده است) را با هم مقایسه کنیم و همچنین می توانیم بررسی کنیم که پاسخگویان در مجموع چه نمره‌ای در شایستگی علمی اساتید دارند.

جدول ۲-۵

در هر کدام از ویژگی‌های زیر، از نمره ۰ تا ۱۰، چه نمره‌ای به استاد درس آموزش SPSS می‌دهید؟

ویژگی‌های استاد	نمره (دور نمره مدنظر خود خط بکشید)
۱. تسلط علمی	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸ ۹ ۱۰
۲. ارائه جذاب مباحث درسی	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸ ۹ ۱۰
۳. استقبال از سوالات دانشجویان	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸ ۹ ۱۰
۴. ارائه تکالیف درسی مناسب	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸ ۹ ۱۰
۵. استفاده از وسایل کمک آموزشی (ویدئو پرژکتور، پاورپوینت)	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸ ۹ ۱۰

مقیاس سازی در SPSS را می توان به دو صورت انجام داد: جمع کردن و میانگین گیری.

در روش جمع کردن، نمرات پاسخگویان در تمامی سوالات با یکدیگر جمع می شود و نمره حاصل شده که مجموع نمرات پاسخگو در تمامی سوالات مقیاس است به عنوان نمره مقیاس پاسخگو در نظر گرفته می شود. در روش میانگین گیری بعد از جمع کردن تمامی نمرات، نمره حاصل شده بر تعداد سوالات تقسیم می شود. در مثال بالا که مربوط به شایستگی علمی اساتید است نمرات مجموع و میانگین به ترتیب ۲۰ و ۴ است.

روش میانگین‌گیری بر جمع کردن برتری دارد. در روش جمع‌کردن، دامنه نمرات مقیاس متکی بر تعداد سوالات (یا متغیرهای) تشکیل دهنده آن است. در نتیجه، تفسیر نمرات به ویژه در مقایسه با نمرات مقیاس‌های دیگر سخت می‌شود. اما در روش میانگین‌گیری با تقسیم نمره مجموع بر تعداد سوالات، نمره تعدیل شده و در دامنه مشخصی قرار می‌گیرد و می‌توان آن را با مقیاس‌های دیگر مقایسه کرد.

همچنین در روش جمع‌کردن، اگر حتی یکی از سوالات تشکیل دهنده مقیاس ناقص باشد و فرد به یکی از سوالات پاسخ نداده باشد، در برنامه SPSS نمره مقیاس برای آن فرد ساخته نمی‌شود و فرد در آن مقیاس نمره ای نمی‌گیرد. اما در روش میانگین‌گیری، مقیاس‌سازی بر اساس تعداد سوالات معتبر (فاقد داده ناقص) ساخته می‌شود، در نتیجه حتی اگر فرد فقط به یکی از سوالات نیز پاسخ داده باشد نمره مقیاس برای او ساخته خواهد شد. در نتیجه ما در روش میانگین‌گیری با مسأله ناقص بودن و عدم محاسبه نمره مقیاس مواجه نخواهیم شد.

مثال

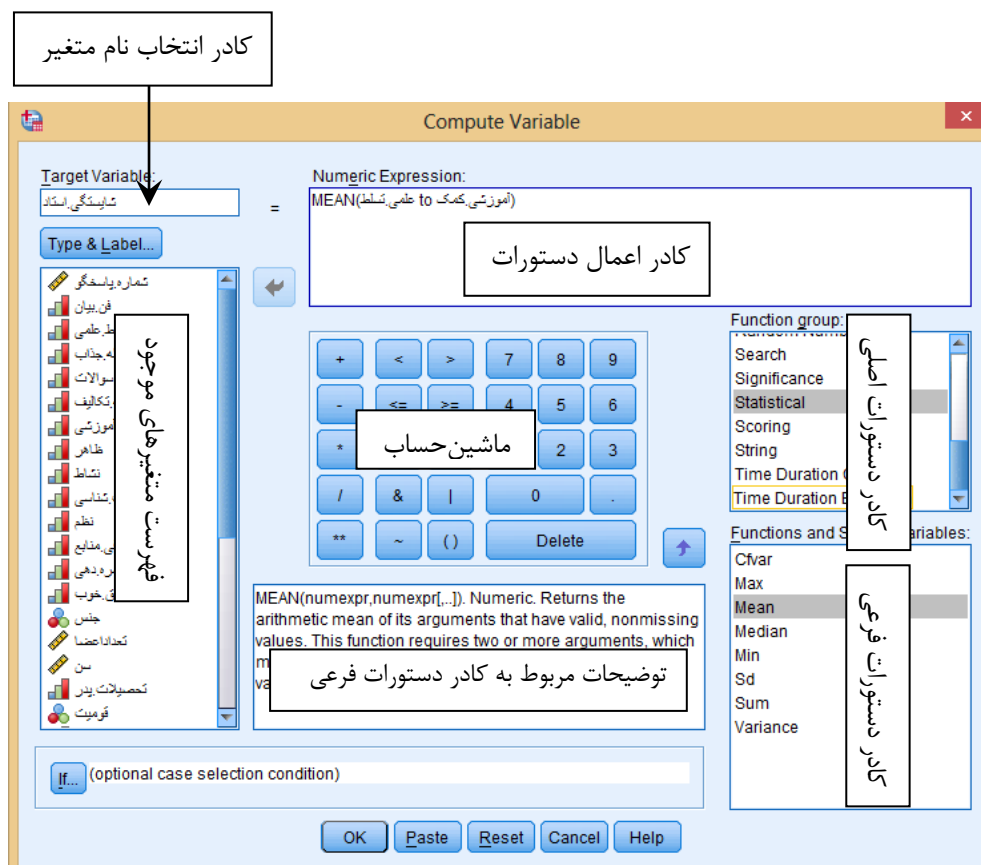
محاسبه نمره مقیاس شایستگی علمی استاد با روش میانگین‌گیری:
ما شایستگی علمی استاد را با ۱۳ سوال سنجیده‌ایم و می‌خواهیم بدانیم که هر کدام از پاسخگویان در این مقیاس چه نمره‌ای می‌گیرند.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Transform ---> Compute Variable

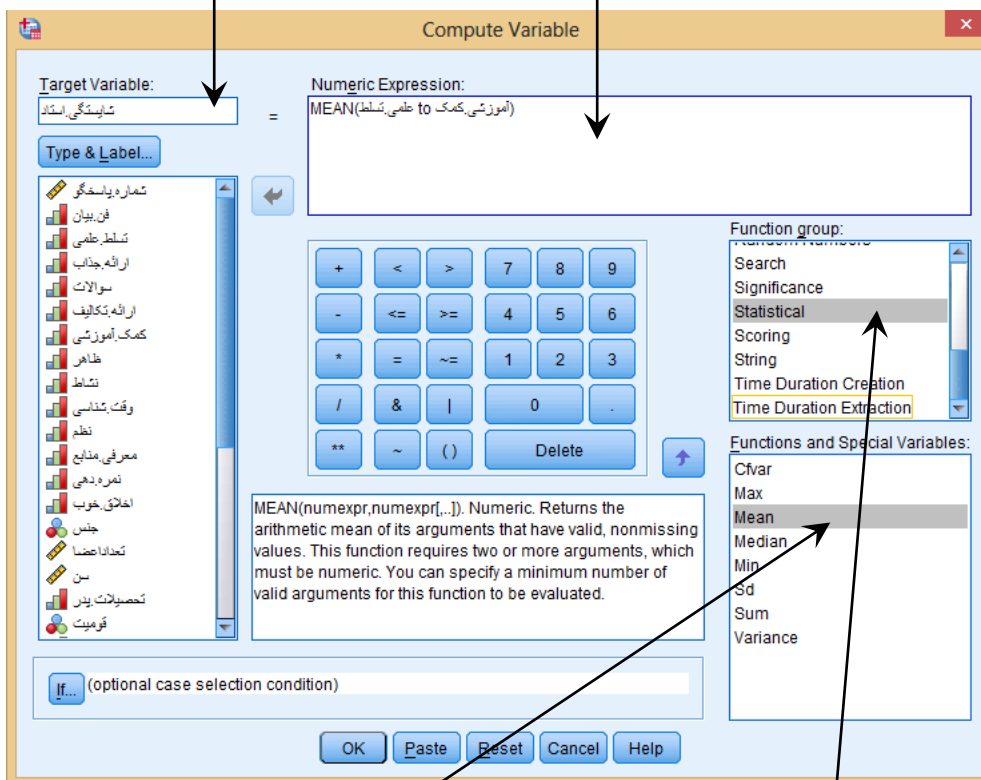
با اجرای دستور بالا، پنجره Compute تشکیل می‌شود که در ادامه به توضیح هر کدام از بخش‌های آن پرداخته‌ایم:



در ادامه با ذکر مثال به چگونگی محاسبه میانگین متغیر شایستگی استاد می پردازیم:

۴- در ابتدای کار کل کلمات و دستورات داخل پرانتز را پاک می‌کنیم. در مرحله بعد ابتدا اولین متغیر (تسلط علمی) را وارد پرانتز می‌کنیم. سپس یکبار کلید Space صفحه کلید را می‌زنیم تا یک فاصله ایجاد شود، سپس کلمه to را می‌نویسیم. در ادامه مجدداً کلید Space را یکبار می‌زنیم و متغیر آخر (وسایل کمک آموزشی) را وارد می‌کنیم. در انتها گزینه OK را انتخاب می‌کنیم.

۱- در ابتدا نامی برای متغیر جدیدی که قصد داریم ایجاد کنیم انتخاب می‌کنیم.



۳- در این کادر دستور Mean یا میانگین گرفتن را انتخاب می‌کنیم و با استفاده از دستور  به کادر اعمال دستورات منتقل می‌کنیم

۲- در این کادر گزینه Statistical یا دستورات آماری را انتخاب می‌کنیم.

نتیج:

با اجرای دستور میانگین‌گیری متغیر جدیدی به نام شایستگی استاد در پنجره داده‌ها تشکیل می‌شود که در واقع نشان‌دهنده نمره مقیاس هر کدام از پاسخگویان است. دامنه این متغیر مانند دامنه هر کدام از سوالات بوده و از ۰ تا ۱۰ است.

برای میانگین‌گیری به غیر از روش فوق روش دیگر نیز وجود دارد. در این روش در کادر Numeric Expression و در داخل پرانتز مربوط به دستور میانگین مراحل زیر را اجرا می‌کنیم:

چنانچه متغیرها در پنجره داده‌ها پشت‌سرهم نیامده باشند و بین آن‌ها متغیرهای دیگری باشد باید از این روش استفاده کرد: فرض کنید می‌خواهیم تنها با سه متغیر A و B و C مقیاس جدید بسازیم و این سه متغیر در پنجره داده‌ها به دنبال هم نیامده‌اند و بین این سه متغیر متغیرهای دیگری هم آمده است. در این روش در داخل پرانتز ابتدا متغیر A را وارد می‌کنیم و سپس بدون فاصله بعد از آن ویرگول را قرار می‌دهیم و سپس متغیر B را وارد کرده و ویرگول (ویرگول انگلیسی) را قرار می‌دهیم و در انتها متغیر C را وارد می‌کنیم:

Mean(A,B,C)

☑ نکته: چنانچه بخواهیم سوالات یا متغیرها را با هم جمع کنیم و نمره مجموع آن‌ها را حساب کنیم، کافیست در کادر Functions and Special Variables به جای گزینه Mean، گزینه Sum (مجموع) را وارد کادر Numeric Expression کنیم و سایر مراحل را مانند میانگین‌گیری ادامه دهیم. در این وضعیت تمامی سوالات با هم جمع می‌شوند اما بر تعداد سوالات تقسیم نمی‌شوند و نمره به دست آمده نمره مجموع تمامی سوالات است.

تعریف

اطلاعات زیر را وارد برنامه کنید:

وزن	میزان ورزش	تحصیلات	سن	جنس	مورد	وزن	میزان ورزش	تحصیلات	سن	جنس	مورد
۵۸	۳۶۰	۱	۲۸	۲	۱۶	۵۵	۳۰۰	۱	۲۰	۱	۱
۶۷	۳۳۰	۱	۳۱	۲	۱۷	۵۱	۲۸۰	۱	۲۴	۱	۲
۷۳	۳۵۰	۱	۳۲	۲	۱۸	۵۷	۲۷۰	۱	۲۳	۱	۳
۷۹	۲۷۰	۲	۳۵	۲	۱۹	۶۱	۲۵۰	۱	۲۹	۱	۴
۸۵	۲۹۰	۲	۳۵	۲	۲۰	۶۲	۲۳۰	۲	۳۵	۱	۵
۸۰	۲۴۰	۲	۳۶	۲	۲۱	۶۵	۲۰۰	۲	۴۵	۱	۶
۸۱	۲۴۰	۲	۳۸	۵	۲۲	۲۱	۲۰۰	۲	۴۲	۱	۷
۸۶	۲۱۰	۲	۳۹	۲	۲۳	۶۹	۲۱۰	۲	۴۴	۱	۸
۸۹	۲۰۰	۲	۴۲	۲	۲۴	۶۸	۱۸۰	۲	۵۱	۱	۹
—	۱۸۰	۲	۵۲	۲	۲۵	۷۱	۱۰۰	۲	۳۲	۱	۱۰
۸۸	۱۶۰	۲	۴۱	۲	۲۶	۷۲	۱۲۰	۲	۴۱	۱	۱۱
۹۳	۱۵۰	۳	۵۲	۲	۲۷	۷۴	۵۰	۳	۴۷	۱	۱۲
۹۷	۱۲۰	۳	۵۵	۲	۲۸	۷۶	۵۰	۳	۴۱	۱	۱۳
۹۵	۶۰	۳	۵۶	۲	۲۹	۷۶	۳۰	۳	۴۴	۱	۱۴
۹۹	۳۰	۳	۶۴	۲	۳۰	۹۴	۳۰	۳	۵۵	۱	۱۵

نکاتی در ارتباط با جدول فوق:

نکته: در جدول بالا، در متغیر جنس کد ۱ نشان دهنده زنان و کد ۲ نشان دهنده مردان است.

نکته: در متغیر میزان تحصیلات، کد ۱ نشان دهنده میزان تحصیلات دیپلم و پایین‌تر، کد ۲ نشان دهنده میزان تحصیلات کاردانی و کارشناسی و کد ۳ نشان دهنده میزان تحصیلات کارشناسی ارشد و دکترا است.

نکته: میزان ورزش برحسب دقیقه در طول یک هفته است و وزن بر حسب کیلوگرم گزارش شده است.

تمرین:

۱- متغیرهای جنس و وزن را از نظر وجود داده‌های پرت بیازمایید (داده پرت متغیر جنس را با عدد ۲ جایگزین کنید).

۲- بعد از یافتن داده پرت متغیر وزن، دو داده ناقص متغیر وزن که یکی در مردان و یکی در زنان است را با هر کدام از روش‌های زیر جایگزین کنید:

الف) ابتدا میانگین کل متغیر وزن را جایگزین داده ناقص کنید (روش میانگین متغیر).

ب) میانگین داده‌های ناقص را برحسب متغیر جنس افراد به دست آورده و جایگزین کنید (روش میانگین خرده گروه‌ها).

۳- متغیر میزان ورزش را کدگذاری مجدد کرده و افراد را برحسب میزان ورزش در طول هفته (به دقیقه) در سه دسته قرار دهید: الف) فعالیت زیاد: بیشتر از ۲۴۰ دقیقه، ب) فعالیت متوسط: بین ۱۲۰ تا ۲۴۰ دقیقه و ج) فعالیت کم: کمتر از ۱۲۰ دقیقه.

۴- افراد را برحسب میزان تحصیلات در دو گروه قرار دهید قرار دهید و کدگذاری مجدد کنید:

الف) تحصیلات غیردانشگاهی: دیپلم و فوق دیپلم، ب) تحصیلات دانشگاهی: کاردانی، کارشناسی، کارشناسی ارشد و دکترا

۵- به جدول صفحه ۱۸۸ مراجعه کرده و مقیاس اعتماد اجتماعی را با روش میانگین‌گیری بسازید. اعتماد اجتماعی از ترکیب چهار شاخص به دست آمده است: اعتماد به اعضای خانواده، اعتماد به بستگان نزدیک، اعتماد به همسایگان و اعتماد به همکاران.



واژه نامه فصل دوم			
Count	شمارش	Decimal	اعشار
Classification	طبقه بندی	Label	برچسب
Numeric	عددی	Output window	پنجره برون داد
Width	عرض	Align	ترازبندی - چیدمان
Mahalanobis Distance	فاصله مَهالانوبیس	Functions	توابع
Recode	کدگذاری مجدد	Replace	جایگزینی
Compute	محاسبه کردن	Pairwise	حذف زوجی
Sort cases	مرتب کردن موارد	List wise	حذف فهرستی
Value	مقدار، ارزش، کد	String	حرفی، الفبایی
Scale	مقیاس	Grid Lines	خطوط چهارخانه
Box plot	نمودار جعبه ای	Outlier data	داده پرت
		Missing data	داده های ناقص (گمشده)

توصیف داده‌ها و پیش‌فرض‌های آماری

محتوای فصل

«بخش اول: توصیف داده‌ها: توزیع فراوانی، شاخص‌های گرایش به مرکز،

شاخص‌های پراکندگی، نمودارها»

«بخش دوم: پیش‌فرض‌های آماری: آزمون نرمال بودن داده‌ها، خطی بودن

رابطه، هم‌خطی چندگانه، یکسانی پراکندگی، تبدیل داده‌ها»

«بخش سوم: تقسیم یا گزینش داده‌ها: تقسیم داده‌ها، انتخاب موردها»

بخش اول:

توصیف داده‌ها

آمارهای توصیفی، مجموعه‌ای از ابزارهای آماری هستند که توصیف دقیق حجم زیادی از داده‌ها را ممکن می‌سازند. آمارهای توصیفی رایج عبارتند از شاخص‌های توزیع فراوانی، اندازه‌های گرایش به مرکز، اندازه‌های پراکندگی و نمودارها. هر گزارش پژوهش همیشه باید شامل آمار توصیفی باشد. برای دادن اطلاعات در مورد نمونه و برای توصیف داده‌ها قبل از انجام آزمون‌های آمار استنباطی باید مورد استفاده قرار گیرد (بریس و همکاران: ۱۳۹۱: ۱۰۰). جدول زیر رایج‌ترین و پرکاربردترین روش‌هایی را که در توصیف داده‌ها مورد استفاده قرار می‌گیرد نشان می‌دهد.

جدول ۳-۱- رایج‌ترین روش‌های توصیف داده‌ها (آمار توصیفی)

نوع آمار	نوع توصیف	آماره‌های رایج
	توزیع فراوانی‌ها	فراوانی - درصد فراوانی - چارک*
	شاخص‌های گرایش به مرکز	میانگین* - میانه - مد
آمار توصیفی	شاخص‌های پراکندگی	انحراف استاندارد* - واریانس* دامنه تغییرات* - ضریب تغییرات*
	نمودارها	دایره ای - ستونی - هیستوگرام*

☑ نکته: آماره‌های ستاره‌دار (*) مختص داده‌های فاصله‌ای/نسبی است و نمی‌توان آن‌ها را برای داده‌های اسمی/ترتیبی به کار برد.

☑ نکته: تمامی آماره‌های بالا و نمودارها (به غیر از ضریب تغییرات) را می‌توان از طریق دستور Frequencies در منوی Analyze به دست آورد.

توزیع فراوانی‌ها

(۱) فراوانی و درصد فراوانی

اولین کاری که بعد از گردآوری و وارد کردن داده‌ها انجام می‌دهیم شمارش تعداد افرادی است که پاسخ‌های معینی به هر پرسش داده‌اند. با این کار نحوه پراکندگی یا توزیع نمونه را در طبقات مختلف هرمتغیر بررسی می‌کنیم. نتیجه این شمارش توزیع فراوانی است.

توزیع فراوانی به طور متداول شامل فراوانی و درصد فراوانی است که عموماً برای متغیرهای اسمی یا ترتیبی به کار می‌رود. کاربرد آن برای متغیرهای فاصله‌ای/نسبی فقط زمانی که تعداد طبقات متغیر کم باشد و یا متغیر کمی را به متغیری ترتیبی تبدیل کنیم، مناسب است. مثلاً در مورد تعداد اعضاء خانواده (بعد خانوار)، استفاده از آماره فراوانی فقط زمانی مناسب که است تعداد اعضا شامل حدود ۶ طبقه باشد. همچنین زمانی که درآمد افراد را به صورت عددی می‌سنجیم ولی درآمد افراد را کدگذاری مجدد کرده و در پنج طبقه قرار بدهیم می‌توانیم از توزیع فراوانی برای متغیر درآمد بهره بگیریم. در مجموع می‌توان گفت که در وضعیتی که طبقات تشکیل‌دهنده یک متغیر فاصله‌ای/نسبی محدود باشد (کمتر از ۱۰ طبقه باشد)، استفاده از آماره فراوانی مناسب است.

مثال

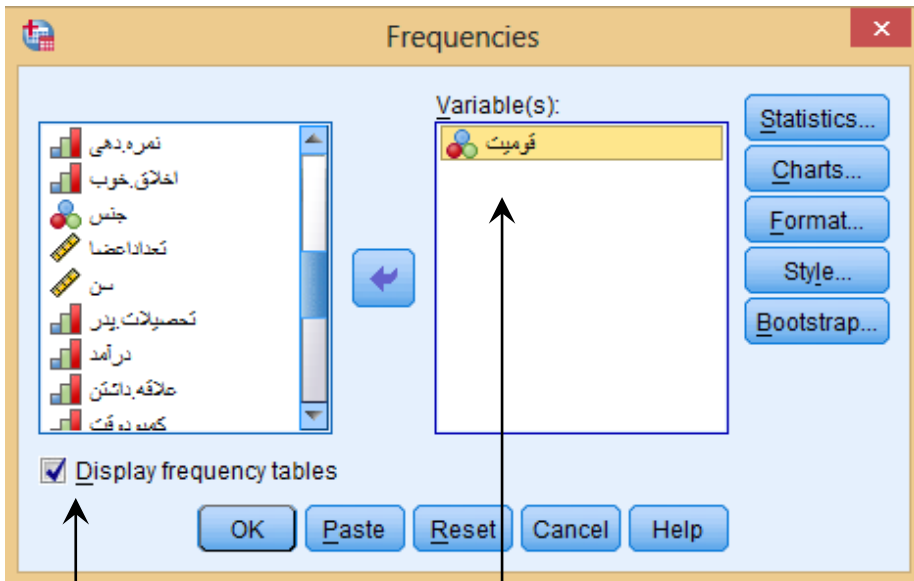
بررسی فراوانی و درصد فراوانی متغیر قومیت پاسخگویان

قومیت در مثال شامل چهار طبقه می‌شود: ترک، فارس، کرد و سایر قوم‌ها. می‌خواهیم بدانیم فراوانی و درصد فراوانی هرکدام از طبقات در نمونه ما چقدر است و به بیان دیگر، چه تعداد فارس هستند، چه تعداد کرد هستند و غیره.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Descriptive Statistics ---> Frequencies



(۲) به‌طور پیش‌فرض گزینه نشان دادن فراوانی‌ها فعال است. گزینه OK را انتخاب می‌کنیم.

(۱) در این کادر متغیر یا متغیرهایی را که می‌خواهیم آمار توصیفی‌شان را بدست بیاوریم وارد می‌کنیم. می‌توانیم یک یا چند متغیر را وارد کادر متغیرها کنیم و فراوانی چندمتغیر را به‌طور همزمان بدست بیاوریم.

نتیجه:

بخش‌های مهم جداول خروجی با رنگ خاکستری متمایز شده‌اند. در جدول اول مشاهده می‌گردد که تعداد داده‌های مورد تایید (Valid) برابر با ۱۰۰ عدد است و هیچ داده ناقص یا بدون پاسخی (Missing) وجود ندارد. به بیان دیگر تمامی ۱۰۰ پاسخگو، به سوال مربوط به قومیت پاسخ داده‌اند.

Statistics

قومیت

N	Valid (داده معتبر)	100
	Missing (داده ناقص)	0

قومیت

	Frequency فراوانی	Percent درصد فراوانی	Valid Percent درصد فراوانی معتبر	Cumulative Percent درصد فراوانی تجمعی
ترک	18	18.0	18.0	18.0
فارس	45	45.0	45.0	63.0
کرد	15	15.0	15.0	78.0
سایر	22	22.0	22.0	100.0
Total	100	100.0	100.0	

در جدول دوم، فراوانی و درصد فراوانی هر کدام از قومیت‌ها نشان داده شده است. در این جدول دو ستون فراوانی (Frequency) و درصد فراوانی تایید شده (Valid Percent) مهم‌تر از دو ستون دیگر هستند. همان‌طور که در ستون فراوانی (Frequency) مشاهده می‌گردد تعداد ۱۸ نفر از پاسخگویان ترک هستند، ۴۵ نفر فارس هستند، ۱۵ نفر کرد هستند و ۲۲ نفر قومیت‌های دیگری دارند. چون تعداد پاسخگویان برابر با ۱۰۰ نفر بوده است در نتیجه مقادیر فراوانی و درصد فراوانی کاملاً با هم یکسان می‌شوند. چنانچه مجموع فراوانی ۱۰۰ نباشد نتایج دو ستون با هم تفاوت خواهد داشت. نتایج ستون درصد فراوانی تایید شده، درصد فراوانی طبقات را بدون در نظر گرفتن داده‌های ناقص نشان می‌دهد که در هنگام گزارش یافته‌ها باید از نتایج این ستون از جدول استفاده کرد (و نه ستون درصد).

۲) چارک

اگر تعداد مشاهدات خیلی زیاد باشد (مثلاً بیشتر از ۲۵ یا ۳۰)، گاهی مفید است که مفهوم میانه را تعمیم دهیم و مجموعه داده‌ها را به چهار قسمت تقسیم کنیم. درست همان‌طور که نقطه تقسیم داده‌ها به دو نیمه میانه خوانده شده، نقاط تقسیم داده‌ها به چهار قسمت را چارک می‌نامند. چارک مختص داده‌های فاصله‌ای/نسبی است.

در آمار توصیفی به هر یک از سه مقداری که یک مجموعه از داده‌های مرتب شده را به چهار بخش مساوی تقسیم می‌کند چارک گفته می‌شود. در این صورت هر کدام از آن

بخش‌ها یک‌چهارم از نمونه یا جمعیت را به نمایش می‌گذارد. مثلاً اگر یازده داده زیر را داشته باشیم و قصد به دست آوردن چارک را داشته باشیم ابتدا داده‌ها را به صورت صعودی و از کم به زیاد مرتب می‌کنیم. سپس اعداد مربوط به هر چارک را پیدا می‌کنیم.

به عنوان مثال داده‌های زیر را در اختیار داریم و می‌خواهیم چارک‌های آن را پیدا کنیم: ۶، ۴۷، ۴۹، ۱۵، ۴۲، ۴۱، ۷، ۳۹، ۴۳، ۴۰، ۳۶. برای به‌دست‌آوردن چارک ابتدا داده‌ها را به صورت صعودی مرتب می‌کنیم: ۶، ۷، ۱۵، ۳۶، ۳۹، ۴۰، ۴۱، ۴۲، ۴۳، ۴۷، ۴۹. اکنون به دنبال عددی می‌گردیم که یک چهارم اعداد دارای آن مقدار یا کمتر هستند (عدد ۱۵)، عددی که نیمی از داده‌ها دارای آن مقدار یا کمتر هستند (۴۰) و عددی که سه‌چهارم داده‌ها دارای آن مقدار یا کمتر هستند (عدد ۴۳). در نتیجه: چارک اول یا Q_1 = ۱۵، چارک دوم یا Q_2 = ۴۰ و چارک سوم یا Q_3 = ۴۳

☑ نکته: چارک مختص داده‌های کمی (فاصله ای/نسبی) است.

قصد داریم چارک‌های مربوط به معدل مقطع کارشناسی دانشجویان را به دست بیاوریم و دانشجویان را بر اساس مقطع کارشناسی به چهار قسمت تقسیم کنیم.

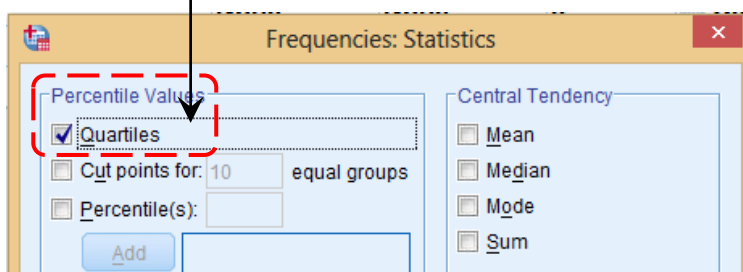
اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Descriptive Statistics ---> Frequencies

مانند دستور فراوانی، متغیر موردنظر را وارد کادر Variables می‌کنیم و سپس گزینه Statistics را انتخاب می‌کنیم.

در کادر مقادیر صدک (Percentile Values) گزینه چارک‌ها (Quartiles) را انتخاب می‌کنیم.



نتایج:

جدول بعد چارک‌های به دست آمده را نشان می‌دهد. در این جدول عدد ۲۵ نشان دهنده چارک اول، عدد ۵۰ نشان دهنده چارک دوم و عدد ۷۵ نشان دهنده چارک سوم است. در نتیجه معدل کارشناسی یک چهارم پایینی دانشجویان ۱۵.۳۳ یا کمتر است. معدل نیمی از دانشجویان مساوی یا کمتر از ۱۶.۷۵ است و معدل سه چهارم دانشجویان برابر با ۱۷.۲۵ یا کمتر است: چارک اول یا $Q_1 = 15.33$ ، چارک دوم یا $Q_2 = 16.75$ و چارک سوم یا $Q_3 = 17.25$

Statistics

معدل کارشناسی

N	Valid (داده یا مورد های معتبر)	100
	Missing (داده های ناقص یا بدون پاسخ)	0
Percentiles (چارک ها)	25	15.33
	50	16.75
	75	17.25

شاخص‌های گرایش به مرکز

یکی دیگر از شیوه‌های تلخیص (خلاصه‌کردن) داده‌ها، استفاده از شاخص‌های گرایش به مرکز است. منظور از شاخص گرایش به مرکز، هر عدد یا معیار عددی است که معرف مرکز مجموعه‌ای از داده‌ها باشد. زمانی که داده‌ها کمی باشد و یا زمانی که قصد مقایسه توزیع یک متغیر را در دو جمعیت داشته باشیم، شاخص‌های گرایش به مرکز در کنار شاخص‌های پراکندگی بسیار مفید خواهند بود. در مثال کتاب، اگر بخواهیم بهره هوشی تمامی پاسخگویان را در یک عدد خلاصه کنیم می‌توانیم از میانگین بهره بگیریم. همچنین چنانچه بخواهیم بدانیم که کدام قومیت در نمونه ما بیشترین فراوانی و تعداد را دارد می‌توانیم از شاخص مد (نما) استفاده کنیم. در کل می‌توان گفت که شاخص‌های گرایش به مرکز معرف تمرکز داده‌ها در توزیع متغیر هستند. میانگین، میانه و نما (مد) از شاخص‌های مهم گرایش به مرکز هستند.

میانگین

میانگین شناخته‌شده‌ترین و وسیع‌ترین مقدار متوسطی است که مورد استفاده قرار می‌گیرد و توصیف‌کننده مرکز توزیع فراوانی می‌باشد. میانگین نمونه را به صورت نمادین با $\bar{X}$ که "ایکس بار" خوانده می‌شود. در بسیاری از مجله‌ها برای نمایش میانگین از حرف M استفاده می‌کنند (مونرو، ۱۳۸۹: ۵۰).

میانگین عموماً بهترین شاخص گرایش به مرکز است. یکی از امتیازاتش نسبت به میانه و نما، ثبات بیشتر آن است. بنابراین، اگر چند نمونه را که به طور تصادفی از یک جامعه گرفته شده باشند، مطالعه کنیم، میانگین‌های آن‌ها نسبت به میانه به هم نزدیک‌تر خواهند بود (گال و همکاران، ۱۳۸۳: ۲۹۹). میانگین مختص داده‌های کمی (فاصله ای/نسبی) است و نمی‌توان آن را برای داده‌های کیفی (اسمی و ترتیبی) به کار برد. میانگین اساساً تلخیص توزیع یک متغیر کمی جهت مقایسه است (نایبی، ۱۳۸۸: ۸۶). میانگین، مجموع مقادیر، تقسیم بر تعداد آن‌هاست. به عنوان مثال میانگین سه عدد ۱، ۲ و ۶ برابر با ۳ است. بدین صورت که مجموع سه عدد ۱، ۲ و ۶ برابر با ۹ است که چنانچه عدد ۹ را بر تعداد اعداد که ۳ است تقسیم کنیم، میانگین به دست می‌آید که عدد ۳ است.

میانۀ

در متغیرهای کمی، مقداری (عددی) که نیمی (۵۰ درصد) از جمعیت دارای آن مقدار یا کمتر از آن مقدارند و نیمی دیگر دارای مقداری بیشتر از آن هستند، میانۀ خوانده می‌شود. میانۀ، مقداری است که جمعیت را به دو قسمت مساوی تقسیم می‌کند. میانۀ با رتبه‌بندی موردها از پایین به بالا و یافتن نفر وسط به دست می‌آید. هر طبقه‌ای که در برگزیده نفر وسط است طبقه میانۀ است. در جایی که تعداد افراد زوج باشد، نفر وسط واقعی وجود ندارد و میانۀ در نقطه‌ای بین موردهایی است که در دو سمت نفر وسط خیالی قرار دارند. مثلاً چنانچه بعد خانوار پنج خانواده به این صورت باشد: ۳، ۴، ۵، ۶ و ۳؛ در این صورت عدد ۴ نشان‌دهنده میانۀ بعد خانوار در بین این پنج خانواده است. چنانچه بعد خانوار چهار خانواده به صورت زیر باشد: ۲، ۳، ۴، ۶. در این صورت میانۀ برابر است با مجموع دو عدد ۳ و ۴ تقسیم بر دو. یعنی میانۀ بین دو مورد ۳ و ۴ قرار دارد و میانۀ عدد ۳.۵ است.

برای میانۀ فرمول جبری وجود ندارد و فقط دارای یک روش محاسبه است:

- ۱- مقادیر را به ترتیب منظم می‌کنیم.
- ۲- اگر تعداد کل اعداد فرد باشد، میانۀ عددی است که در وسط قرار دارد و نصف اعداد قبل از آن و نصف دیگر اعداد بعد از آن قرار دارند. اگر در وسط چند عدد مشابه قرار گرفته باشند، میانۀ برابر همان عدد است.
- ۳- اگر تعداد اعداد زوج بود، میانگین اعداد وسط را محاسبه کنید.

مد

رایج‌ترین پاسخ یا طبقه‌ای که بیشترین پاسخ‌ها را به خود اختصاص داده است مد (نما) خوانده می‌شود. این شاخص را می‌توان برای تمامی متغیرهای اسمی، ترتیبی، فاصله‌ای و نسبی محاسبه کرد. به عنوان مثال اگر از بین ۲۰ دانشجوی یک کلاس، ۴ نفر دارای قومیت فارس، ۶ نفر قومیت کرد و ۱۰ نفر دارای قومیت ترک باشند؛ در این صورت قومیت ترک طبقه نما (مد) به حساب می‌آید، یا در بین اعداد ۸، ۸، ۷ و ۵، عدد ۸ نما به حساب می‌آید چرا که بیشتر از سایر اعداد تکرار شده است.

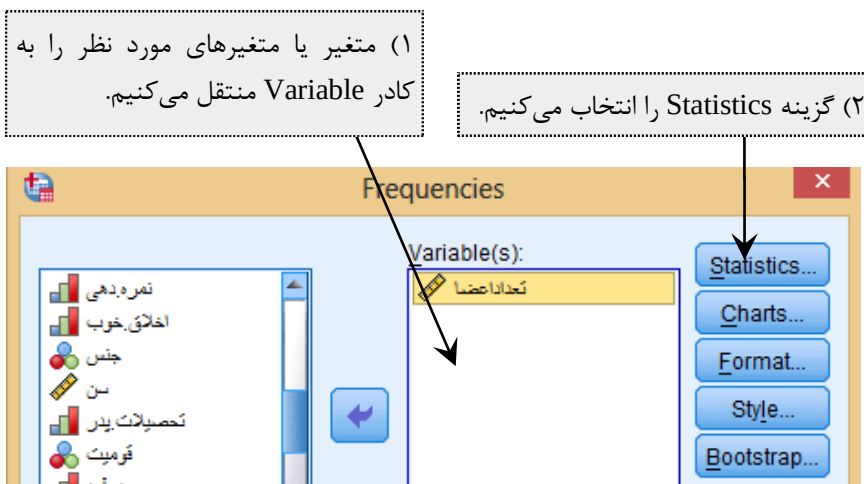
مثال

قصد داریم میانگین، میانه و مد متغیر بعد خانوار (تعداد اعضای خانواده) را در بین نمونه ۱۰۰ نفری از شرکت‌کنندگان به دست بیاوریم.

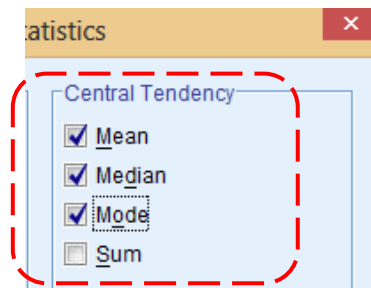
نحوه اجرا:

مانند مثال فراوانی مسیر زیر را دنبال می‌کنیم:

Analyze ---> Descriptive Statistics ---> Frequencies



در کادر Central Tendency یا گرایش‌های مرکزی، آماره‌های میانگین (Mean)، میانه (Median) و مد (Mode) را انتخاب می‌کنیم. در انتها بر روی گزینه Continue و سپس OK کلیک می‌کنیم.



نتایج:

خروجی به دست آمده توسط برنامه SPSS در ادامه آمده است. مقادیر میانگین، میانه و مد در جدول متمایز شده است. میانگین تعداد اعضای خانواده در بین ۱۰۰ نمونه برابر با ۵.۸ نفر است. میانه برابر با ۵ است و مد طبقه یا عدد ۴ است که نشان می‌دهد خانواده‌هایی که دارای ۴ عضو هستند بیشترین تعداد را در نمونه ما دارند. برای تعیین تعداد و درصد فراوانی طبقات می‌توانیم از جدول فراوانی (مثال قبل) بهره بگیریم.

Statistics

تعداد اعضا

N	Valid	100
	Missing	0
Mean (میانگین)		5.8
Median (میانه)		5
Mode (مد یا نما)		4

شاخص‌های پراکندگی

شاخص‌های پراکندگی، میزان پراکندگی (تغییرات) مقادیر هر متغیر را در اطراف میانگین نشان می‌دهند. زمانی که در شاخص‌های گرایش به مرکز، داده‌ها را در یک اندازه واحد خلاصه کنیم، طبیعتاً بخشی از جزئیات و اطلاعات حذف خواهد شد. از این رو باید به دنبال شاخص‌هایی برای اندازه‌گیری تفاوت موردها در یک متغیر باشیم تا از نحوه و میزان پراکندگی و تغییر داده‌ها در اطراف میانگین (یا میانه و مد) مطلع شویم.

به عنوان مثال نمرات درس‌های دو دانشجو را با یکدیگر مقایسه می‌کنیم. معدل (میانگین نمرات) هر دو دانشجو یکسان و برابر با ۱۵ است. نمرات دانشجوی الف بدین صورت است: ۱۷، ۱۶، ۱۵، ۱۴ و ۱۳. و نمرات دانشجوی ب بدین صورت است: ۲۰، ۱۸، ۱۶، ۱۱ و ۱۰ است. نگاهی به نمرات دو دانشجو نشان می‌دهد که دامنه نمرات دانشجوی الف محدودتر است و داده‌ها از نمره میانگین (۱۵) فاصله کمتری دارند. نتیجه می‌گیریم در حالتی که میانگین، یک شاخص گرایش به مرکز است در هر دو دانشجو برابر با ۱۵ است اما پراکندگی نمرات در دانشجوی ب بیشتر از دانشجوی الف است.

شاخص‌های گرایش به مرکز اطلاعاتی از نحوه پراکندگی داده‌ها و نحوه توزیع‌شان به ما نمی‌دهند و جهت اطلاع از نحوه پراکندگی داده‌ها، باید از شاخص‌های پراکندگی استفاده کنیم. از مهم‌ترین شاخص‌های پراکندگی می‌توان به انحراف استاندارد، واریانس، ضریب تغییرات و دامنه تغییرات اشاره کرد.

انحراف استاندارد

شاخص انحراف استاندارد مهم‌ترین شاخص پراکندگی است. انحراف استاندارد (یا انحراف معیار)، یک شاخص پراکندگی است که غالباً در پژوهش‌ها گزارش می‌شود. اساساً، انحراف استاندارد، اندازه‌ای از میزان انحرافی است که نمره‌های یک توزیع از میانگین خود دارند. منطق انحراف استاندارد، نشان دادن متوسط میزان فاصله هر مورد از میانگین است. تفسیر مقدار انحراف استاندارد بدین صورت است که هر چه مقدار انحراف استاندارد بیشتر باشد، پراکندگی نمرات از میانگین هم بیشتر است. انحراف استاندارد را با علامت S یا Sd نشان می‌دهند.

واریانس

با به توان دو رساندن انحراف استاندارد، واریانس به دست می‌آید (به بیان صحیح‌تر، جذر واریانس برابر با انحراف استاندارد است). اما چون تفسیر واریانس کمی دشوار است در گزارش‌ها و پژوهش‌ها از انحراف استاندارد به جای واریانس استفاده می‌شود. انحراف استاندارد بزرگ‌نمایی حاصل از مجذور کردن (به توان دو رساندن) را تا حد زیادی خنثی می‌سازد و تفسیر میزان پراکندگی را آسان‌تر و قابل فهم‌تر می‌سازد.

ضریب تغییرات

ضریب تغییرات (پراکندگی) شاخصی است که برای اندازه‌گیری توزیع پراکندگی داده‌های آماری به کار می‌رود. از ضریب تغییرات برای مقایسه پراکندگی دو یا چند صفت (متغیر) استفاده می‌کنند و کاربرد اصلی آن مقایسه متغیرهایی است که واحدهای سنجش متفاوتی دارند. ضریب تغییرات میزان پراکندگی را به ازای یک واحد از میانگین بیان می‌کند. این شاخص تنها برای سطح سنجش نسبی کاربرد دارد. ضریب تغییرات را با علامت CV نشان می‌دهند.

معمولاً ضریب تغییرات را در عدد ۱۰۰ ضرب می‌کنند تا عدد نهایی برحسب درصد به دست بیاید. ضریب تغییرات از تقسیم انحراف استاندارد بر میانگین به دست می‌آید و فرمول آن بدین صورت است:

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

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

$$\begin{aligned} 10 &= 100 \times .10 \rightarrow .10 = \frac{17}{170} = \text{ضریب تغییرات قد} \\ 10 &= 100 \times .10 \rightarrow .10 = \frac{5}{50} = \text{ضریب تغییرات سن افراد} \end{aligned}$$

نتایج نشان می‌دهد که ضریب تغییرات قد برابر با ۱۰ و سن هم برابر با ۱۰ به دست آمده است. در نتیجه میزان پراکندگی متغیر قد و سن افراد یکسان است. مقایسه مقادیر انحراف استاندارد دو متغیر قد و سن نشان می‌دهد که انحراف استاندارد قد برابر با ۱۷ و سن برابر با ۵ است که اگر واحدهای سنجش (مقیاس) دو متغیر را نادیده بگیریم نشان می‌دهد که میزان پراکندگی قد بسیار بیشتر از سن است. اما شاخص ضریب تغییرات با خنثی کردن نقش واحدهای سنجش در مقایسه میزان پراکندگی متغیرها نشان داد که میزان پراکندگی و تغییرات در دو متغیر سن و قد برابر است.

دامنه تغییرات

اختلاف بین بزرگ‌ترین داده از کوچک‌ترین داده در یک توزیع مشخص، دامنه تغییرات آن توزیع نام دارد. دامنه را با علامت R نشان می‌دهند. برای به دست آوردن دامنه تغییرات، بزرگ‌ترین داده (عدد) را منهای کوچک‌ترین داده می‌کنیم. مثلاً دامنه تغییرات اعداد ۱، ۲ و ۵ عدد ۴ است که اگر عدد ۵ (بزرگ‌ترین داده) را از عدد ۱ (کوچک‌ترین داده) کم کنیم، عدد ۴ که معادل دامنه تغییرات است به دست می‌آید.

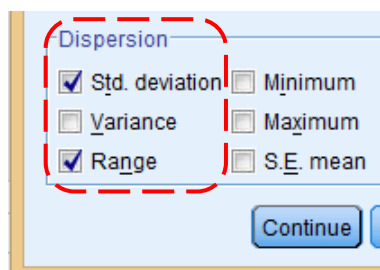
مثال

مثال قبل را که مربوط به محاسبه شاخص‌های گرایش به مرکز متغیر بعد خانوار است را ادامه می‌دهیم و انحراف استاندارد و دامنه تغییرات متغیر بعد خانوار را محاسبه کنیم.

اجرا:

تمامی مراحل مانند شاخص‌های گرایش به مرکز است و فقط در پنجره Statistics در کادر پراکندگی یا Dispersion گزینه انحراف استاندارد (Std.deviation) و دامنه تغییرات (Range) را انتخاب می‌کنیم.

گزینه‌های انحراف استاندارد و دامنه تغییرات را انتخاب می‌کنیم.
در انتها گزینه Continue و OK را انتخاب می‌کنیم.



نتایج:

خروجی به دست آمده از برنامه در ادامه آمده است. مشاهده می‌شود که مقدار انحراف استاندارد متغیر بعد خانوار ۱.۸۹ به دست آمده است که بدین معناست که بعد خانوار هر پاسخگو به طور متوسط، ۱.۸۹ از میانگین فاصله دارد. البته معمولاً کم یا زیاد بودن مقدار انحراف استاندارد در مقایسه بین گروه‌ها یا جمعیت‌های مختلف است که معنی پیدا می‌کند و به تنهایی ارزیابی نمی‌شود. دامنه تغییرات برابر با عدد ۹ نفر است که نشان می‌دهد اختلاف بین بیشترین تعداد اعضای خانواده از کمترین اعضای خانواده برابر با ۹ است.

Statistics

تعداد اعضا

N	Valid	100
	Missing	0
Std. Deviation (انحراف استاندارد)		1.88
Range (دامنه تغییرات)		9



تعریف

به جدول صفحه ۹۶ رجوع کنید:

- ۱- فراوانی و درصد فراوانی متغیرهای جنس و میزان تحصیلات را به دست بیاورید.
- ۲- چارک متغیر وزن را به دست بیاورید و تفسیر کنید. افراد با چه وزنی در هر کدام از چارک‌ها قرار می‌گیرند؟
- ۳- میانگین، میانه و مد را برای دو متغیر وزن افراد و میزان ورزش به دست بیاورید.
- ۴- انحراف استاندارد و دامنه تغییرات متغیر وزن و میزان ورزش را به دست بیاورید. هر فرد به طور متوسط چند کیلو از میانگین وزن نمونه فاصله دارد؟
- ۵- ضریب تغییرات دو متغیر وزن و میزان ورزش را به دست آورده و پاسخ دهید که پراکندگی متغیر وزن بیشتر است یا متغیر میزان ورزش؟

نمودارها

شیوه دیگر توصیف متغیرها، نشان دادن توزیع متغیر با نمودار است. نمودارها شیوه‌ای جذاب برای ارائه نتایج و اطلاعات هستند و امکان به دست آوردن اطلاعات را در فرصتی کوتاه فراهم می‌آورند. نمودارها، به ویژه برای ارائه گزارش در رسانه‌های عمومی (روزنامه‌ها، تلویزیون) بسیار مفید هستند، چون فهم آن برای افراد غیرمتخصص، از جدول توزیع فراوانی ساده‌تر است. نمودار توزیع متغیر مانند جدول توزیع متغیر است با این تفاوت که در جدول توزیع متغیر، تعداد یا فراوانی با عدد ارائه می‌شود اما در نمودار با خط یا سطح نشان داده می‌شود. نمودارها را می‌توان به دو دسته نمودارهای تک‌متغیره و نمودارهای چند متغیره تقسیم کرد. نمودارهای دایره‌ای، ستونی و هیستوگرام از رایج‌ترین نمودارهای تک متغیره هستند. رایج‌ترین نمودارهای متغیرهای کیفی نمودار دایره‌ای و نمودار ستونی است و نمودار هیستوگرام از رایج‌ترین نمودارها برای متغیر کمی است.

(۱) نمودار دایره‌ای

نمودار دایره‌ای جهت بررسی درصد فراوانی یک متغیر اسمی یا ترتیبی به کار می‌رود. مساحت دایره به ۱۰۰ درجه (درصد) تقسیم شده است و متغیرها بر حسب درصد فراوانی‌شان بخشی از مساحت دایره را پرخواهند کرد. در مثال کتاب، می‌توان برای متغیرهای جنس و قومیت نمودار دایره‌ای رسم کرد. البته برای سوالات میزان تحصیلات و درآمد هم می‌توان نمودار دایره‌ای رسم کرد اما چون این متغیرها ترتیبی هستند نمودار ستونی برای آن‌ها مناسب‌تر است.

مثال

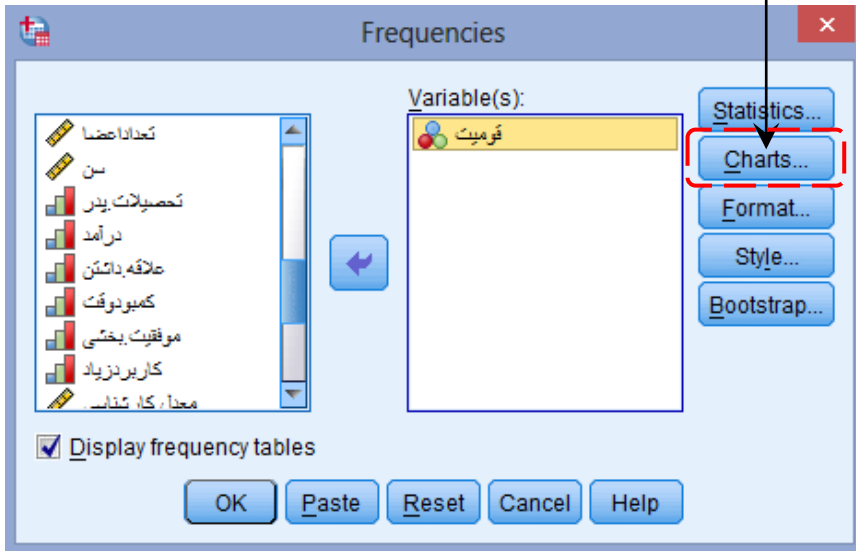
قومیت متغیری اسمی است و در مثال کتاب شامل چهار طبقه می‌شود: فارس، ترک، کرد و سایر قومیت‌ها. در اینجا قصد داریم نمودار دایره‌ای متغیر قومیت را برحسب درصد فراوانی به دست بیاوریم.

اجرا:

دستور فراوانی را اجرا می‌کنیم:

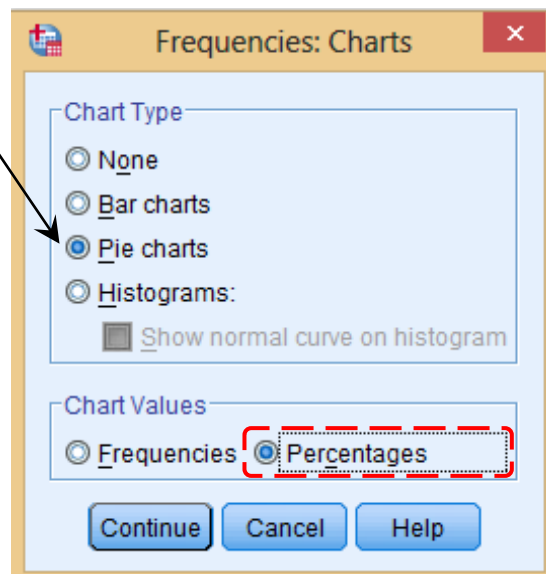
Analyze ---> Descriptive Statistics ---> Frequencies

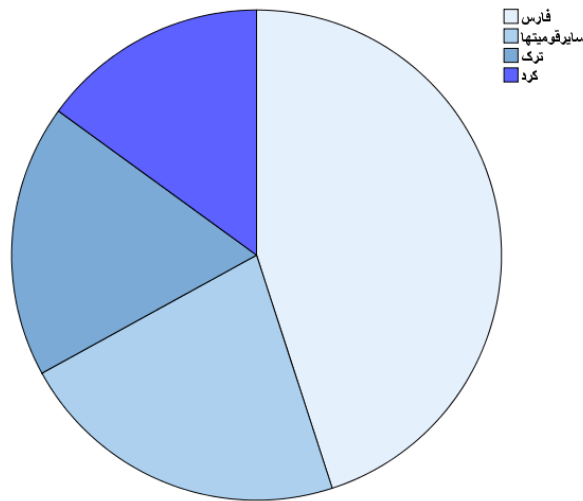
متغیر قومیت را وارد کادر Variable کرده و گزینه Charts (نمودارها) را انتخاب می‌کنیم.



گزینه Pie Charts (نمودار دایره‌ای) را انتخاب می‌کنیم.

گزینه Percentages (درصد)، نمودار را برحسب درصد نشان می‌دهد و گزینه Frequencies (فراوانی) نمودار را بر حسب فراوانی طبقات. بهتر است نمودارها را بر حسب درصد رسم نماییم (البته شکل نمودار دایره‌ای و سطح اختصاص یافته به هر طبقه از متغیر، در هر دو حالت یکسان است).





شکل ۳-۱- نمودار دایره‌ای قومیت دانشجویان (درصد)

نمودار دایره‌ای متغیر قومیت برحسب فراوانی قومیت‌ها به صورت نزولی مرتب شده است و نشان می‌دهد دانشجویان با قومیت فارس بیشترین درصد فراوانی را دارند و کمترین درصد فراوانی متعلق به دانشجویان با قومیت کرد است. مطابق راهنمای نمودار، ترتیب طبقات برحسب رنگ از روشن به تیره مرتب شده است.

☑ نکته: برای ویرایش و انجام تغییرات در نمودارها کافیست در پنجره خروجی (SPSS Viewer) بر روی نمودار دوبار کلیک کنیم. با این عمل، صفحه ویرایش نمودار (Chart Editor) باز می‌شود. این پنجره امکانات و دستورات متنوعی دارد که با استفاده از ابزارها و دستورات آن می‌توان تغییرات مورد نظر را بر روی نمودار انجام داد.

۲) نمودار ستونی (میله‌ای)

نمودار ستونی شامل مجموعه‌ای از ستون‌هاست که با فاصله یکنواختی در کنار هم قرار می‌گیرند و هر ستون مختص یک طبقه از متغیر و طول آن متناسب با فراوانی یا درصد آن طبقه است.

نمودار ستونی برای متغیرهای کیفی (اسمی و ترتیبی) به کار می‌رود، اما برای متغیر کمی که تعداد طبقات آن کم باشد هم قابل استفاده است. به عنوان مثال چنانچه

متغیری به نام تعداد فرزند داشته باشیم و در نمونه نهایی تعداد فرزندان هر خانواده از ۱ تا ۵ فرزند باشد و در واقع تعداد فرزندان شامل ۵ طبقه باشد بازهم می‌توانیم از این نمودار استفاده کنیم. در مثال کتاب، می‌توان برای متغیرهای تحصیلات پدر، قومیت و درآمد اقدام به ترسیم نمودار ستونی نمود. همچنین برای متغیرهایی که در قالب طیف لیکرت سنجیده شده باشند می‌توان از نمودار ستونی استفاده کرد.

مثال

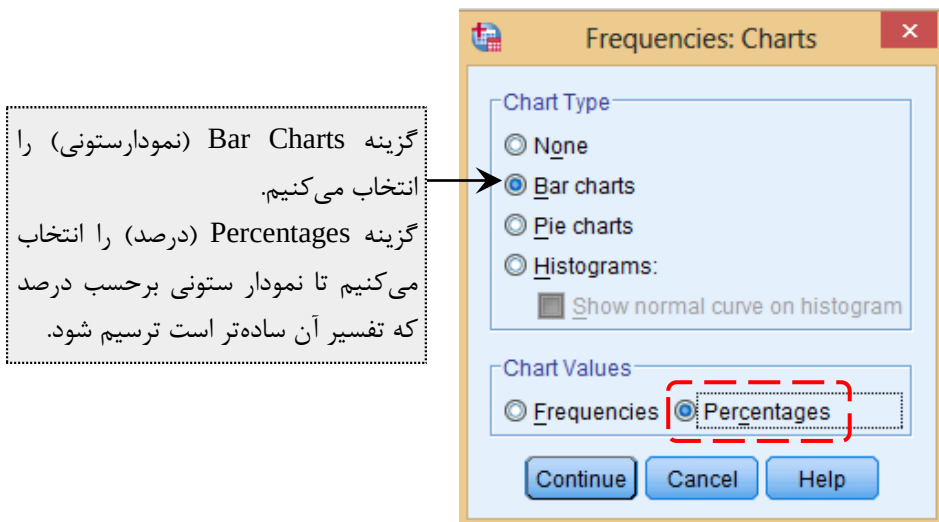
قومیت متغیری اسمی است و شامل چهار طبقه می‌شود: ترک، فارس، کرد و سایر قومیت‌ها. در اینجا قصد داریم نمودار ستونی (میله‌ای) متغیر قومیت را به دست بیاوریم.

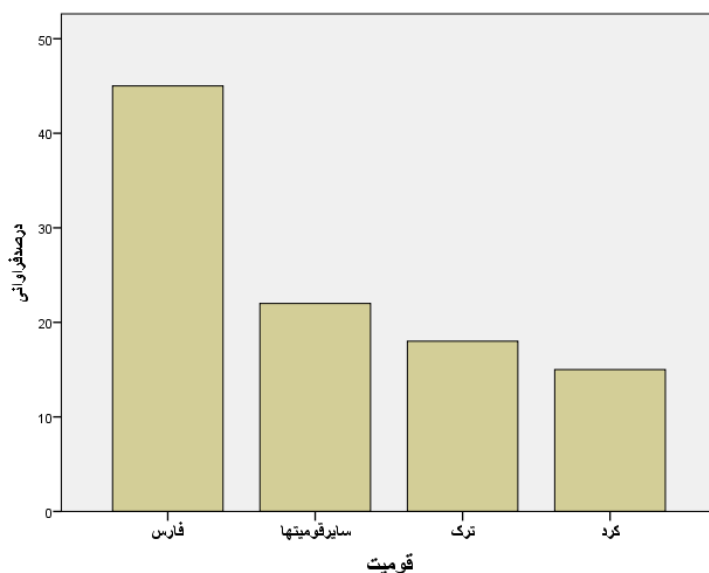
اجرا:

مانند مثال قبل دستور فراوانی را اجرا می‌کنیم:

Analyze ---> Descriptive Statistics ---> Frequencies

متغیر قومیت را وارد کادر Variable کرده و گزینه Charts (نمودارها) را انتخاب می‌کنیم.





شکل ۳-۲- نمودار ستونی قومیت دانشجویان (درصد)

محور افقی طبقات متغیر قومیت را نشان می‌دهد که شامل قومیت‌های فارس، ترک، کرد و سایر قومیت‌ها می‌شود. محور عمودی درصد فراوانی هر کدام از طبقات یا قومیت‌ها را نشان می‌دهد. همان طور که مشاهده می‌گردد دانشجویان فارس، دارای بیشتری درصد فراوانی بوده و تعدادشان از سایر قومیت‌ها بیشتر است.

۳) نمودار هیستوگرام

نمودار هیستوگرام برای متغیرهای کمی (فاصله‌ای/نسبی) کاربرد دارد. این نمودار، فراوانی یا درصد فراوانی هر کدام از طبقات را به صورت ستون نشان می‌دهد. همچنین با استفاده از نمودار هیستوگرام و رسم منحنی توزیع نرمال، می‌توانیم از شکل توزیع (نرمال بودن، وضعیت چولگی و کشیدگی) متغیر مدنظر اطلاع پیدا کنیم. در مثال کتاب، می‌توانیم برای متغیرهای سن، بهره هوشی، معدل و ... نمودار هیستوگرام رسم کرد.

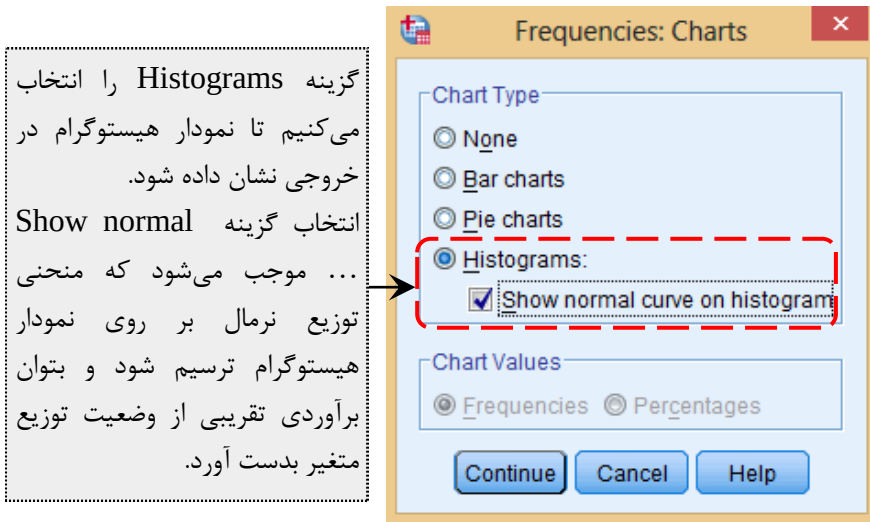
مثال

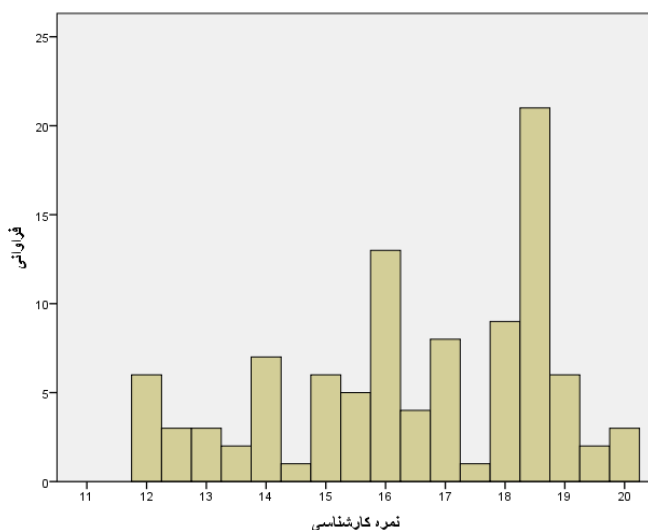
می‌خواهیم نمودار هیستوگرام متغیر نمره درس SPSS در مقطع کارشناسی دانشجویان را به دست بیاوریم.

اجرا:

Analyze ---> Descriptive Statistics ---> Frequencies

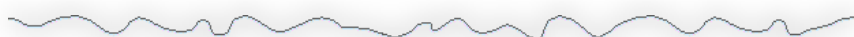
متغیر نمره کارشناسی را وارد کادر Variable کرده و گزینه Charts (نمودارها) را انتخاب می‌کنیم.





شکل ۳-۳- نمودار هیستوگرام متغیر نمره درس SPSS دانشجویان

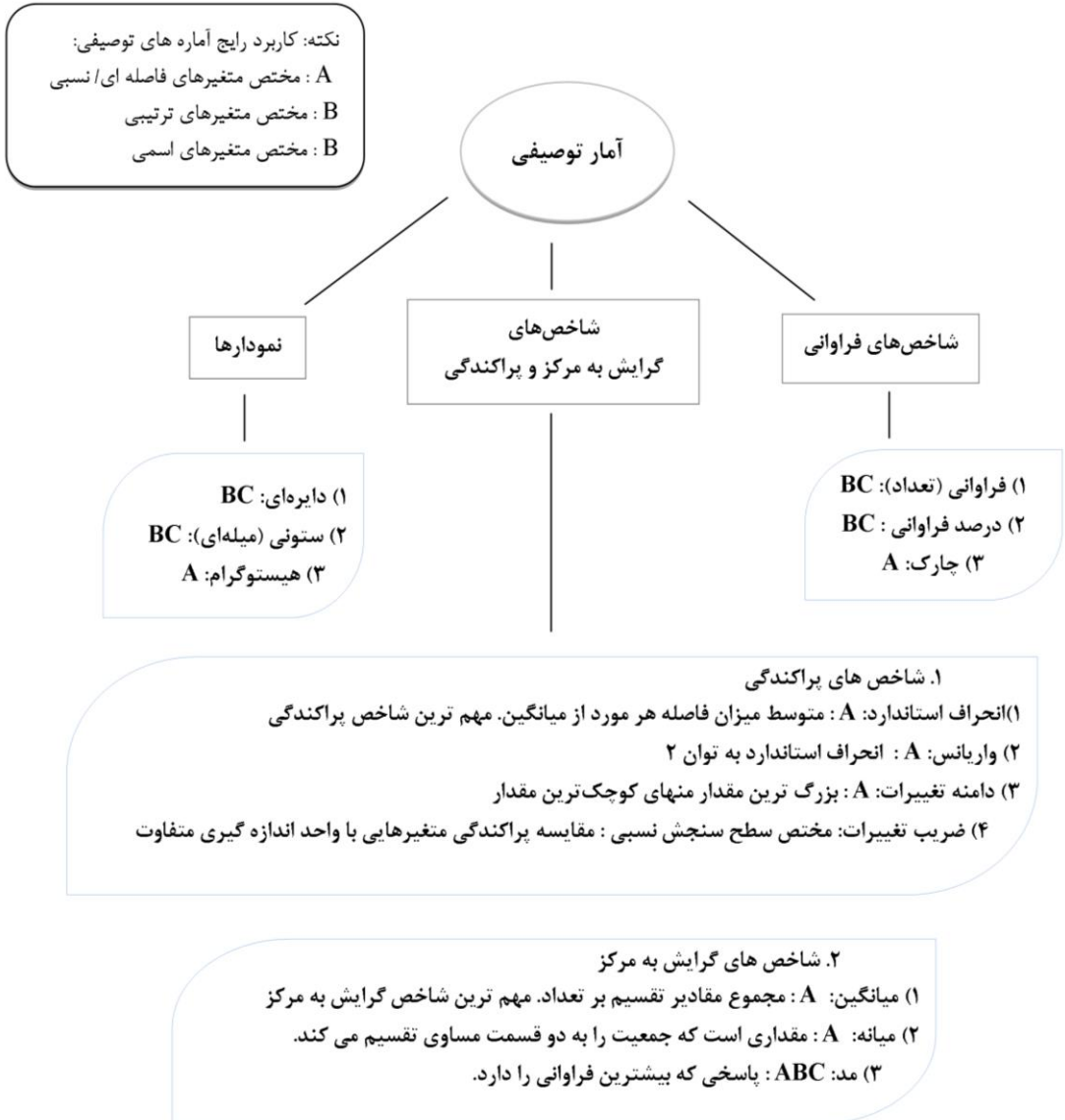
نمودار هیستوگرام متغیر نمره SPSS در مقطع کارشناسی رسم شده است. در محور افقی، نمره‌های پاسخگویان درج شده است و در محور عمودی درصد فراوانی هر کدام از نمرات آمده است. همان طور که مشاهده می‌شود درصد فراوانی افرادی که نمره ۱۹ گرفته‌اند بیشتر از سایرین است و ستون مربوط به نمره ۱۹ بیشترین ارتفاع را دارد.



تعریف

به جدول صفحه ۹۶ رجوع کنید:

- ۱- نمودار دایره‌ای را برای متغیرهای جنس و میزان تحصیلات ترسیم کنید. یکبار نمودار را برحسب فراوانی و بار دیگر برحسب درصد ترسیم کنید.
- ۲- نمودار ستونی متغیر تحصیلات را برحسب درصد به دست بیاورید. بیشترین فراوانی در بین طبقات متغیر تحصیلات مربوط به کدام طبقه و با چه تعداد فراوانی و درصد فراوانی است.
- ۳- نمودار هیستوگرام را برای متغیرهای وزن و میزان ورزش در هفته به دست بیاورید. منحنی توزیع نرمال را برای متغیرهای مدنظر به دست بیاورید.



بخش دوم:

پیش فرض های آماری

پیش فرض های آماری، پایه بسیاری از آزمون های آماری تک متغیری و چندمتغیری است. شرایط مهم و اساسی برای تحلیل داده های چندمتغیری، برقراری پیش فرض های نرمال بودن، خطی بودن و یکسانی پراکندگی داده ها است. چنانچه یک یا چندتا از این مفروضه ها نادیده گرفته شود، در این صورت، در نتایج آماری سوگیری یا تحریف رخ می دهد (میزر و همکاران، ۱۳۹۱: ۱۰۷). قبل از انجام تحلیل های آماری تک متغیره و چندمتغیره (که پیش فرض های نرمال بودن، خطی بودن، یکسانی پراکندگی و نبود هم خطی چندگانه بین متغیرهای مستقل در مورد آن ها باید صدق کند) باید برقراری پیش فرض های آماری را بیازماییم.

اگر انحراف از پیش فرض های آماری ناچیز باشد می توان با کمی تسامح و تساهل این انحراف را نادیده گرفت و به ادامه تحلیل پرداخت. اگر انحراف از پیش فرض ها قابل توجه باشد باید یا از روش تبدیل داده ها برای برقرار کردن مجدد پیش فرض ها استفاده کنیم یا از آزمون های جایگزین استفاده کنیم که پیش فرض های فوق را مطرح نمی کنند (آزمون های ناپارامتریک).

در ادامه به توضیح چهار پیش فرض آماری نرمال بودن، خطی بودن رابطه ها، یکسانی پراکندگی و نبود هم خطی چندگانه می پردازیم و روش آزمون هر کدام را توضیح می دهیم.

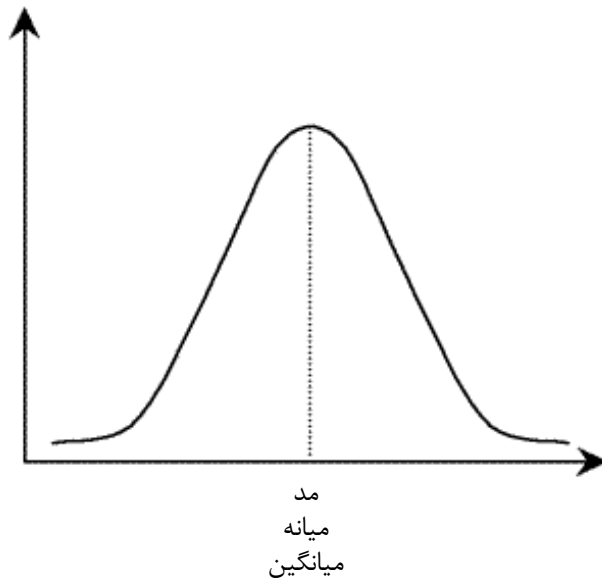
آزمون نرمال بودن داده‌ها

مفهوم توزیع نرمال در مورد داده‌های پارامتری صدق می‌کند (نه داده‌های ناپارامتری). آزمون نرمال بودن، با ایجاد یک نمودار احتمال نرمال بودن (به شکل زنگوله و نیز متقارن نسبت به میانگین)، به آزمون این فرض می‌پردازد که آیا مشاهدات پژوهش از توزیع نرمال تبعیت می‌کنند یا خیر. بسیاری از مشخصه‌های انسانی مانند هوش، نگرش-ها و شخصیت در جمعیت (جامعه) دارای توزیع نسبتاً نرمال هستند. البته توزیع نرمال به معنای توزیع استاندارد یا مطلوب نیست.

نرمال بودن، اساسی‌ترین پیش‌فرض تحلیل چندمتغیره است. اگر این فرض برقرار نباشد، برخی آزمون‌های آماری مشخص، غیرمعتبر بوده و قابل استفاده نیستند (هیر و دیگران، ۲۰۱۰). اهمیت آشنایی و سنجش نرمال بودن توزیع داده‌ها در این است که برخی از روش‌های آماری مانند همبستگی پیرسون، آزمون‌های t و آزمون تحلیل واریانس بر فرض نرمال بودن توزیع داده‌ها (در جامعه) استوارند. همچنین برآورد پارامتر جمعیت نیز با اتکاء به نرمال بودن توزیع متغیر در جمعیت صورت می‌گیرد.

توزیع نرمال دارای ویژگی‌های زیر است:

- ۱- متقارن بوده، حداکثر ارتفاع در میانگین قرار دارد. نیمی از نمره‌ها در بالای میانگین و نیمی دیگر در پایین میانگین قرار دارد.
- ۲- مقادیر نما، میانه و میانگین برابر است.
- ۳- منحنی توزیع شبیه زنگوله است (ساعی، ۱۳۸۸: ۷۹)



شکل ۳-۴- منحنی توزیع نرمال

☑ نکته: چنانچه توزیع یک متغیر در جمعیت نرمال باشد و ما از این جمعیت نمونه‌ای با حجم ۳۰ یا بیشتر ($n \geq 30$) به طور تصادفی ساده انتخاب کنیم، در این صورت توزیع متغیر در نمونه ما نیز نرمال خواهد بود.

برای تشخیص وضعیت توزیع داده‌ها (وضعیت نرمال بودن یا بررسی چولگی و کشیدگی) روش‌های متعددی وجود دارد. در جدول ۲-۳ برخی روش‌های تشخیص وضعیت نرمال بودن توزیع داده‌ها نشان داده شده است.

جدول ۳-۲: روش‌های عددی و تصویری ارزیابی وضعیت توزیع داده‌ها (آزمون‌های سنجش نرمال بودن)

انواع روش‌ها		ماهیت روش‌ها
روش‌های عددی	روش‌های گرافیکی	
کشیدگی کجی	نمودار ساقه و برگ نمودار جعبه‌ای نمودار هیستوگرام	توصیفی
آزمون کولموگروف-اسمیرنف آزمون شاپیرو-ویلک	نمودار احتمال - احتمال (P-P) نمودار چارک - چارک (Q-Q)	مبتنی بر نظریه

(اقتباس از حبیب پور و صفری، ۱۳۸۸: ۴۰۸)

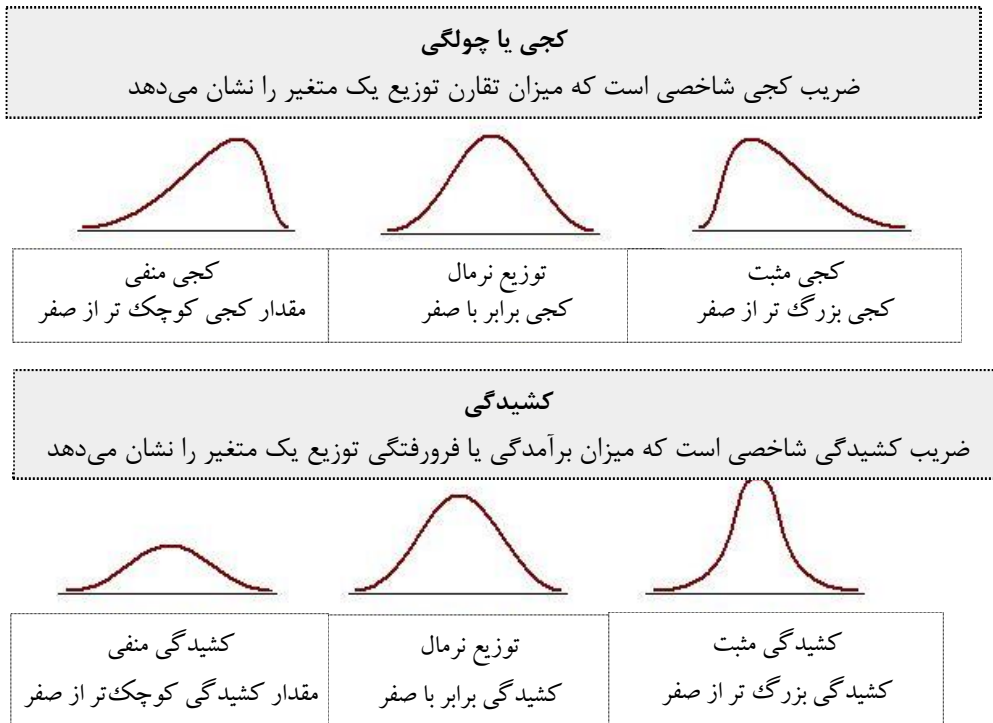
در این کتاب از بین روش‌های ذکر شده در جدول ۳-۲، از روش‌های عددی کجی و کشیدگی و آزمون‌های کولموگروف-اسمیرنف و شاپیرو-ویلک بهره می‌گیریم.

۱) کجی و کشیدگی

کجی یا چولگی در آمار، مقداریست از تقارن توزیع یک متغیر در اطراف میانگین. در واقع کجی، انحراف منحنی نرمال از حالت تقارن است. مقدار کجی می‌تواند مثبت یا منفی باشد. در حالت کجی مثبت، میانگین بزرگتر از میانه و میانه بزرگتر از مد است و در حالت کجی منفی، مد بزرگتر از میانه و میانه بزرگتر از میانگین است.

کشیدگی بیان‌کننده نحوه انباشته‌شدن نمره‌ها در مرکز توزیع یک متغیر است. در واقع کشیدگی، برآمدگی یا فرورفتگی منحنی توزیع نرمال است. در حالت کشیدگی تقارن توزیع حفظ می‌شود و دو نیمه منحنی متقارن هستند، اما نقطه اوج منحنی نرمال دچار تغییر می‌شود. مقدار کشیدگی می‌تواند مثبت یا منفی باشد. در حالت کشیدگی منفی، منحنی توزیع متغیر فررفته‌تر از حالت نرمال می‌شود و نقطه اوج توزیع متغیر، پایین‌تر از توزیع نرمال است. در حالت کشیدگی مثبت، منحنی توزیع متغیر برآمده‌تر از منحنی نرمال می‌شود و نقطه اوج توزیع متغیر، بالاتر از توزیع نرمال است.

در شکل شماره ۳-۵ منحنی توزیع نرمال و انواع حالت‌های کشیدگی و چولگی نمایش داده شده است.



شکل ۳-۵- مقایسه منحنی توزیع نرمال با منحنی‌های دارای کجی و کشیدگی

در مورد کجی و کشیدگی، برخی آمارشناسان تفسیر ± 1 را در مورد کجی، کشیدگی یا هر دو ترجیح می‌دهند (میزر و دیگران، ۱۳۹۱: ۸۵). به این معنا که چنانچه مقدار کجی (تقارن توزیع) و کشیدگی در دامنه -1 تا $+1$ نباشد (یعنی یا بزرگتر از $+1$ و یا کوچکتر از -1 باشد)، به عنوان شاخص مهمی برای انحراف از نرمال بودن تلقی می‌شود.

۲) آزمون‌های آماری نرمال بودن

آزمون‌های آماری جهت سنجش نرمال بودن شامل آزمون کولموگروف-اسمیرنوف^۷ و آزمون شاپیرو-ویلک^۸ است. اگر چه این هر دو آزمون می‌توانند به طور مؤثری به کار

^۷ Kolmogorov-Smirnov

^۸ Shapiro-Wilk

گرفته شوند، استیونس (۲۰۰۲) استفاده از آزمون شاپیرو-ویلک را پیشنهاد می‌کند، زیرا به نظر می‌رسد «در تشخیص انحراف از نرمال قوی‌تر است». معنی‌داری آماری این شاخص‌ها به طور آرمانی در سطح آلفای ($P < .001$)، بیانگر تخطی از نرمال بودن تک متغیری است (میزر، گامست و گارینو، ۱۳۹۱:۱۰۷). یعنی اگر معنی‌داری این آزمون‌ها کمتر از مقدار $.001$ به دست بیاید نشان می‌دهد که توزیع متغیر نرمال نیست و اگر سطح معنی‌داری بیشتر از $.001$ شود، نشان می‌دهد توزیع متغیر نرمال یا نزدیک به نرمال است.

❑ نکته: آزمون‌های نرمال بودن مختص متغیرهای کمی (فاصله‌ای/نسبی) است.

مثال

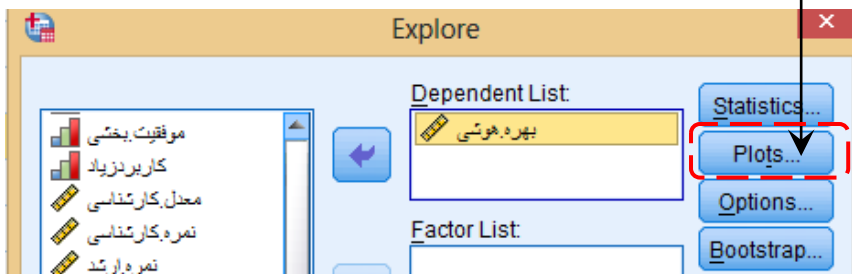
می‌خواهیم نرمال بودن توزیع متغیر بهره هوشی دانشجویان را بررسی کنیم.

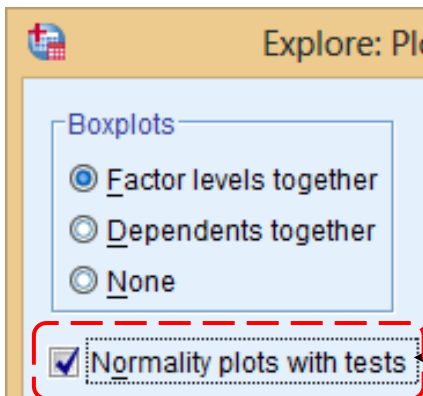
اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Descriptive Statistics ---> Explore

در کادر Dependent List متغیر بهره هوشی که می‌خواهیم نرمال بودن آن را بیازمایی وارد می‌کنیم. بر روی گزینه Plots (نمودارها) کلیک می‌کنیم.





گزینه آزمون‌ها و نمودارهای نرمال بودن (Normality Plots....) را فعال می‌کنیم. با انتخاب این گزینه هم نمودارها و هم آزمون‌های نرمال بودن را بدست می‌آوریم.

نتایج:

در جدول بعد نتیجه آزمون نرمال بودن داده‌ها گزارش شده است. همان‌طور که مشاهده می‌شود برنامه دو آزمون نرمال بودن را ارائه می‌دهد: آزمون شاپیرو-ویلک (Shapiro-Wilk) و آزمون کولموگروف-اسمیرنوف (Kolmogorov-Smirnov).

جهت بررسی نرمال بودن داده‌ها، میزان Sig (سطح معنی‌داری) متناظر با آن‌ها را بررسی می‌کنیم. چنانچه میزان Sig برای هر کدام از متغیرها بیشتر از ۰/۰۱ باشد، نتیجه می‌گیریم که داده‌ها از توزیع نرمال برخوردارند (به بیان صحیح‌تر، توزیع داده‌ها از توزیع نرمال انحراف قابل توجهی ندارد).

☑ نکته: بیشتر نویسندگان محتاط‌ترند و معتقدند برای این که نرمال بودن داده‌ها را بپذیریم؛ سطح معنی‌داری آزمون‌های کولموگروف-اسمیرنوف و شاپیرو-ویلک باید بیشتر از ۰/۰۵ باشد ($P > 0.05$).

☑ نکته: آزمون‌های کولموگروف-اسمیرنوف و شاپیرو-ویلک به انحراف از توزیع نرمال بسیار حساس هستند؛ در نتیجه در نمونه‌های با حجم زیاد، انحراف کوچکی از توزیع نرمال ممکن است موجب شود این آزمون‌ها توزیع متغیرها را غیر نرمال نشان دهند. در نتیجه زمانی که حجم نمونه زیاد است (حدوداً بیشتر از ۱۰۰)، بهتر است از روش‌های دیگر بررسی نرمال بودن هم استفاده شود. از جمله این روش‌ها می‌توان به شاخص‌های کجی و کشیدگی و نمودارهایی مانند هیستوگرام، ساقه و برگ و نمودار نرمال بودن Q-Q اشاره کرد.

نتایج آزمون شاپیرو-ویلک نشان می‌دهد که توزیع متغیر بهره هوشی در نمونه آماری نرمال نیست ($P < .001$). سطح معنی‌داری به دست آمده کمتر از مقدار $.001$ است و نشان می‌دهد که توزیع متغیر بهره هوشی در نمونه نرمال نیست.

Tests of Normality

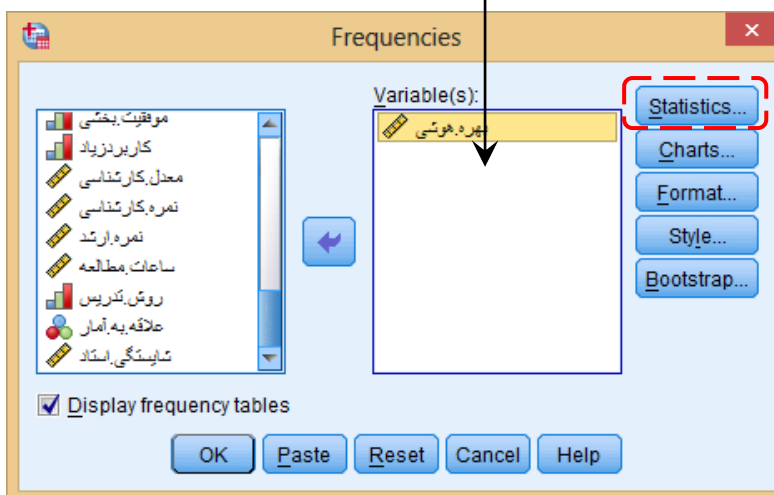
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
	آماره	درجه آزادی	سطح معنی‌داری	آماره	درجه آزادی	سطح معنی‌داری
IQ	.181	100	.000	.636	100	.000

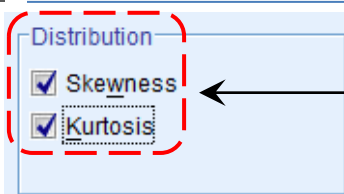
اجرا:

قصد داریم نرمال بودن متغیر بهره هوشی را با شاخص‌های کجی و کشیدگی آزمون کنیم. مسیر زیر را دنبال می‌کنیم:

Analyze ---> Descriptive Statistics ---> Frequencies

در کادر Variables متغیری که می‌خواهیم کجی و چولگی آن را محاسبه کنیم (بهره هوشی) وارد می‌کنیم. سپس گزینه Statistics را انتخاب می‌کنیم.





در کادر شکل توزیع (Distribution) گزینه‌های کجی (Skewness) و کشیدگی (Kurtosis) را انتخاب می‌کنیم. گزینه Continue و OK را انتخاب می‌کنیم.

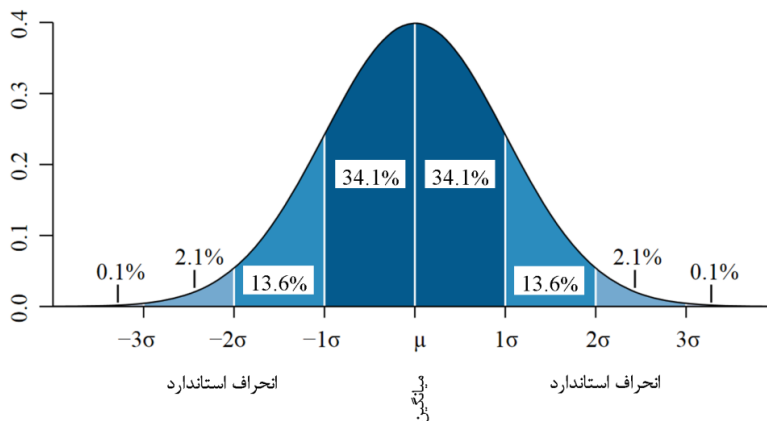
نتایج:

جدول زیر مقادیر شاخص‌های کجی و کشیدگی را نشان می‌دهد. همان طور که توضیح داده شد، مقادیر کجی و کشیدگی باید بین $+1$ تا -1 باشد تا بتوانیم نرمال بودن توزیع متغیر را بپذیریم. مقدار کجی به دست آمده برابر با -1.32 است که کمتر از مقدار -1 (منفی یک) است و نشان از این دارد که توزیع متغیر با توزیع نرمال فاصله دارد. مقدار کشیدگی هم برابر با 26.69 است که بسیار بیشتر از مقدار $+1$ (مثبت یک) است که نشان از توزیع غیرنرمال متغیر بهره هوشی دارد. توزیع متغیر دارای کجی منفی و کشیدگی مثبت است و از توزیع نرمال برخوردار نیست. پیشنهاد می‌شود یا از روش‌های ناپارامتریک برای آزمون‌های آماری استفاده شود و یا با استفاده از روش تبدیل داده‌ها، سعی شود که توزیع متغیر، نرمال یا نزدیک به نرمال شود.

Statistics

بهره هوشی

N	Valid	100
	Missing	0
Skewness (کجی)		-1.32
Std. Error of Skewness (خطای استاندارد کجی)		.24
Kurtosis (کشیدگی)		26.69
Std. Error of Kurtosis (خطای استاندارد کشیدگی)		.47



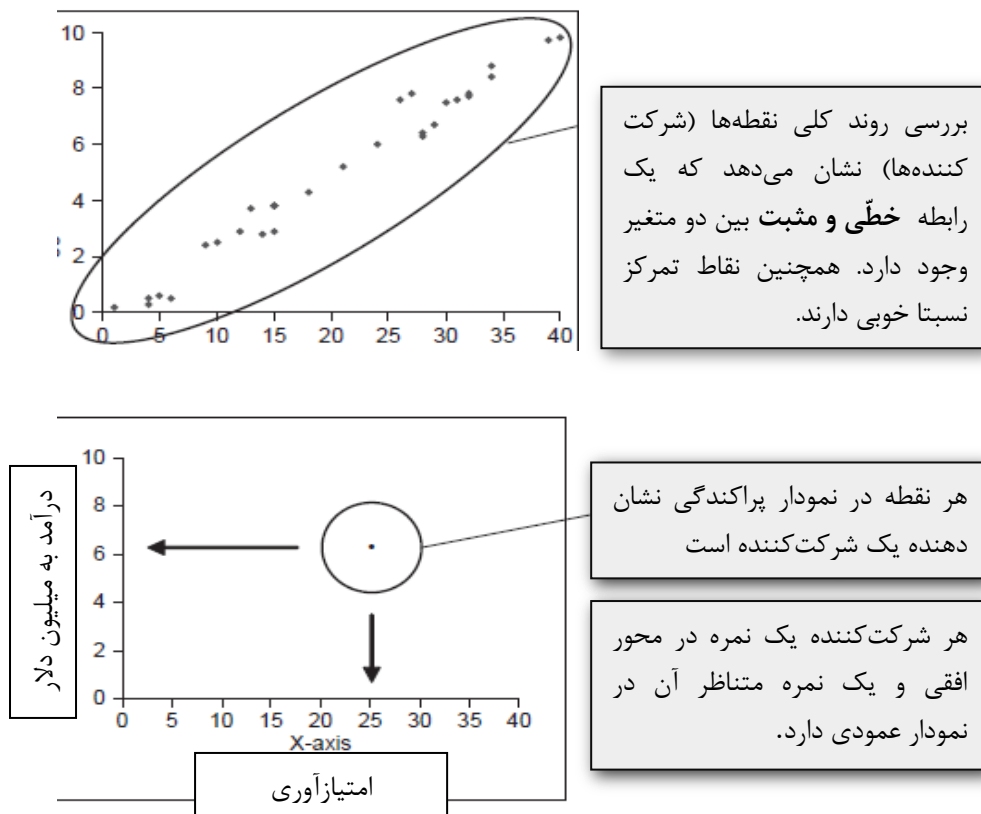
• خطی بودن رابطه (نمودار پراکندگی)

تحلیل‌های چندمتغیره (وهمبستگی دومتغیره) براین فرض استوار هستند که رابطه علی بین متغیر مستقل و وابسته خطی است (هیر و دیگران، ۲۰۱۰). چون پیش‌فرض برخی آزمون‌های آماری (مانند همبستگی پیرسون و رگرسیون خطی)، خطی بودن رابطه متغیرها با یکدیگر است؛ در نتیجه قبل از اجرای این آزمون‌ها باید از خطی بودن رابطه متغیرها مطمئن شد. یکی از روش‌های پرکاربرد برای بررسی خطی بودن رابطه، استفاده از نمودار پراکندگی است.

نمودار پراکندگی، روشی مفید برای ترسیم ارتباط بین داده‌هاست. نمودار پراکندگی یک از ساده‌ترین روش‌ها برای بررسی همبستگی و ارتباط دو متغیر است. این نمودار، نوع و جهت رابطه را به طور بصری ارائه می‌دهد و می‌توان با مشاهده نمودار از نوع رابطه بین دو متغیر و جهت (خطی یا غیر خطی و مثبت یا منفی) و شدت رابطه آگاهی تقریبی یافت. دقت شود که این نوع نمودار، معنی‌داری رابطه را نشان نمی‌دهد و برای بررسی دقیق معنی‌داری و شدت رابطه باید از آزمون‌های همبستگی (همبستگی پیرسون یا اسپیرمن) استفاده کرد. این نوع نمودار تنها مختص داده‌های فاصله‌ای/نسبی است.

با استفاده از نمودار پراکندگی می‌توان به اطلاعات دیگری هم دست یافت. مثلاً نمودار پراکندگی مواردی که اثر شدیدی بر خطرگرسین داشته باشند و به عبارتی موارد دور از محدوده (پرت یا دورافتاده) را نیز نشان می‌دهد. همچنین رابطه منحنی یا هر نوع رابطه غیر خطی را می‌توان از نمودار پراکندگی استنباط کرد.

نمودار پراکندگی نموداری است شامل دو محور افقی و عمودی و نقاطی که معرف داده‌هاست. محور افقی (محور X) معرف مقادیر متغیر مستقل است و محور عمودی (محور Y) معرف مقادیر متغیر وابسته. مختصات هر نقطه معرف مقادیر دو متغیر در هر مورد است. نمودار بعد (شکل ۳-۶) مربوط به لیگ حرفه‌ای بسکتبال آمریکا یا NBA است و نمودار پراکندگی متغیرهای درآمد سالانه به میلیون دلار و میانگین امتیازآوری در هر بازی را نشان می‌دهد. امتیازآوری را متغیر وابسته و درآمد سالانه را متغیر وابسته در نظر گرفته‌ایم. تعداد شرکت‌کنندگان برابر با ۳۰ نفر است.

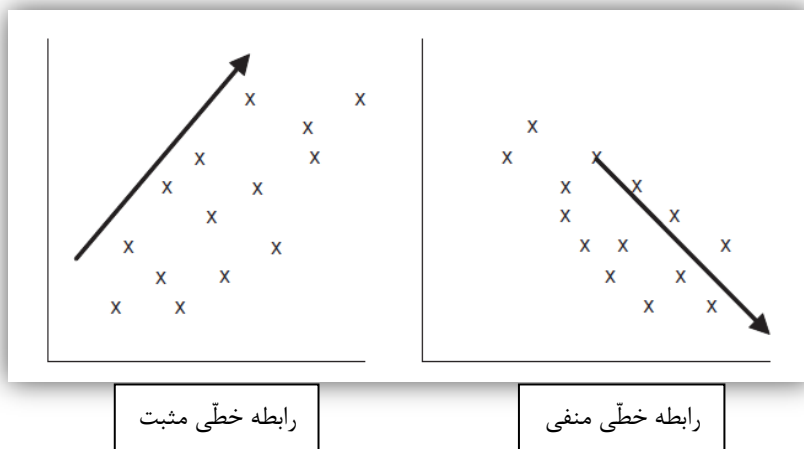


شکل ۳-۶- آشنایی با اجزاء نمودار پراکندگی

نحوه پراکندگی نقاط در نمودار پراکندگی؛ بیانگر شدت، جهت و نوع رابطه بین دو متغیر است:

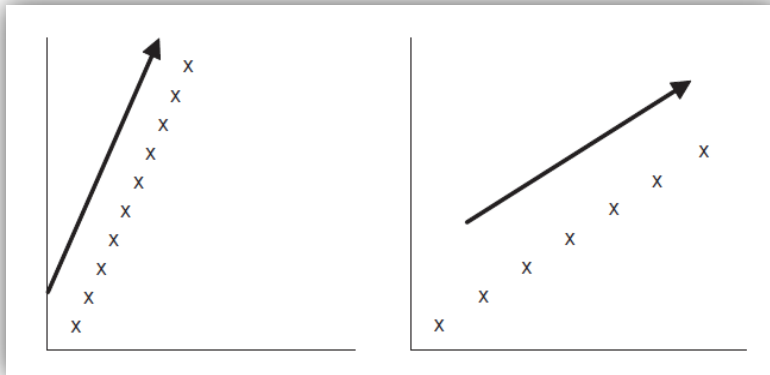
- **نوع:** آن جا که نقاط تقریباً به صورت یک خط‌اند رابطه خطی است. آن جا که تجمع نقاط به شکل منحنی است رابطه منحنی شکل است. آن جا که الگوی وجود ندارد رابطه‌ای نیز وجود ندارد.
- **جهت:** اگر شیب نقاط از چپ به راست رو به بالا (صعودی) باشد رابطه مثبت است. چنان چه این شیب رو به پایین باشد (نزول از چپ به راست) رابطه منفی است.

• شدت: هر چه تراکم نقاط بیشتر باشد و نقاط تمرکز بیشتری در اطراف خط رگرسیونی داشته باشند و فاصله نقاط از یکدیگر کمتر باشد، رابطه قوی‌تر است. به عنوان مثال، دو نمودار بعد (شکل ۷-۳) نشان دهنده رابطه خطی بین دو متغیر هستند. خطی بودن بدین دلیل است که جهت کلی و روند کلی نقطه‌ها (پاسخگویان) به طور مستقیم و در یک جهت است. نمودار سمت چپ، نشان دهنده یک رابطه مثبت بین دو متغیر است، و نمودار سمت راست یک رابطه منفی را نشان می‌دهد. در رابطه منفی، با افزایش مقدار یک متغیر، مقدار متغیر دیگر کاهش می‌یابد. در این نمودار می‌توان شدت رابطه بین دو متغیر را به طور تقریبی ملاحظه کرد. نقاط در اطراف یک خط به طور کامل متمرکز نیستند. در نتیجه بین دو متغیر رابطه کامل یا خیلی قوی وجود ندارد. یک رابطه مثبت کامل را می‌توان در نمودار بعد (شکل ۸-۳) مشاهده کرد.



شکل ۷-۳- تعیین نوع و جهت رابطه در نمودار پراکندگی

شکل ۸-۳ نشان دهنده یک رابطه خطی و مثبت بین دو متغیر است. همان‌طور که مشاهده می‌گردد تمامی نقاط روی یک خط قرار دارند و هیچ نوسان یا انحرافی بین نقاط دیده نمی‌شود. هر دو رابطه مشاهده شده یک رابطه کامل بوده و شدت رابطه در هر دو متغیر کامل و برابر با ۱ است.



رابطه کامل (جهت مثبت)

شکل ۳-۸- رابطه خطی کامل در نمودار پراکندگی

شکل شماره ۳-۹ وجود رابطه غیرخطی یا رابطه منحنی را بین دو متغیر نشان می‌دهد.



شکل ۳-۹- انواع روابط منحنی در نمودار پراکندگی

شکل ۳-۱۰ عدم وجود رابطه بین دو متغیر را نشان می‌دهد که بدین معناست که دو متغیر با هم رابطه و همبستگی ندارند.



شکل ۳-۱۰- عدم وجود رابطه در نمودار پراکندگی

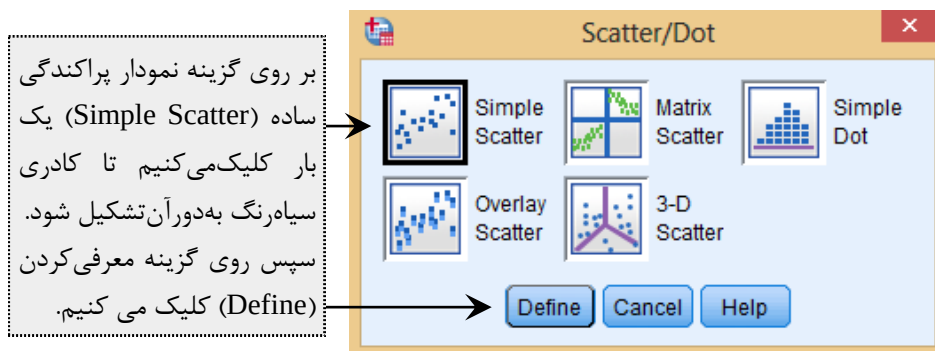
مثال

می‌خواهیم نمودار پراکندگی دو متغیر نمره کارشناسی و نمره ارشد درس SPSS را در بین دانشجویان بررسی کنیم. حدس ما این است که باید بین دو متغیر رابطه خطی مثبت وجود داشته باشد و دانشجویانی که نمره بالاتری در مقطع کارشناسی در این درس گرفته‌اند، در مقطع ارشد هم باید نمره بالاتری کسب کنند. هر دو متغیر نیز در سطح کمی (فاصله‌ای/نسبی) هستند. چنانچه بین دو متغیر رابطه خطی مشاهده شود می‌توانیم از آزمون همبستگی پیرسون یا رگرسیون خطی بین دو متغیر استفاده کنیم.

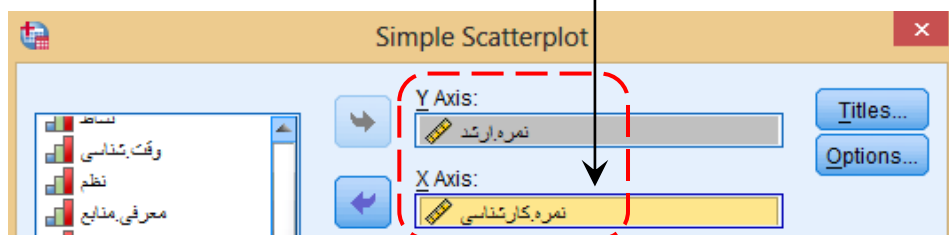
اجرا:

مسیر زیر را دنبال می‌کنیم:

Graphs ---> Legacy Dialogs ---> Scatter/Dot



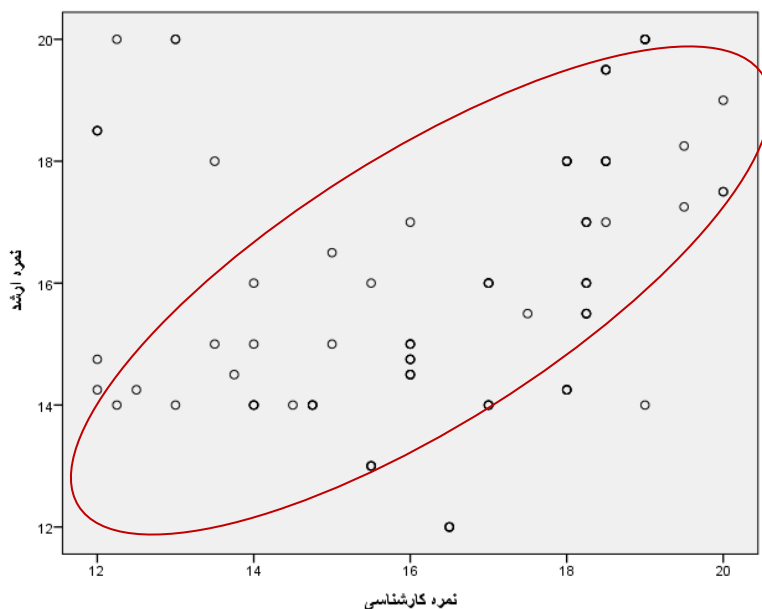
یکی از متغیرها (ترجیحا متغیر مستقل یا پیش بین) را وارد کادر محور افقی (X Axis) و متغیر دیگر را وارد کادر محور عمودی (Y Axis) می‌کنیم. درانتها گزینه Ok را انتخاب می‌کنیم.



نتایج:

نمودار پراکندگی بین نمره کارشناسی و نمره ارشد دانشجویان در درس SPSS را در ادامه (شکل ۳-۱۱) مشاهده می‌کنیم. به دور نقاط (پاسخگویان) کادر بسته‌ای کشیده ایم تا بتوان آسان‌تر رابطه بین دو متغیر را درک کرد. همان‌طور که مشاهده می‌شود بین دو متغیر یک رابطه خطی مثبت وجود دارد. یعنی افزایش نمره کارشناسی با افزایش نمره ارشد، و کاهش نمره کارشناسی با کاهش نمره ارشد همراه است.

همان‌طور که مشاهده می‌شود یک روند کلی مثبت بین دو متغیر مشاهده می‌شود اما تمرکز نقاط روی یک خط نیست، در نتیجه احتمالاً شدت رابطه بین دو متغیر قوی نیست و در نتیجه دقت پیش‌بینی بالا نیست. یعنی، اگرچه ما می‌توانیم نمره مقطع ارشد افراد را برحسب نمره کارشناسی آنان پیش‌بینی کنیم اما دقت این پیش‌بینی بالا نیست، چرا که تمامی نقاط در اطراف یک خط متمرکز نشده‌اند. در نمودار پراکندگی بعدی که خط رگرسیون آن نیز رسم شده است بهتر می‌توان عدم تمرکز نقاط را مشاهده کرد.



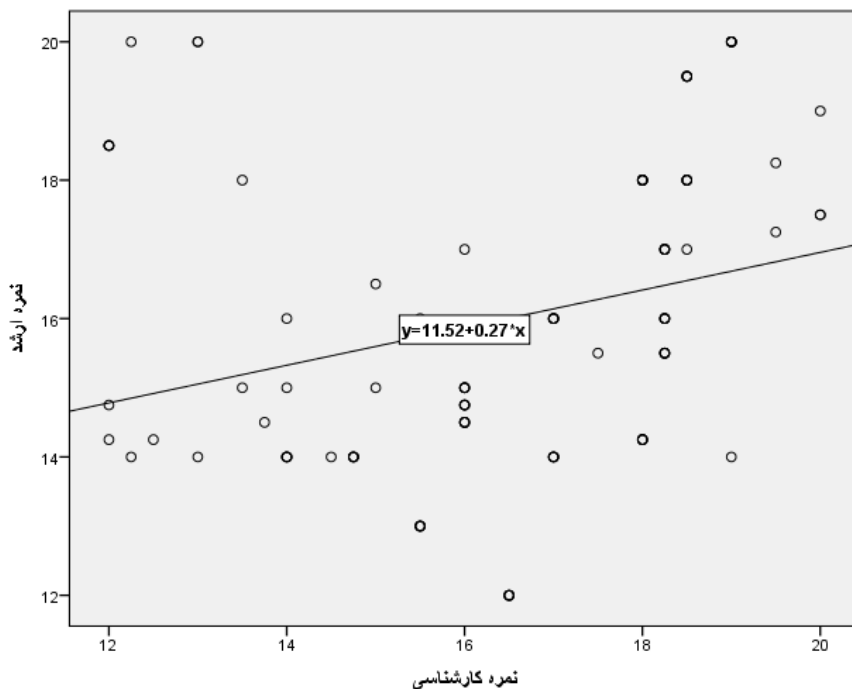
شکل ۳-۱۱- نمودار پراکندگی متغیرهای نمره کارشناسی و نمره ارشد درس SPSS

جهت افزودن خط رگرسیون به نمودار پراکندگی، تنها کافیست بر روی نمودار به دست آمده دابل کلیک (دوبار کلیک) کنیم و در پنجره جدیدی که با عنوان ویرایش گر داده



(Chart Editor) تشکیل می شود بر روی ابزار Add Fit Line at نام Total کلیک کنیم.

نمودار پراکندگی به همراه خط رگرسیون در ادامه (شکل ۳-۱۲) ارائه شده است. شیب خط رگرسیون نیز نشان دهنده وجود یک رابطه خطی مثبت بین دو متغیر است. همان طور که مشاهده می شود نقاط در اطراف خط رگرسیون به طور دقیق متمرکز نیستند، در نتیجه این عدم تمرکز از شدت رابطه بین دو متغیر می کاهد و دقت پیش بینی را نیز کاهش می دهد. هرچه نقاط به خط رگرسیون نزدیک تر باشند رابطه قوی تر است.



شکل ۳-۱۲- خط رگرسیون و معادله رگرسیون در رابطه متغیرهای نمره کارشناسی و نمره ارشد

هم‌خطی چندگانه

هم‌خطی حالتی است که در آن بین دو متغیر پیش‌بین، همبستگی قوی وجود دارد. هم‌خطی چندگانه حالتی است که در آن بیش از دو متغیر پیش بین همبستگی قوی با یکدیگر دارند. هم‌خطی چندگانه می‌تواند تفسیر نتایج رگرسیون چندگانه را تحریف کند (میزر و دیگران، ۱۳۹۱: ۲۳۶).

هم‌خطی چندگانه زمانی اتفاق می‌افتد که متغیرهای مستقل چیز یکسانی را اندازه بگیرند. زمانی که دو یا چند متغیر مستقل همبستگی بالایی داشته باشند امکان وجود "هم‌خطی چندگانه" زیاد است. در ارتباط با گویه‌های یک متغیر نیز مساله هم‌خطی چندگانه وجود دارد. گویه‌هایی که مفهوم یکسانی را اندازه می‌گیرند باید با یکدیگر همبستگی بالایی داشته باشند، اما همبستگی‌های بزرگتر از ۰/۹ بین هرگویه می‌تواند موجب مشکلات آماری شود (تاباچینک و فیدل، ۲۰۰۱).

وقتی که هدف پژوهش تنها بیشینه کردن ضریب تعیین (R^2) نیست، بلکه فهم تأثیر متقابل متغیرهای پیش‌بین است، هم‌خطی چندگانه ممکن است موجب مشکلات متعددی در تحلیل شود. یکی از مشکلاتی که از وجود هم‌خطی چندگانه ناشی می‌شود آن است که مقادیر ضرایب رگرسیون متغیرهای مستقل که همبستگی بالایی دارند تحریف می‌شود. اغلب، این ضرایب تا حدودی پایین هستند و حتی ممکن است از نظر آماری معنی‌دار هم نباشند. دومین مشکل آن است که خطاهای استاندارد وزن‌های رگرسیون متغیرهای پیش‌بین که هم‌خطی چندگانه دارند ممکن است متورم^۹ شود که در نتیجه فواصل اطمینان آن‌ها بزرگ خواهد شد، تا جایی که بعضی وقت‌ها این فواصل اطمینان، صفر را نیز در بر می‌گیرد. اگر چنین باشد به یقین نمی‌توان گفت که افزایش اندازه متغیر پیش‌بین موجب افزایش یا کاهش متغیر ملاک می‌شود. مشکل سوم آن است که اگر هم‌خطی چندگانه به اندازه کافی زیاد باشد عملیات ریاضی درونی خاص (مانند معکوس کردن ماتریس) مختل شده و برنامه آماری با سوت کشیدن^{۱۰} متوقف می‌شود (میزر و دیگران، ۱۳۹۱: ۲۳۷).

^۹ inflated

^{۱۰} screeching halt

دلایل بروز هم‌خطی چندگانه

یک علت رایج هم‌خطی چندگانه این است که پژوهش‌گران هم خرده‌مقیاس‌های یک پرسشنامه و هم نمره کل پرسشنامه را به عنوان متغیرهای پیش‌بین به کار می‌برند. بسته به این که خرده‌مقیاس‌ها چگونه محاسبه شده باشند، این امکان وجود دارد که ترکیب آن‌ها با نمره کل پرسشنامه دارای همبستگی تقریباً کامل باشد. بنابراین باید فقط خرده‌مقیاس‌ها یا نمره کل را پرسشنامه را به کار ببریم نه این که همه آن‌ها را با هم وارد تحلیل کنیم. علت رایج دیگر هم‌خطی چندگانه این است که متغیرهایی در تحلیل وجود دارند که سازه مشابهی را می‌سنجند. ما باید همه این متغیرها به جز یک متغیر را از تحلیل کنار بگذاریم، یا احتمال ترکیب آن‌ها را با یک روش معقول و منطقی مدنظر قرار بدهیم.

تشخیص هم‌خطی چندگانه

زمانی که فقط دو متغیر بررسی می‌شوند، پیدا کردن رابطه هم‌خطی از طریق وجود همبستگی بالا آسان است. فقط کافی است که همبستگی‌های پیرسون بین متغیرهای موجود در تحلیل را به عنوان پیش‌درآمد (پیش‌شرط) آزمون‌هایی مانند رگرسیون چندگانه بررسی کنیم. وقتی رابطه بین دو متغیر بسیار قوی است پرچم قرمز برافراشته می‌شود. به عنوان یک قاعده کلی سرانگشتی پیشنهاد می‌کنیم زمانی که همبستگی بین دو متغیر ≥ 0.7 یا بالاتر است احتمالاً نباید آن‌ها را در رگرسیون یا هر تحلیل چندمتغیری دیگری با هم به کار برد. اگر همبستگی بالاتر از 0.8 باشد، تقریباً بدون شک مشکل خواهیم داشت، اما با مقادیر کمتر از این مقدار هم ممکن است دشواری‌هایی وجود داشته باشد (میزر و دیگران، ۱۳۹۱: ۲۳۸).

زمانی که بیشتر از دو متغیر بررسی می‌شوند از پارامتر تحمل استفاده می‌کنیم. این پارامتر موجب می‌شود با بیرون‌راندن متغیرهای پیش‌بینی که با سایر متغیرهای مستقل همبستگی بسیار بالایی دارند روش را در برابر خطر هم‌خطی چندگانه حفظ کند. از نظر مفهومی، پارامتر تحمل، مقدار واریانس متغیر پیش‌بین است که به وسیله سایر متغیرهای پیش‌بین تبیین نشده است (R^2 بین متغیرهای پیش‌بین - ۱). دامنه مقادیر پارامتر تحمل از 0 تا 1 است و اندازه‌های پایین‌تر پارامتر تحمل نشان می‌دهد که بین متغیرهای پیش‌بین روابط قوی‌تری وجود دارد (میزر و دیگران، ۱۳۹۱: ۲۳۹). اگر اندازه

پارامتر تحمل در دامنه ۰/۴۰. باشد جای نگرانی دارد، اما اگر در دامنه ۰/۱۰. باشد مشکل آفرین است (استیونس، ۲۰۰۲). هر چه مقدار پارامتر تحمل به عدد ۱ نزدیک‌تر باشد نشان می‌دهد احتمال وجود هم‌خطی چندگانه کمتر است. لازم به ذکر است که پارامتر تحمل برای تمام متغیرهای پیش‌بین ارائه می‌شود و پارامتر تحمل هر متغیر را باید جداگانه ارزیابی کرد.

مثال

از مثال رگرسیون چندمتغیره برای یافتن هم‌خطی چندگانه استفاده می‌کنیم. متغیرهای مستقل عبارتند از: تعداد ساعات مطالعه، علاقه به درس آمار و نمره درس SPSS در مقطع کارشناسی، متغیر وابسته هم نمره درس SPSS در مقطع ارشد است. در نتیجه ما سه متغیر پیش‌بین (مستقل) و یک متغیر ملاک (وابسته) خواهیم داشت. می‌خواهیم بدانیم که آیا بین متغیرهای مستقل هم‌خطی چندگانه وجود دارد یا خیر.

جدول ۳-۳- متغیرهای ملاک و پیش‌بین

متغیرهای پیش‌بین	متغیر ملاک
ساعات مطالعه	
علاقه به آمار	نمره درس SPSS (ارشد)
نمره درس SPSS (کارشناسی)	

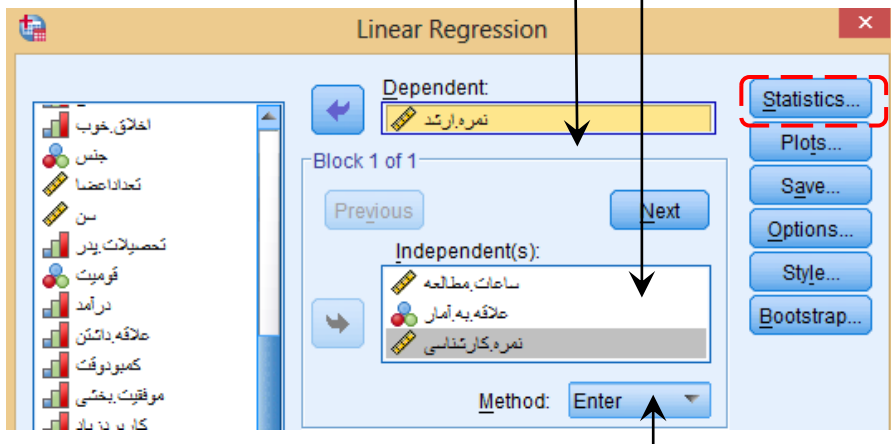
اگر تنها دو متغیر مستقل وجود داشت از روش همبستگی پیرسون استفاده می‌کردیم و چنانچه ضریب همبستگی بین دو متغیر مستقل بیشتر از ۰/۷۰. بود نشان از وجود هم‌خطی بین متغیرهای مستقل داشت و باید تنها یکی از متغیرهای مستقل را در تحلیل رگرسیون استفاده می‌کردیم. اما در مثال کتاب سه متغیر مستقل داریم و بهتر است از دستور بررسی هم‌خطی چندگانه که در دستور رگرسیون خطی وجود دارد استفاده کنیم. می‌خواهیم بررسی کنیم که آیا مسأله هم‌خطی چندگانه بین سه متغیر مستقل یا پیش‌بین تعداد ساعات مطالعه، علاقه به درس آمار و نمره درس SPSS (کارشناسی) وجود دارد یا خیر.

اجرا:

مسیر زیر را دنبال می‌کنیم:

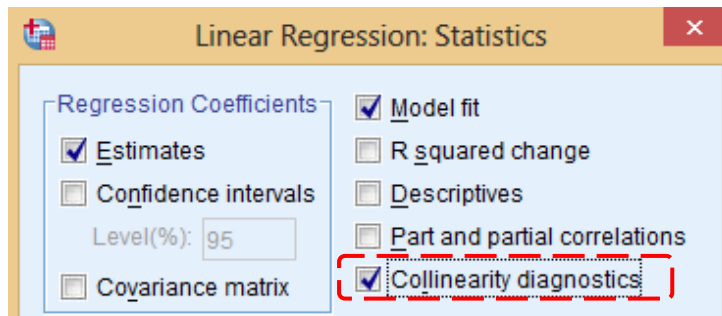
Analyze ---> Regression ---> Linear

در کادر Dependent متغیر ملاک یا وابسته را وارد می‌کنیم، یعنی متغیر نمره ارشد.
در کادر Independent متغیرهای مستقل را وارد می‌کنیم: متغیرهای ساعات مطالعه،
نمره کارشناسی و علاقه به آمار. سپس گزینه Statistics را انتخاب می‌کنیم.



روش‌های متعددی برای انجام رگرسیون وجود دارد. روش همزمان یا Enter از رایج‌ترین آن‌هاست.

در کادر Statistics گزینه تشخیص هم‌خطی یا Collinearity diagnostics را انتخاب می‌کنیم.



نتایج:

از میان خروجی‌ها و جداول ارائه شده، جدول ضرایب را جهت بررسی هم‌خطی متغیرهای مستقل بررسی می‌کنیم. در این جدول نتایج ستون آماره تحمل یا Tolerance نشان دهنده میزان هم‌خطی متغیرهای مستقل است.

همان‌طور که ذکر شد از نظر مفهومی، پارامتر تحمل، مقدار واریانس متغیر پیش‌بین است که به وسیله سایر متغیرهای پیش‌بین تبیین نشده است. اگر اندازه پارامتر تحمل در دامنه ۴۰٪ باشد جای نگرانی دارد، اما اگر در دامنه ۱۰٪ باشد مشکل‌آفرین است (استیونس، ۲۰۰۲). مطابق نتایج به‌دست‌آمده میزان آماره تحمل در بین سه متغیر مستقل بیشتر از مقدار ۴۰٪ است و حداقل مقدار آماره تحمل برابر با ۶۳٪ است. نتایج به‌دست‌آمده نشان می‌دهد که میزان هم‌خطی بین متغیرهای مستقل نگران‌کننده نیست.

آماره تحمل متغیر نمره کارشناسی ۶۳۲٪ است که بدین معناست که حدود ۶۳ درصد از واریانس متغیر نمره کارشناسی ارشد توسط دیگر متغیرهای مستقل (ساعات مطالعه و علاقه آمار) تبیین نشده است. چنانچه میزان آماره تحمل متغیری حدود ۴۰٪ یا کمتر باشد باید اقدام به بازنگری این متغیر کرده و یا آن را از تحلیل حذف کرده و یا اقدام به ترکیب این متغیر با متغیرهایی که دارای ارتباط با آن‌هاست، نمود.

Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics
	B ضریب غیراستاندارد	Std. Error خطای استاندارد	Beta ضریب استاندارد شده			Tolerance آماره تحمل
1 (Constant)	13.166	1.369		9.620	.000	
ساعات مطالعه	.352	.049	.633	7.115	.000	.638
علاقه به آمار	1.512	.318	.343	4.750	.000	.969
نمره کارشناسی	-.029	.092	-.028	-.312	.756	.632

یکسانی پراکندگی

پیش فرض یکسانی پراکندگی ایجاب می کند که متغیرهای وابسته کمی در سرتاسر دامنه متغیرهای مستقل (پیوسته یا طبقه ای)، از سطوح پراکندگی یکسانی برخوردار باشند. تخطی از این پیش فرض موجب ناهمگنی پراکندگی می شود. معمولاً وقتی یک متغیر به روش نرمال توزیع نشده باشد یا روال تبدیل داده ها به توزیع غیرقابل انتظار منجر شود، ناهمگنی پراکنش به وجود می آید (میزر و دیگران، ۱۳۹۱). در بافت تحلیل واریانس تک متغیری یا آنووا (با یک متغیر وابسته کمی و یک یا چند متغیر مستقل طبقه ای)، پیش فرض یکسانی پراکندگی به عنوان یکسانی واریانس تلقی می شود که در آن پیش فرض برابری واریانس های متغیر وابسته در همه سطوح متغیرهای مستقل برقرار است (میزر و دیگران، ۱۳۹۱: ۱۱۰).

برای آزمودن یکسانی واریانس می توان از آزمون آزمون لوین استفاده کرد. آزمون لوین (یا لوِن)، فرضیه آماری یکسانی پراکندگی واریانس ها را در تمام سطوح متغیر مستقل مورد سنجش قرار می دهد. رد فرضیه صفر (در $P < .05$) نشان دهنده تخطی از پیش فرض نابرابری پراکندگی واریانس است. شاخص آماری F بیشینه را می توان از طریق روش آنووا به دست آورد و آزمون لوین را می توان با روش های تحلیل واریانس یک راهه (آنووا) آزمون تی گروه های مستقل و Explore انجام داد.

زمانی که پیش فرض نرمال بودن چند متغیره برقرار باشد، روابط بین متغیرها از یکسانی پراکندگی برخوردار است. بنابراین، یکسانی پراکندگی ارتباط قدرتمندی با پیش فرض توزیع نرمال دارد (تاباچینک و فیدل، ۲۰۰۱). از این رو، شاید بهتر باشد که پیش از بررسی برابری واریانس ها، ابتدا تخطی از نرمال بودن را بررسی و در صورت امکان اصلاح کرد. در صورت ناهمگنی پراکندگی می توان آن را با تبدیل داده ها اصلاح کرد.

مثال

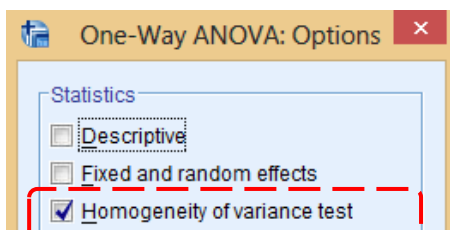
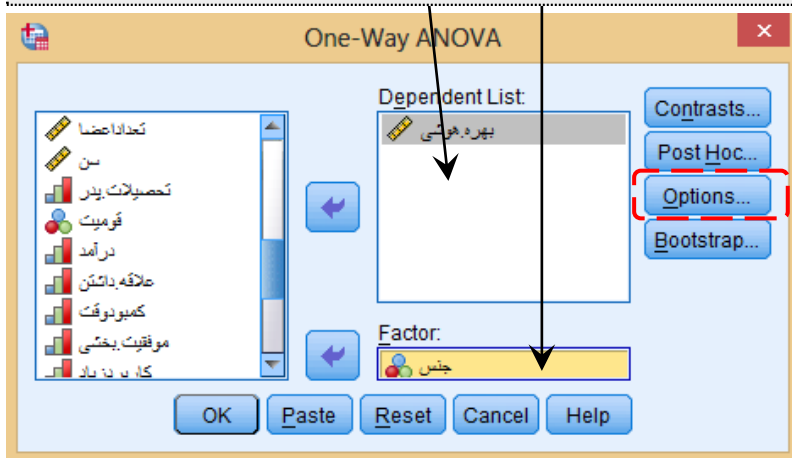
قصد داریم برقرار بودن پیش فرض یکسانی پراکندگی متغیر بهره هوشی (متغیر وابسته) را در بین زنان و مردان (متغیر مستقل) بیازماییم.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Compare Means ---> One Way ANOVA

متغیر وابسته یا بهره هوشی را وارد کادر متغیرهای وابسته می‌کنیم.
متغیر مستقل را وارد کادر عامل (Factor) می‌کنیم. گزینه Options را انتخاب می‌کنیم.



در کادر گزینه‌ها (Options) گزینه آزمون همگنی واریانس را فعال می‌کنیم. در انتها گزینه Continue و OK را انتخاب می‌کنیم.

نتایج:

جدول زیر نتیجه آزمون لوین جهت سنجش برابری واریانس‌هاست که نشان از یکسانی واریانس‌ها دارد. شاخص آماری لوین به لحاظ آماری معنی‌دار نیست ($P > 0.05$)، که یکسانی واریانس‌های متغیر بهره هوشی را در همه سطوح متغیر جنس نشان می‌دهد. لازم به ذکر است زمانی یکسانی واریانس‌ها تایید می‌شود که معنی‌داری بیشتر از ۰.۰۵ باشد.

Test of Homogeneity of Variances

Levene Statistic	df1	df2	Sig.
آماره لون	درجه آزادی		سطح معنی داری
1.255	1	98	.265



تبدیل داده‌ها

گاهی داده‌های خامی که برای تحلیل داریم مناسب گروهی از آزمون‌های آماری نیستند و برای این‌که بتوانیم از این دسته آزمون‌های آماری استفاده کنیم و همچنین دقت تحلیل را بالا ببریم باید در داده‌های خام تغییراتی ایجاد کنیم. یکی از این تغییرات، تبدیل داده‌ها نام دارد. تبدیل داده‌ها، روش‌هایی ریاضی است که برای تعدیل متغیرهایی به کار می‌رود که از مفروضه‌های آماری نرمال بودن، خطی بودن و یکسانی پراکندگی پیروی نمی‌کنند یا الگوهایی با داده‌های پرت غیرمعمول دارند (میزر و دیگران، ۱۳۹۱: ۱۱۱).

در مجموع زمانی که پیش شرط‌های آزمون‌های چندمتغیره برقرار نباشد، باید داده‌های به دست آمده را تبدیل کنیم تا امکان استفاده از برخی آزمون‌های مدنظر (عموما پارامتریک) فراهم شود.

در ابتدا باید میزان تخطی و تفاوت داده‌ها از پیش‌فرض‌های ذکر شده را تعیین کرد و در صورتی که پیش‌فرض‌ها یا پیش‌شرط‌های آماری به دست آمده دارای تفاوت قابل اعتنایی با مقدار معیار باشند از روش تبدیل داده‌ها استفاده کرد. تبدیل داده‌ها با هدف تعدیل متغیرها از جنبه علمی روشی پذیرفته شده است. البته زمانی که اختلاف داده‌ها با پیش‌فرض‌های آماری اندک باشد و به طور تقریبی مفروضات آماری برقرار باشد می‌توان از تبدیل داده‌ها صرف نظر کرد.

باید توجه داشت که تبدیل داده‌ها تا اندازه‌ای مانند شمشیر دولبه است. حسن این روش این است که می‌تواند دقت معنی‌داری تحلیل‌های آماری را افزایش دهد و عیب آن این است که ممکن است تفسیر داده‌ها را دشوارتر کند. در نتیجه باید از روش تبدیل داده‌ها به شیوه‌ای مدبرانه استفاده کرد.

دشوار کردن تفسیر داده‌ها بدین معناست که وقتی داده‌ها را تبدیل می‌کنیم، مقدار حداقل و حداکثر و شیوه توزیع متغیر و تمامی شاخص‌های میانگین و انحراف استاندارد تغییر می‌کند و با حالت معمول و عادی تفاوت پیدا می‌کند. مثلاً اگر سن افراد که به صورت کمی (نسبی) سنجیده شده است را به توان دو برسانیم شاخص‌های آماری سن افراد تغییر می‌کند و با سن‌های غیر عادی مثل ۲۵۰، ۳۰۰ و غیره مواجه می‌شویم. یا

وقتی متغیری مانند اعتماد اجتماعی داریم و با ۱۰ سوال این متغیر را سنجیدیم و دامنه میانگین این متغیر بین ۱ تا ۵ باشد، لگاریتم گرفتن از این متغیر دامنه نمرات را تغییر می‌دهد و توضیح و تفسیر متغیر را با مشکل مواجه می‌کند. یکی از راه‌های رفع این مشکل این است که هنگام گزارش یافته‌های توصیفی و شاخص‌های آماری (مانند میانگین، انحراف استاندارد و مقدار حداقل و حداکثر)؛ یافته‌ها و شاخص‌های آماری را هم به صورت عادی (قبل از تبدیل داده‌ها) و هم بعد از تبدیل داده‌ها گزارش کنیم.

انواع روش‌های تبدیل داده

روش‌های مختلفی برای تبدیل داده‌ها وجود دارد که برخی از روش‌های متداول عبارتند از: ریشه دوم^{۱۱}، لگاریتم^{۱۲}، وارون^{۱۳} و مجذور کردن^{۱۴}. بین آمارشناسان در مورد استفاده از انواع روش‌های تبدیل داده‌ها در شرایط خاص اتفاق نظر وجود ندارد. با وجود این، یک راهبرد اساسی برای به‌کاربردن روش‌های تبدیل این است که برحسب جدی بودن تخطی از پیش‌فرض‌های آماری، راهبردهای پیشرفته‌تر (سطح بالاتر) به کار بسته شود. برای مثال، تاباچینک و فیدل (۲۰۰۱) و مرتلر واناتا (۲۰۰۱) برای پیشرفت تبدیل داده‌ها از ریشه دوم (برای اصلاح تخطی متوسط)، لگاریتم (برای تخطی اساسی‌تر) و سپس ریشه دوم وارون (برای رسیدگی به تخطی شدید) جانبداری می‌کنند. مجذور کردن یک متغیر در رابطه دو متغیری غیرخطی می‌تواند مشکل غیرخطی بودن را به طور مؤثر کاهش بدهد (میزر و دیگران، ۱۳۹۱: ۱۱۳).

به طور کلی زمانی که داده‌ها نرمال نیستند از روش‌های ریشه دوم، لگاریتم و وارونه کردن برای تبدیل داده‌ها استفاده می‌شود.

^{۱۱} Square root

^{۱۲} Logarithm or Log

^{۱۳} Invers

^{۱۴} Square

جدول ۳-۴- انواع روش‌های تبدیل داده‌ها برحسب نوع انحراف از پیش‌فرض‌ها

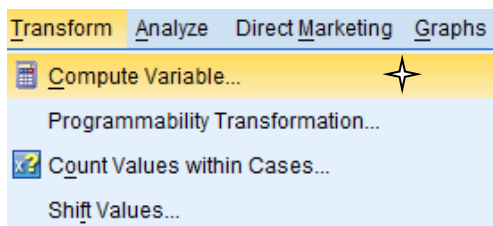
نوع مشکل در پیش‌فرض‌ها	نوع تبدیل پیشنهادی
نقض پیش‌فرض‌ها به طور کلی	ریشه دوم (برای اصلاح تخطی متوسط) لگاریتم (برای تخطی اساسی تر) ریشه دوم وارون (برای رسیدگی به تخطی شدید)
غیرخطی بودن رابطه	مجذور کردن (به توان دو رساندن)
چولگی متوسط (مثبت یا منفی)	ریشه دوم
چولگی شدید (مثبت یا منفی)	لگاریتم (Log 10)

مثال

نتایج آزمون‌های نرمال بودن نشان داد که یکی از متغیرهای مثال کتاب از توزیع نرمال برخوردار نیست. بدین جهت که قصد استفاده از آزمون‌های پارامتریک را داریم نیازمند این هستیم که توزیع این متغیر را به توزیع نرمال نزدیک کنیم. در نتیجه از تبدیل داده‌ها بهره می‌گیریم.

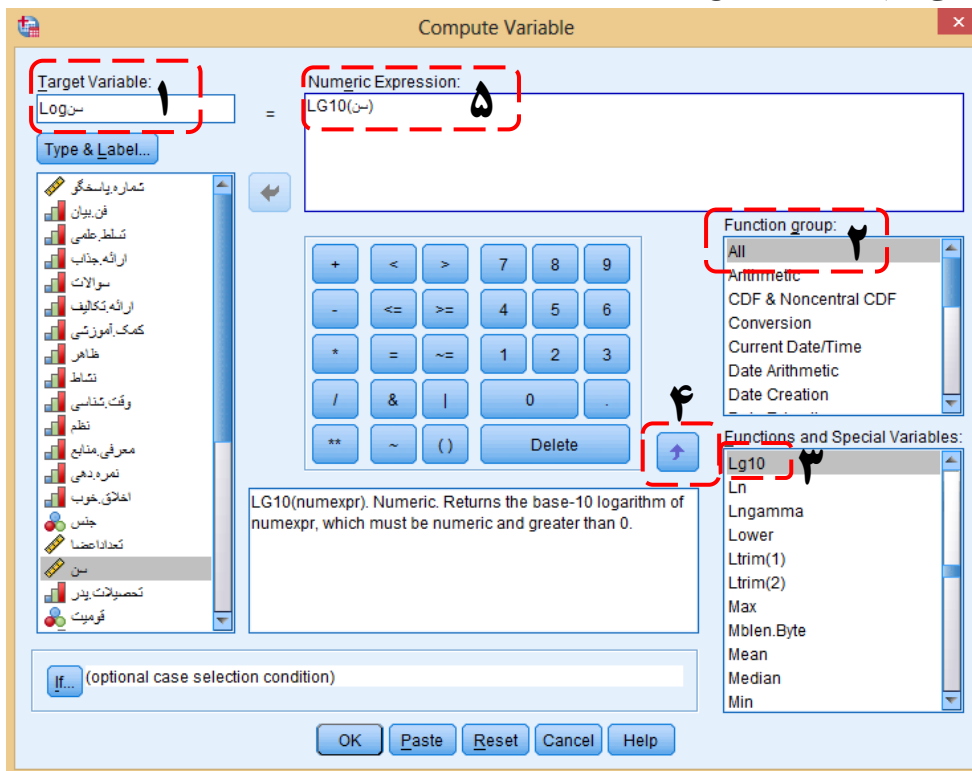
اجرا:

تبدیل داده‌ها در برنامه SPSS از طریق دستور Compute انجام می‌شود. این دستور در منوی Transform در دسترس است.



۱- در پنجره ایجادشده (Compute Variable) و در کادر Target Variable نامی برای متغیری که می‌خواهیم ایجاد کنیم انتخاب می‌کنیم (نام انتخاب شده باید با نام متغیرهای موجود تفاوت داشته باشد).

۲- برای گرفتن لگاریتم، در کادر Function group بر روی گزینه All و در کادر پایین آن (Function and Special Variables) بر روی گزینه Lg10 کلیک می‌کنیم و با استفاده از دستور  آن را وارد کادر بیان عددی (Numeric Expression) می‌کنیم. متغیر مورد نظر را از کادر متغیرها در سمت چپ انتخاب می‌کنیم و با دستور  در داخل پرانتز کادر Numeric Expression قرار می‌دهیم و در انتها گزینه OK را می‌زنیم. دستورات مطرح شده در شکل بعد نشان داده شده است.



عبارت داخل کادر بیان عددی باید بدین صورت باشد.

Numeric Expression:
LG10(سن)

چگونگی انجام سایر روش‌های تبدیل

نکته: برای مجذور کردن یا به توان رساندن متغیر کافیسیت در کادر Numeric Expression متغیر مد نظر را وارد کرده و سپس علامت * (ضرب) را از کادر ماشین حساب در کادر قرار دهیم و دوباره همان متغیر را وارد کادر کنیم.

Numeric Expression:

سن * سن

نکته: برای وارون یا معکوس کردن متغیر، عدد ۱ را در کادر Numeric Expression وارد کنیم و علامت / (تقسیم) را از کادر ماشین حساب انتخاب کرده و سپس متغیر مد نظر را وارد کادر کنیم.

Numeric Expression:

1 / سن

نکته: برای به‌دست آوردن ریشه دوم (جذر گرفتن) کافیسیت که گزینه SQRT را از کادر Function and... را انتخاب کنیم و مانند مورد لگاریتم عمل کنیم.

Numeric Expression:

SQRT(سن)

نتایج:

بعد از انجام تبدیل متغیر، متغیری جدید در صفحه داده‌ها ایجاد می‌شود که در تحلیل داده‌ها از آن استفاده می‌کنیم.

متغیر قبل از انجام تبدیل داده

متغیرهای تبدیل شده با روش‌های گوناگون

سن	لگاریتم	توان دوم	ریشه دوم	معکوس
22.0	1.34	484.00	4.69	.05
26.0	1.41	676.00	5.10	.04
29.0	1.46	841.00	5.39	.03
31.0	1.49	961.00	5.57	.03
32.0	1.51	1024.00	5.66	.03
25.0	1.40	625.00	5.00	.04
24.0	1.38	576.00	4.90	.04
26.0	1.41	676.00	5.10	.04

☑ نکته: زمانی که شیوه تعریف داده‌ها طوری باشد که برخی از موارد با کد ۰ (صفر) نشان داده شده باشند، قبل از تبدیل داده‌ها باید تغییراتی در داده‌های خام ایجاد کرد. بدین صورت که به همه داده‌های خام متغیر مورد نظر عدد ۱ را اضافه می‌کنیم تا حداقل نمره در داده‌ها برابر با ۱ باشد و بتوان به نتایج معنی‌دار و مناسبی از تبدیل داده‌ها دست یافت. برای افزودن عدد ۱ به همه داده‌ها کافیت در دستور Compute و در کادر Numeric Expression ابتدا نام متغیر را وارد کنیم و سپس علامت + (جمع) را وارد کنیم و عدد یک را بزنیم، مانند شکل پایین.

Numeric Expression:

1 + سن

برقرار کردن پیش فرض‌های آزمون‌های آماری هنگام
انحراف داده‌ها (از توزیع نرمال، رابطه خطی و یکسانی پراکندگی)

انجام برخی دستورات ریاضی بر روی داده‌های خام

تبدیل داده‌ها

مزایا: افزایش دقت تحلیل‌های آماری
معایب: دشوارتر شدن تفسیر نتایج

به توان دو رساندن — برای خطی کردن رابطه دو متغیره
ریشه دوم و لگاریتم — نرمال کردن توزیع متغیر

بخش سوم:

تقسیم و گزینش داده‌ها

گاهی در تحلیل داده‌ها نیاز داریم که تحلیل را به صورت جداگانه و برحسب طبقات یک متغیر ارائه کنیم و یا این که می‌خواهیم گروه خاصی از پاسخگویان (موردها) را انتخاب کرده و تحلیل را فقط بر روی این دسته از موردها انجام دهیم و سایر موردها را به طور موقت از تحلیل کنار بگذاریم. در این صورت باید از دستورات تقسیم کردن داده‌ها^{۱۵} و یا گزینش موردها^{۱۶} استفاده کنیم.

تقسیم کردن داده‌ها

دستور تقسیم کردن داده‌ها، داده‌ها را به گروه‌های مختلف تقسیم می‌کند و در نتایجی که در خروجی ارائه می‌شود تمام نتایج به تفکیک گروه‌ها و به صورت مجزا ارائه خواهد شد. با این فرمان، فایل داده‌ها بر حسب یک متغیر (به تعداد طبقات آن متغیر) تقسیم شده و تحلیل‌های آماری متعاقب آن برای هر طبقه آن متغیر، به طور جداگانه محاسبه می‌شود. مثلاً ما می‌توانیم تحلیل‌های آماری را برای شرکت‌کنندگان زن و مرد و یا قومیت‌های گوناگون (ترک، فارس و کرد) به صورت جداگانه به دست بیاوریم. با این دستور تمامی آزمون‌هایی که اجرا می‌شود و تمامی دستوراتی که در پنجره خروجی ارائه می‌شود به طور مجزا محاسبه می‌شود و می‌توانیم تمامی آماره‌های توصیفی (نمودار، فراوانی، میانگین، انحراف استاندارد و...) و تمامی آماره‌های استنباطی (همبستگی، رگرسیون و...) را به طور جداگانه و به تعداد طبقات متغیر مدنظر به دست بیاوریم.

¹⁵ Split File

¹⁶ Select cases

مثال

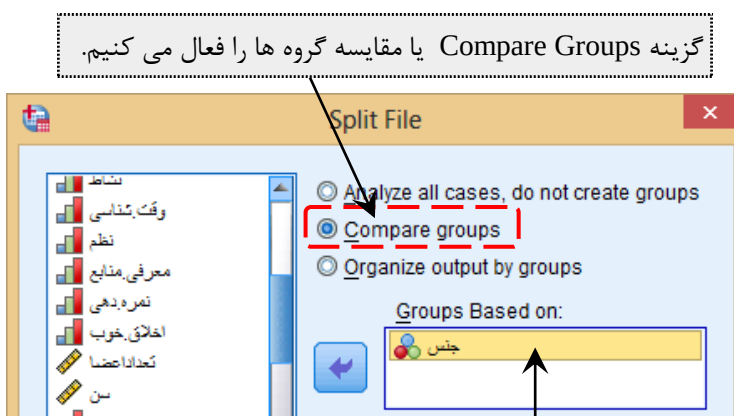
می‌خواهیم میانگین نمره درس SPSS دانشجویان ارشد را به تفکیک زن و مرد محاسبه کنیم. یعنی میانگین نمره درس SPSS کل دانشجویان را نمی‌خواهیم، بلکه می‌خواهیم میانگین دانشجویان دختر و میانگین دانشجویان پسر را به طور مجزا به دست بیاوریم.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Data ---> Split File

در پنجره ایجاد شده به شیوه زیر عمل می‌کنیم:



بعد از این که گزینه مقایسه گروه‌ها را فعال کردیم، کادر مبنای تقسیم گروه‌ها (Groups Based on) فعال می‌شود و متغیری (جنس) که می‌خواهیم بر اساس آن داده‌هایمان را تفکیک کنیم وارد کادر می‌کنیم. در انتها بر روی OK کلیک می‌کنیم.

نتایج:

با انجام این کار ما فایل داده‌هایمان را می‌شکنیم و به دو فایل جداگانه تقسیم می‌کنیم و در نتیجه تمام آزمون‌هایی را بر روی داده‌هایمان اجرا می‌کنیم به صورت جداگانه و بر حسب دختر یا پسر بودن دانشجویان می‌آید.

مطابق آموزش‌های قبل، جهت به‌دست آوردن میانگین متغیر نمره درس SPSS از دستور Frequencies استفاده می‌کنیم. جدول بعد میانگین‌های زنان و مردان را نشان می‌دهد. همان‌طور که مشاهده می‌شود میانگین کل ارائه نشده است، بلکه میانگین درس SPSS در بین دانشجویان دختر و پسر به صورت مجزا ارائه شده است.

Statistics

نمره ارشد

زن	N	Valid	50
		Missing	0
		Mean	17.15
مرد	N	Valid	50
		Missing	0
		Mean	14.81

گزینش موردها

جایگزین دیگر برای دستور تقسیم داده‌ها، انتخاب موردهای خاص و استفاده از آن‌ها در تحلیل است. دستور گزینش یا انتخاب موردها مشابه دستور تقسیم داده‌ها عمل می‌کند؛ با این تفاوت که در این دستور، داده‌ها به چند قسمت تقسیم نمی‌شوند بلکه موردهای خاصی از داده‌ها گزینش و انتخاب شده و تحلیل‌ها تنها در مورد داده‌های منتخب اجرا می‌شود. مثلاً ممکن است فقط به پاسخ‌های دانشجویان مرد یا ترک علاقه‌مند باشیم. دستور گزینش موردها به ما امکان می‌دهد تا فقط داده‌های مربوط به این دسته از دانشجویان را مورد بررسی قرار دهیم. با گزینش مواردی که قومیت‌شان ترک است تمام تحلیل‌های بعدی فقط بر روی داده‌های مربوط به دانشجویان ترک انجام خواهد شد و سایر داده‌ها (یا قومیت‌ها) به طور موقت کنار گذاشته می‌شوند.

مثال

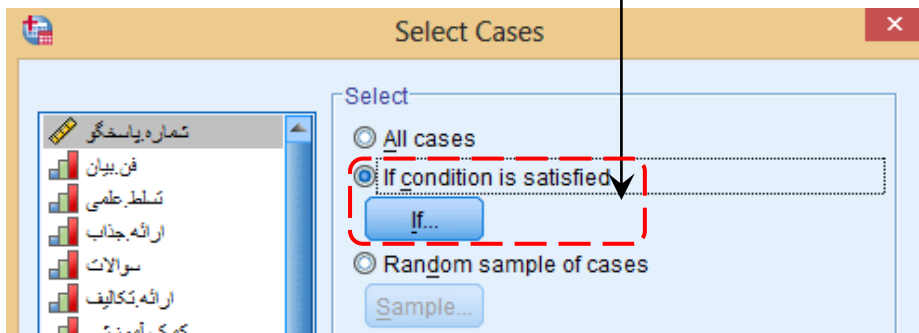
می‌خواهیم فقط نمرات دانشجویان ترک را مورد تحلیل قرار دهیم و دانشجویان سایر قومیت‌ها را به طور موقت از تحلیل‌هایمان کنار بگذاریم.

اجرا:

مسیر زیر را دنبال می‌کنیم:

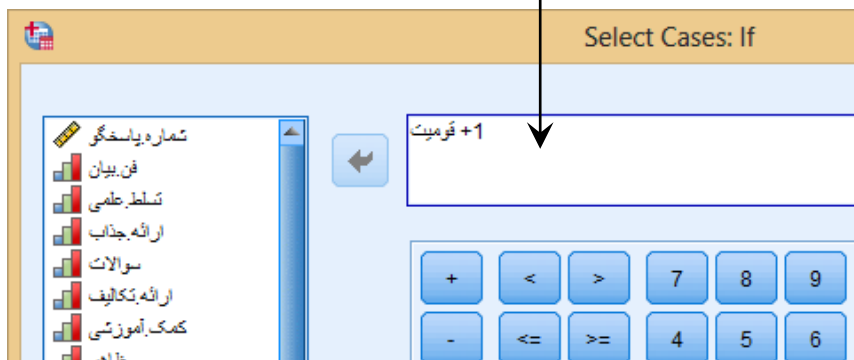
Data ---> Select cases

گزینه If condition is satisfied را فعال می‌کنیم و سپس بر روی گزینه If... کلیک می‌کنیم تا وارد پنجره گزینش موردهای موردنظر شویم.



می‌خواهیم تحلیل‌ها را فقط روی دانشجویان ترک انجام بدهیم. متغیر قومیت را از کادر سمت چپ وارد کادر می‌کنیم و از کادر دستورات و محاسبات ریاضی، گزینه مساوی (=) را انتخاب می‌کنیم و سپس گزینه ۱ را وارد می‌کنیم. کد دانشجویان ترک برابر با ۱ است، در نتیجه بعد از مساوی عدد ۱ را می‌گذاریم تا برنامه دانشجویان ترک را انتخاب کند و فقط بر روی این دانشجویان تحلیل‌های بعدی را انجام دهد.

در انتها روی گزینه Continue و سپس OK را انتخاب می‌کنیم.



بعد از این کار هر تحلیلی که بر روی داده‌ها انجام دهیم صرفاً بر روی داده‌های دانشجویان ترک انجام می‌شود و فقط تحلیل‌های مربوط به دانشجویان با قومیت ترک در خروجی برنامه نشان داده می‌شود.

حالت‌ها یا روش‌های مختلفی برای انتخاب موردهای خاص وجود دارد که در ادامه به برخی از این روش‌ها اشاره کرده‌ایم:

روش اول: اگر بخواهیم فقط یکی از قومیت‌ها را انتخاب کنیم، کد قومیت مدنظر را بعد از علامت مساوی می‌گذاریم.

مانند: $Ghomiat = 1$ یا $Ghomiat = 3$. در حالت اول تنها دانشجویان ترک را انتخاب کردیم و در حالت دوم تنها دانشجویان کرد انتخاب می‌شوند.

روش دوم: چنانچه بخواهیم کدهای کمتر از ۳ را انتخاب کنیم (دانشجویان ترک با کد ۱ و دانشجویان فارس با کد ۲) و سایر قومیت‌ها را از تحلیل کنار بگذاریم، باید از علامت بزرگتر یا کوچکتر استفاده کنیم.

مانند: $Ghomiat < 3$. در این حالت قومیت‌هایی که کد آن‌ها کمتر از ۳ است (یعنی کد ۱ و ۲) انتخاب می‌شوند.

روش سوم: می‌خواهیم دانشجویان ترک با کد ۱ و دانشجویان کرد با کد ۳ را انتخاب کنیم. یکی از روش‌های اجرای این دستور این است که قومیت‌ها را مجدداً کدگذاری کنیم و به دانشجویان کرد کد ۲ و به دانشجویان فارس کد ۳ بدهیم و مانند مثال قبل عمل کنیم.

هرکدام از حالت‌ها را که انتخاب کنیم، در تحلیل‌های بعدی فقط همان گروه‌ها یا داده‌ها انتخاب می‌شوند و مورد تحلیل قرار می‌گیرند و سایر داده‌ها به طور موقت از تحلیل کنار گذاشته می‌شوند.

تعریف

به جداول صفحات ۹۶ و ۱۸۸ رجوع کنید:

- ۱- نرمال بودن متغیر وزن و اعتماد اجتماعی (نمره کل) را بیازمایید. از آزمون شاپیرو-ویلک و کولموگروف اسمیرنوف و شاخص‌های چولگی و کشیدگی استفاده کنید. نرمال بودن متغیرهای فوق را با استفاده از نمودار هیستوگرام و رسم منحنی توزیع نرمال هم بررسی کنید.
- ۲- آیا رابطه بین متغیرهای میزان ورزش و وزن افراد و همچنین بین سن و وزن، خطی است؟ (بررسی نمودار پراکندگی)
- ۳- اگر سن و میزان ورزش را متغیرهای پیش‌بین (مستقل) در نظر بگیریم و وزن را متغیر ملاک (وابسته)، آیا بین متغیرهای سن و میزان ورزش هم‌خطی وجود دارد؟ (روش همبستگی)
- ۴- متغیر جنس افراد را تصنعی کنید. بدین صورت که به زنان کد ۰ (صفر) و به مردان کد ۱ (یک) بدهید. سپس با استفاده از آماره تحمل بررسی کنید که آیا بین متغیرهای سن، میزان ورزش و جنس افراد که بر وزن تاثیر دارند، هم‌خطی چندگانه وجود دارد.
- ۵- آیا واریانس میزان ورزش زنان و مردان برابر است؟ به بیان دیگر آیا پیش‌فرض یکسانی واریانس میزان ورزش در بین زنان و مردان برقرار است؟

به جدول صفحه ۹۶ رجوع کنید:

- ۱- میانگین و انحراف استاندارد متغیر سن را به تفکیک زنان و مردان به دست بیاورید.
- ۲- همبستگی اسپیرمن بین میزان تحصیلات و میزان ورزش را به دست بیاورید. همبستگی اسپیرمن را به تفکیک زنان و مردان جداگانه محاسبه کنید. آیا میزان همبستگی بین تحصیلات و ورزش در زنان و مردان متفاوت است.
- ۳- میانگین سن زنان را به دست بیاورید. از دستور گزینش موردها استفاده کرده و طوری عمل کنید که در خروجی برنامه تنها میانگین سن مربوط به زنان ارائه شود.

واژه نامه فصل سوم

Kurtosis	کشیدگی	Ascending	افزایشی
Central Tendency	گرایش به مرکز	Standard Deviation	انحراف استاندارد
Valid	معتبر	Tolerance parameter	پارامتر تحمل
Mean	میانگین	Dispersion	پراکندگی
Arithmetic Mean	میانگین حسابی	Data Transformation	تبدیل داده‌ها
Median or Middle	میانه	Normal distribution	توزیع نرمال
Normality	نرمال بودن	Quartile	چارک
Percentage Point	نقطه درصد	Linear	خطی
Mode	نما (مد)	Linearity	خطی بودن
Chart or Graph	نمودار	Range	دامنه تغییرات
Scatter Plot	نمودار پراکندگی	Percentage	درصد
Pie Chart	نمودار دایره ای	Decile	دهک
Bar Chart	نمودار ستونی	Percentile	صدک
Histogram	نمودار هیستوگرام	Coefficient of Variation	ضریب تغییرات
Multi - Collinearity	هم خطی چندگانه	Frequency	فراوانی
Variance	واریانس (تغییر)	Cumulative Frequency	فراوانی تجمعی
homoscedasticity	یکسانی پراکندگی (واریانس)	Descending	کاهشی
		Skewness	کجی (چولگی)

سنجش اعتبار و پایایی

محتوای فصل

« بخش اول: تحلیل عاملی اکتشافی »

« بخش دوم: ضریب آلفای کرونباخ »

بخش اول:

تحلیل عاملی اکتشافی

تحلیل عاملی^{۱۷} یک روش رایج آماری است که بسیاری از آمارشناسان آن را بکار می‌برند. تحلیل عاملی برای ایجاد ارتباط بین متغیرهای مشاهده‌شده و تعداد کمتری از متغیرهای مفهومی بنیادی طراحی شده است. تحلیل عاملی روشی است برای: ۱- کاهش تعداد متغیرها (بافتن سازه‌های پنهان یا عوامل زیربنایی) و ۲- ارزیابی اعتبار سازه^{۱۸}.

روش ساده برای شروع توضیح تحلیل عاملی این است که خاطرنشان سازیم همبستگی علّیت را نمی‌رساند: همیشه متغیر سومی می‌تواند وجود داشته باشد که رابطه بین دو متغیر را تبیین کند. در تحلیل عاملی، عملاً می‌توانیم بگوییم که دنبال آن متغیر «سوم» هستیم، اکنون می‌توان آن عامل را زیربنای همبستگی‌های بین دو یا تعداد بیشتری متغیر نامید. اما این توضیح ساده بدین معنی نیست که همه همبستگی‌ها به وسیله عوامل توضیح داده می‌شوند. تحلیل عاملی روشی است که در آن شما بررسی می‌کنید که آیا عواملی ممکن است وجود داشته باشند؛ گفته می‌شود عوامل را از متغیرها استخراج می‌کنند (بریس و همکاران، ۱۳۹۱: ۴۳۰).

به عنوان مثال در یک نظرسنجی، از مردم درباره مهم‌ترین صفاتی که کودکان باید داشته باشند (مثلاً ادب، اطاعت، پاکیزگی، قدرت تخیل، عدم وابستگی و خویش‌داری) نظرخواهی شد. ممکن است گروهی از مردم ادب، اطاعت و پاکیزگی را با اهمیت دانسته، اما ارزش چندان برای تخیل، عدم وابستگی و خویش‌داری قائل نباشند، به عبارت دیگر گروهی از متغیرها دور هم جمع می‌شوند. با کمک تحلیل عاملی می‌توان به وجود چنین الگوهایی در کل پاسخ‌های ارائه شده به پرسش‌ها پی برد. در همین مثال، با وجود این که شش پاسخ (صفات کودکان) وجود دارد، پاسخ‌ها بازتاب یا معلول دو عامل یا دو جنبه نگرشی عام‌ترند. بدین ترتیب می‌توان اولین مجموعه از متغیرها (ادب، اطاعت و پاکیزگی) را معرف بُعد دنباله‌روی و مجموعه دوم متغیرها را (قدرت تخیل، عدم

¹⁷ Factor Analysis

¹⁸ Construct validity

وابستگی و خویشن‌داری) بازتاب بعد استقلال دانست. پاسخ‌های افراد به این شش پرسش، معلول این دو عامل بنیادی (دنباله‌روی و استقلال) است.

تحلیل عاملی اکتشافی چیست؟

معمولا در پژوهش‌های اقتصادی و اجتماعی به دلیل ماهیت کار و مقیاس متغیرهای مورد سنجش، با حجم زیادی از متغیرها (گویه‌ها) روبرو هستیم. از طرفی، پژوهش‌گر برای تحلیل بهتر و دقیق‌تر داده‌ها و رسیدن به نتایجی علمی‌تر و درعین حال عملیاتی‌تر، به دنبال آن است که از یک طرف حجم داده‌ها را کاهش دهد و از طرف دیگر ساختار جدیدی را برای داده‌های خود تشکیل دهد و بر اساس فضای مفهومی دیگری، غیر از آنچه خود در بخش چارچوب نظری در نظر گرفته است، به تحلیل داده‌ها و تفسیر نتایج بپردازد. به عبارتی، در اینجا، پژوهش‌گر درصدد آزمون انطباق بین سازه نظری و سازه تجربی می‌باشد.

هدف اصلی روش‌شناسی تحلیل عاملی اکتشافی^{۱۹}، بررسی ساختار موجود در داده‌های چندمتغیره است. در این تحلیل متغیرهایی که همبستگی بالایی (چه مثبت و چه منفی) باهم دارند، احتمالا تحت تأثیر عامل‌های یکسانی هستند، اما متغیرهایی که نسبت به هم تقریبا همبستگی ندارند، از عامل‌های متفاوتی تأثیر می‌پذیرند. تحلیل عاملی اکتشافی تکنیکی آماری است که برای برآورد عامل‌ها یا متغیرهای پنهان (مکنون) از یک طرف و کاهش تعداد زیادی متغیر به تعداد کمتری عامل از طرف دیگر به کار می‌رود. بنابراین روش تحلیل عاملی با این هدف به کار برده می‌شود که از تعداد زیادی متغیر مشاهده شده، شمار معدودی عامل بیرون کشیده شود که هر یک از این عوامل از روی متغیرها و معنی آن‌ها تفسیر می‌شوند. پژوهش‌گر در این روش، هیچ نظریه اولیه‌ای ندارد و سعی می‌کند تا از بارهای عاملی برای کشف ساختار عاملی داده‌ها استفاده کند.

بنابراین، می‌توان گفت که دو هدف اصلی تحلیل عاملی اکتشافی عبارت است از:

- ✓ کشف و تعیین عامل‌های نهایی (کشف متغیرهای پنهانی که زیربنای تعداد بیشتری متغیر آشکار هستند)
- ✓ اختصاص دادن متغیرهای آشکار به عامل‌ها

¹⁹ Exploratory Factor Analysis (EFA)

نکته: به عامل، متغیر پنهان یا سازه زیربنایی می‌گویند و به متغیر آشکار، متغیر اندازه‌گیری شده یا مشاهده شده می‌گویند. در عمل منظور از متغیر آشکار، سوالات یا گویه‌های پرسشنامه است و یا هر متغیری که در برنامه SPSS تعریف می‌کنیم (در این حالت به هر ستون از داده‌ها در پنجره مشاهده داده‌ها یا Data View، متغیر آشکار می‌گویند).

دو نوع تحلیل عاملی وجود دارد که بسته به هدف پژوهش می‌توان یک یا هر دو را در تحلیل به کار برد: تحلیل عاملی اکتشافی و تحلیل عاملی تاییدی^{۲۰}. در ادامه به تعریف و مقایسه دو نوع تحلیل عاملی می‌پردازیم.

مقایسه تحلیل عاملی تاییدی و اکتشافی

تحلیل عاملی اکتشافی بیانگر یک رویکرد استقرایی است که در آن پژوهش‌گران با استفاده از یک راهبرد از پایین به بالا، از مشاهده‌های خاص نتیجه‌گیری می‌کنند. این نتیجه‌گیری در واقع همان تفسیر عامل است که بر پایه متغیرهای اندازه‌گیری شده‌ای که با آن همبستگی قوی دارند صورت می‌گیرد. بدین ترتیب این متغیرهای اندازه‌گیری شده، نشان‌گرهای عامل‌هایی می‌شوند که از تحلیل آماری به دست آمده‌اند. تحلیل عاملی اکتشافی عموماً با نرم افزار SPSS انجام می‌شود.

اگر چه راهبرد تحلیل عاملی اکتشافی به ظاهر بر پایه کارهای تجربی (استقرایی) استوار است، اما پژوهش‌ها در خلاء انجام نمی‌گیرند و پژوهش‌گران همیشه از قبل درباره ساختار زیربنایی متغیرهای تحت مطالعه، اندیشه‌هایی در ذهن خود دارند. برای مثال، پژوهش‌گران علاقه‌مند به مطالعه سازه خاص، برای تدوین پرسشنامه جدید خود، سوالات (گویه‌های) پرسشنامه را به طور اتفاقی انتخاب نمی‌کنند، بلکه انتخاب محتوای سوالات بر اساس مطابقت آن‌ها با آن چه که پژوهش‌گران درباره سازه می‌دانند (یا باور دارند که می‌دانند) هدایت می‌شود. در این مفهوم، هیچ تحلیل اکتشافی عاری از منطق پژوهش‌گر نیست (میزر و دیگران، ۱۳۹۱: ۶۴۰).

تحلیل عاملی تاییدی یک رویکرد قیاسی است که در آن پژوهش‌گران از طریق پیش‌بینی پیامدی از یک چارچوب نظری، رویکردی از بالا به پایین را به کار می‌گیرند. این پیامد عبارت از مشخص کردن مدل قبل از تحلیل آماری است که در آن متغیرهای اندازه‌گیری شده، نشانگرهای عامل هستند. علاوه بر این، فرمول‌بندی یا مشخص کردن

²⁰ Confirmatory Factor Analysis (CFA)

روابط بین عامل‌ها و روابط بین عبارت‌های خطا نیز باید توسط پژوهش‌گران از پیش تعیین شود (میزر و دیگران، ۱۳۹۱: ۶۴۰). تحلیل عاملی اکتشافی، به عنوان یک آزمون رسمی، معنی‌داری فرضیه‌ها را نمی‌آزماید. بلکه می‌توان گفت امکان وجود یک عامل ساختاری زیربنایی متغیرها را کشف می‌کند. برای آزمون فرضیه‌ها از تحلیل عاملی تاییدی استفاده می‌شود (بریس و همکاران، ۱۳۹۱: ۴۳۱). تحلیل عاملی تاییدی را با نرم افزارهایی مانند LISREL و AMOS انجام می‌دهند. در جدول ۴-۱ به طور اجمالی به مقایسه دو نوع تحلیل عاملی پرداخته‌ایم.

جدول ۴-۱- مقایسه تحلیل عاملی اکتشافی و تاییدی

نوع تحلیل عاملی	رویکرد	کاربرد	نرم افزار
اکتشافی	استقرایی	کشف عوامل پنهان (تعداد و نوع عامل‌ها از قبل مشخص نیست)	SPSS
تاییدی	قیاسی	آزمون فرضیه‌ها و آزمون مدل‌های تدوین شده	LISREL و AMOS و PLS و EQS

مفروضات تحلیل عاملی اکتشافی

- ۱- سطح سنجش: متغیرها باید حداقل در سطح اندازه‌گیری ترتیبی (رتبه‌ای) باشند.
- ۲- شکل توزیع: متغیرها باید دارای توزیع نرمال باشند.
- ۳- نوع رابطه: رابطه بین متغیرها باید تا حدودی خطی باشد.
- ۴- حجم نمونه: تعداد شرکت‌کنندگان برای یک تحلیل عاملی مورد قبول، حداقل ۱۰۰ نفر شرکت‌کننده و بعضی می‌گویند ۲۰۰ نفر یا بیشتر است (بریس و همکاران، ۱۳۹۱: ۴۴۰). تعداد شرکت‌کنندگان باید از تعداد متغیرها بیشتر باشد. کلین (۱۹۹۴) می‌گوید حداقل نسبت باید ۲ به ۱ باشد، اما هر چه این نسبت بزرگتر باشد بهتر است. بنابراین اگر می‌خواهید ساخت عاملی زیربنایی پرسشنامه‌ای ۶۰ سوالی را کشف کنید، باید حداقل ۱۲۰ نفر آزمودنی داشته باشید (همان). برخی (هیر و دیگران، ۲۰۱۰) نسبت مورد به متغیر را نسبت ۵

به ۱ پیشنهاد کرده اند که می تواند تضمینی برای ثبات و پایایی در استفاده از روش تحلیل عاملی اکتشافی باشد.

در جدول ۲-۴ برخی از اصلاحات مهم و پرکاربرد در تحلیل عاملی اکتشافی به همراه توضیح مختصری از هر کدام از آنان ارائه شده است.

جدول ۲-۴- تعریف برخی از اصطلاحات مهم در تحلیل عاملی اکتشافی

مفهوم	معادل انگلیسی	تعریف
عامل	Factor	سازه‌های زیربنایی که در فرآیند تحلیل عاملی اکتشافی به دست می آیند
بار عاملی	Factor Loading	همبستگی هر متغیر با هر عامل
استخراج	Extraction	نام فرآیندی که در آن عامل‌ها شناسایی و کشف می‌شوند
اشتراک	Communality	اندازه ای است از اینکه چه میزانی از واریانس هر متغیر به وسیله تحلیل عاملی اکتشافی تبیین می شود
مقدار ویژه	Eigenvalue	اندازه ای است برای انتخاب و استخراج عامل های نهایی و تعیین می کند چه مقدار واریانس در کل داده ها به وسیله هر عامل تبیین می شود
چرخش	Rotation	تکنیکی ریاضی است که در تحلیل عاملی اکتشافی استفاده می‌شود الگوی بارهای عاملی را به دست می‌دهد

مثال

در پژوهشی به بررسی میزان رضایت شغلی کارمندان ادارات تامین اجتماعی استان تهران پرداخته شد. ۳۰۰ نفر از کارمندان این ادارات به طور تصادفی انتخاب شدند و بین آن‌ها پرسشنامه توزیع شد و پرسشنامه‌های تکمیل شده مورد تحلیل قرار گرفت. جهت تعیین میزان رضایت شغلی کارمندان از پرسشنامه‌ای با هفت سوال استفاده شد. سوالات در قالب طیف لیکرت پنج گزینه‌ای و از بسیار کم تا بسیار زیاد درجه‌بندی

شده‌اند. جهت سنجش میزان رضایت کارمندان، از آنان پرسیده شد که از هرکدام از موارد زیر چقدر رضایت دارند:

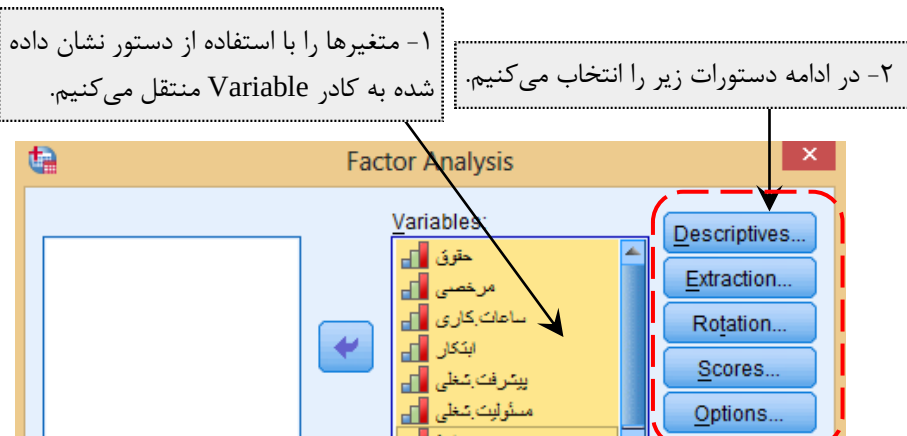
۱. حقوق دریافتی
۲. میزان مرخصی سالیانه
۳. ساعات کاری در هفته
۴. امکان بروز ابتکار و خلاقیت
۵. امکان پیشرفت شغلی
۶. میزان مسئولیت شغلی
۷. روابط بین همکاران

از روش تحلیل عاملی اکتشافی با هدف تعیین سازه‌ها یا ابعاد زیربنایی (عامل‌های) سوالات رضایت شغلی استفاده شد. حدس می‌زنیم که می‌توان هفت متغیر رضایت شغلی را در تعداد کمتری عامل خلاصه کرد اما این که آیا می‌شود سوالات (متغیرهای) رضایت شغلی را تقلیل یا خلاصه کرد و جهت تعیین تعداد عامل‌ها و این که هرکدام از متغیرها مربوط به کدام عامل می‌شوند از تحلیل عاملی اکتشافی استفاده شد. پاسخگویان به سوالات فوق جواب داده و نتایج مورد تحلیل قرار گرفت که نتایج آزمون تحلیل عاملی در ادامه گزارش شده است.

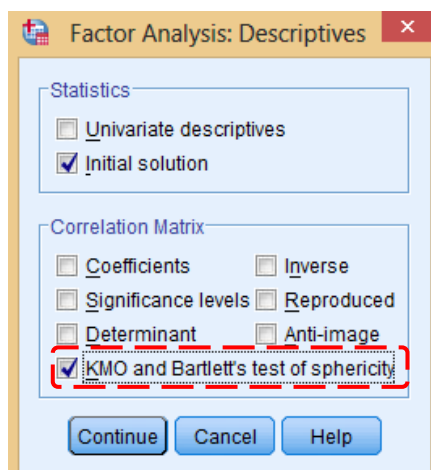
اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Dimension Reduction ---> Factor

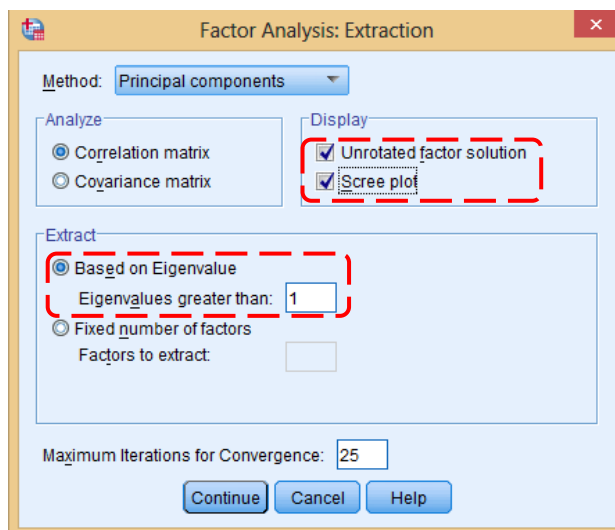


در پنجره Discriptive گزینه KMO... را فعال می‌کنیم. دو آزمون KMO و Bartlett عامل‌پذیر بودن متغیرها یا سوالات را می‌سنجند و مشخص می‌کنند که آیا امکان تحلیل عاملی بر روی متغیرها وجود دارد یا خیر. چنانچه میزان شاخص KMO بیشتر از ۰/۷۰ باشد و شاخص بارتلت در سطح معنی‌داری کمتر از ۰/۰۵ ($P < ۰/۰۵$) قرار داشته باشد، نشان می‌دهد که امکان انجام تحلیل عاملی معنی‌دار بر روی داده‌ها وجود دارد و متغیرها عامل‌پذیر هستند.

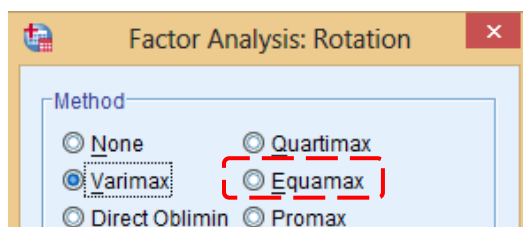


در پنجره Extraction مطابق شکل گزینه‌ها را فعال می‌کنیم. گزینه Based on Eigenvalue شاخصی به نام مقدار ویژه (Eigenvalue) را ارائه می‌دهد که برای شناسایی تعداد عامل‌ها از آن استفاده می‌شود. مبنا و معیار عدد ۱ بوده و به طور رایج عامل‌هایی که مقدار ویژه بیشتر از ۱ داشته باشند به عنوان عامل‌های نهایی و مناسب انتخاب می‌شوند.

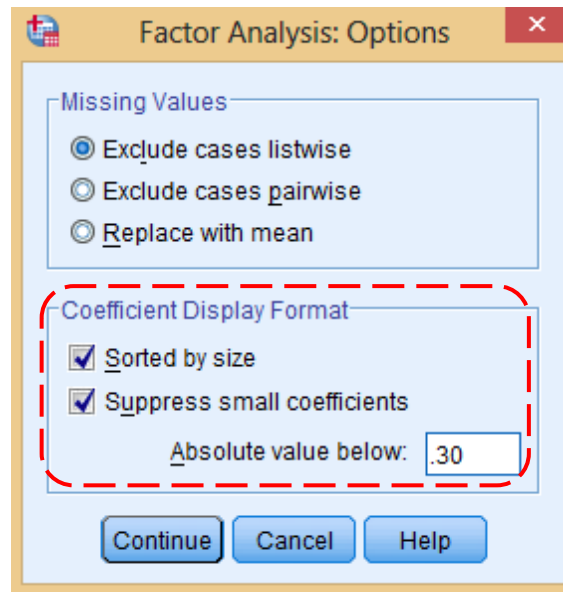
گزینه Scree plot نموداری را در خروجی به ما می‌دهد که به نام نمودار سنگریزه معروف است و می‌توان جهت تعیین تعداد عامل‌های استخراج شده مناسب از آن استفاده کرد.



پنجره Rotation نوع چرخش انتخابی را نشان می‌دهد. جهت دستیابی به نتایج مناسب و تفسیر آسان‌تر نتایج، بر روی عامل‌ها چرخش انجام می‌شود. رایج‌ترین روش چرخش عامل‌ها روش Varimax است.



فعال کردن گزینه‌های زیر در پنجره Options منجر به ارائه جداول خروجی به صورت ساده‌تر و قابل فهم‌تر می‌شود و اعداد و یافته‌ها و موارد اضافی را از برخی جداول حذف می‌کند. انتخاب گزینه Sorted by size موجب می‌شود که متغیرهای مربوط به هر عامل برحسب بار عاملی‌شان و از بزرگ به کوچک مرتب شوند و انتخاب گزینه Suppress small coefficients و قراردادن عدد ۰/۳۰، در کادر مربوطه، منجر به این می‌شود که در جدول خروجی مربوط به بارهای عاملی، تنها بارهای عاملی بزرگتر از ۰/۳۰ ارائه شوند و بارهای عاملی پایین‌تر حذف شوند.



- آزمون KMO یا اندازه کفایت نمونه‌گیری، آزمون مقدار واریانس درون داده‌هاست که می‌تواند توسط عوامل تبیین شود. این آزمون نشان می‌دهد که آیا تحلیل عاملی برای آن مجموعه متغیرها مناسب است یا خیر. به بیان دیگر این آزمون عامل‌پذیر بودن داده‌ها را نشان می‌دهد.

دامنه این آماره بین ۰ تا ۱ است. چنانچه مقدار این آماره بیشتر از 0.70 باشد همبستگی‌های موجود به طور کلی برای تحلیل عاملی بسیار مناسب اند. اگر KMO بین 0.50 تا 0.69 باشد باید دقت زیادی به خرج دارد و مقادیر کمتر از 0.50 بدان معناست که تحلیل عاملی برای آن مجموعه از متغیرها مناسب نیست (دواس، ۱۳۷۶: ۲۵۶).

- آزمون بارتلت (Bartlett) توانایی عاملی بودن داده‌ها را آزمون می‌کند و نشان می‌دهد که آیا انجام تحلیل عاملی بر روی مجموعه متغیرها به نتیجه مناسب می‌رسد یا خیر. در آزمون بارتلت اگر مقدار p (Sig) در برنامه SPSS کمتر از 0.05 باشد، توانایی عاملی بودن داده‌ها تایید می‌شود.

نتایج:

خروجی دستور تحلیل عاملی اکتشافی در ادامه گزارش شده است. جدول زیر مقادیر شاخص‌های KMO و بارتلت را نشان می‌دهد. مقدار شاخص KMO بیشتر از مقدار

معیار $.۷۰$ شده است و برابر با $.۷۹$ است. مقدار آزمون بارتلت هم معنی دار است و سطح معنی داری (مقدار P یا در برنامه SPSS مقدار Sig) کمتر از $.۰۵$ است.

مقدار آماره KMO برابر با $.۷۹۳$ و آزمون بارتلت در سطح معنی داری کمتر از $.۰۵$ ($P < .۰۵$) قرار دارد. نتایج به دست آمده نشان می‌دهد که انجام تحلیل عاملی بر روی متغیرها امکان‌پذیر بوده و متغیرها توانایی عاملی شدن را دارند.

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		.793
(آزمون KMO)		
Bartlett's Test of Sphericity	Approx. Chi-Square	1057.52
	(مقدار کای اسکوئر تقریبی)	
	df	21
Sig.		.000

جدول زیر جدول اشتراکات نام دارد. در این جدول مقدار واریانس استخراج شده هر متغیر مشخص شده است. اشتراکات تعیین می‌کنند چه مقدار از واریانس هر متغیر به وسیله عوامل نهایی (عامل‌هایی که مقدار ویژه بیشتر از ۱ دارند) تبیین می‌شود. این مقدار واریانس، واریانسی است که توسط عامل‌های نهایی (استخراج شده) تبیین شده است. هرچقدر میزان واریانس استخراج شده هر متغیر نزدیک‌تر به ۱ باشد، نشان می‌دهد که عامل‌های استخراج شده مناسب‌تر است. عموماً حداقل میزان واریانس استخراج شده برای هر متغیر را $.۵۰$ (یا ۵۰ درصد) در نظر می‌گیرند.

نتایج جدول نشان می‌دهد که عامل‌های نهایی یا استخراج شده توانسته‌اند میزان واریانس بیشتر از $.۵۰$ (یا ۵۰ درصد) را برای هر متغیر تبیین کنند. چنانچه میزان واریانس تبیین شده متغیری کمتر از $.۵۰$ باشد باید با احتیاط با آن رفتار کرده و با در نظر گرفتن بار عاملی آن متغیر و همچنین اهمیت وجود آن متغیر در پژوهش، اقدام به حذف، اصلاح یا حفظ آن متغیر کرد.

Communalities

	Initial	Extraction واریانس استخراج یا تبیین شده
حقوق	1	.678
مرخصی سالیانه	1	.566
ساعات کاری	1	.711
بروز ابتکار	1	.705
پیشرفت شغلی	1	.761
مسئولیت شغلی	1	.840
روابط بین همکاران	1	.826

- **مقدار ویژه:** اندازه ای است که تعیین می کند چه مقدار از واریانس در کل داده ها به وسیله یک عامل تبیین می شود. هرچه مقدار ویژه یک عامل بیشتر باشد مقدار بیشتری از واریانس توسط آن عامل تبیین می شود. اندازه مقدار ویژه را می توان به منظور تعیین این امر به کار برد که آیا به اندازه کافی، آن عامل واریانس تبیین می کند تا عامل، عاملی مفید باشد. (بریس، کمپوسنلگار، ۱۳۹۱: ۴۴۲). عامل هایی که مقدار ویژه بیشتر از ۱ دارند بهترین عامل ها هستند. نکته: اگر مقدار ویژه هر عامل را بر تعداد کل متغیرهایی که وارد تحلیل کرده ایم تقسیم کنیم و نتیجه را در عدد ۱۰۰ ضرب کنیم، درصد واریانس تبیین شده توسط آن عامل بدست می آید.
- **درصد واریانس کل:** پژوهش گران تحلیل عاملی دارند با راه حل هایی کار کنند که آشکارا واریانس بیشتری را تبیین می کنند تا مطمئن شوند که در فرآیند پردازش داده ها، اطلاعات زیادی از دست نرفته است. یک قانون خیلی سخت توسط برخی از نویسندگان پیشنهاد شده است (نظیر تاباچینک و فیدل، ۲۰۰۱) که راه حل حداقل باید ۵۰ درصد واریانس را تبیین کند (میزر، گامست و گارینو، ۱۳۹۱: ۵۹۱). این بدین معناست که درصد واریانس تجمعی برای عامل هایی که مقدار ویژه بیشتر از ۱ دارند، باید بیشتر از ۵۰ درصد باشد.

از جدول بعد (واریانس تبیین شده کل یا Total Variance Explained) برای تعیین تعداد عامل های استخراج شده استفاده می شود. همچنین درصد واریانس کل متغیرهای پژوهش که توسط هر عامل تبیین می شود در جدول گزارش شده است. برای تعیین تعداد عامل ها از مقدار ویژه کمک می گیریم. همان طور که ذکر شد حداقل مقدار ویژه برای انتخاب عامل های نهایی مقدار ۱ است و عامل هایی که مقدار ویژه بیشتر

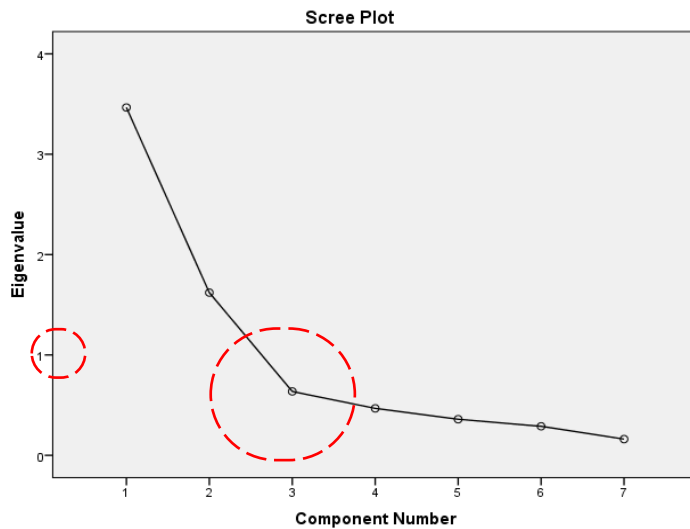
از ۱ داشته باشند جزء عامل‌های نهایی محسوب می‌شوند. نتایج نشان می‌دهد که دو عامل مقدار ویژه بیشتر از ۱ دارند و تعداد عامل‌های استخراج شده ۲ عامل است.

درصد واریانس تبیین شده توسط هر عامل در جدول نشان می‌دهد که عامل اول توانسته است ۴۹.۵ درصد از واریانس تمامی متغیرهای (هفت متغیر) پژوهش را تبیین کند. این مقدار برای عامل دوم ۲۳.۱ درصد است. در مجموع دو عامل استخراج شده نهایی توانسته‌اند ۷۲.۶ درصد از واریانس تمامی متغیرهای پژوهش را تبیین کنند. لازم به ذکر است که برخی نویسندگان حداقل مقدار واریانس تبیین شده تجمعی را ۵۰ درصد پیشنهاد کرده‌اند و منظور این است که عامل‌های استخراج شده نهایی باید بتوانند حداقل ۵۰ درصد از واریانس متغیرهای پژوهش را تبیین کنند (البته این مقدار قراردادی است و بستگی به تصمیم پژوهش‌گر و ملاحظات دیگر دارد و می‌تواند این مقدار را کمتر یا بیشتر در نظر گرفت).

Total Variance Explained

Component عنصر یا عامل‌های استخراج شده	Initial Eigenvalues مقدار ویژه		
	Total مقدار ویژه	% of Variance درصد واریانس تبیین شده	Cumulative % درصد واریانس تجمعی
1	3.466	49.51	49.51
2	1.621	23.15	72.66
3	.636	9.09	81.75
4	.468	6.68	88.43
5	.359	5.13	93.57
6	.289	4.13	97.69
7	.161	2.31	100

نمودار سنگریزه به صورت بصری تعداد عامل‌های استخراج شده را نشان می‌دهد. با در نظر گرفتن مقدار ویژه ۱ در محور عمودی می‌توان تعداد عامل‌های نهایی را مشخص کرد. لازم به ذکر است که نتایج نمودار سنگریزه با نتایج جدول قبل (واریانس تبیین شده کل) یکسان بوده اما به صورت بصری نتایج را ارائه می‌دهد. وجود شیب تند بین عامل‌ها می‌تواند مبنایی تکمیلی جهت گزینش عامل‌های نهایی باشد. همان‌طور که نمودار نشان می‌دهد بین عامل دوم و سوم شیب تندی وجود دارد و بعد از عامل دوم شیب مقدار ویژه به طور محسوسی کاهش پیدا می‌کند.



شکل ۴-۱- نمودار سنگریزه جهت انتخاب تعداد عامل های نهایی

جدول ماتریکس عناصر (Component Matrix) نشان‌دهنده بارهای عاملی متغیرهای پژوهش قبل از چرخش است. بدین علت که قبل از انجام چرخش، متغیرها در روی عامل‌ها به خوبی تفکیک نشده‌اند و تفسیر نتایج دشوار است، در نتیجه در ارتباط با اختصاص‌دادن متغیرها به عامل‌ها، از جدول بارهای عاملی بعد از چرخش استفاده می‌کنیم و از جدول زیر عبور می‌کنیم.

Component Matrix^a

	Component عامل های نهایی	
	1	2
مسئولیت شغلی	.885	
روابط بین همکاران	.857	-.304
پیشرفت شغلی	.850	
بروز ابتکار	.797	
ساعات کاری	.383	.751
مرخصی سالیانه	.325	.678
حقوق	.581	.583

- استخراج مقدماتی عامل ها (قبل از چرخش) مشخص نمی کند که چه متغیرهایی به چه عامل هایی تعلق دارند. غالبا در ماتریس قبل از چرخش، بسیاری از متغیرها بار چند عامل می شوند و برخی عامل ها هم تقریبا حامل تمام متغیرها هستند. برای تشخیص این که چه متغیرهایی به چه عاملی تعلق دارند و نیز برای تفسیرپذیرتر کردن عامل ها وارد مرحله سومی به نام «چرخش عامل» می شویم. به طور مطلوب نتیجه چرخش رسیدن به عامل هایی است که فقط بعضی از متغیرها بار آن ها می شوند و نیز رسیدن به متغیرهایی است که فقط بار یک عامل می شوند (دواس، ۲۵۹۱۳۷۶). پرکاربردترین روش چرخش، **چرخش واریماکس** است.
- حداقل مقدار بار عاملی: بارهای عاملی، همبستگی های متغیرها با عامل هاست. چنانچه این همبستگی ها بیشتر از ۰/۶ باشند (بدون توجه به علامت منفی یا مثبت)، به عنوان بارهای عاملی بالا و چنانچه **بیشتر از ۰/۳** باشند به عنوان بارهای عاملی نسبتا بالا در نظر گرفته می شوند. بارهای کمتر از ۰/۳ را می توان نادیده گرفت (کلاين، ۱۲:۱۳۸۰).

جدول بعد مربوط به بارهای عاملی بعد از چرخش (از نوع واریماکس) است. این جدول که جدول عناصر چرخش یافته نام دارد، جدول اصلی در ارتباط با متغیرهایی است که متعلق به هر عامل هستند. میزان بار عاملی هر متغیر در این جدول گزارش شده است. در وضعیت عادی، هر متغیر با تمامی عامل ها همبستگی دارد و بارعاملی هر متغیر با تمامی عامل ها در جدول ارائه می شود اما با توجه به دستوری که هنگام اجرای دستور تحلیل عاملی به برنامه دادیم، مقادیر بار عاملی کمتر از ۰/۳۰ در خروجی نشان داده نشده است و تنها مقادیر بارهای عاملی قابل قبول (۰/۳۰ و بیشتر) نمایش داده شده اند. همچنین مقادیر بار عاملی به طور نزولی گزارش شده اند.

با استفاده از تحلیل عاملی اکتشافی توانستیم عامل ها یا سازه های پنهان موجود در هفت متغیر شرایط کاری را استخراج کنیم و ۲ عامل که زیربنای هفت متغیر پژوهش هستند را شناسایی کنیم. نتایج گویای این است که ۴ متغیر در عامل اول قرار می گیرند (بار می شوند) و ۳ متغیر در عامل دوم قرار می گیرند. متغیرهای روابط بین همکاران، میزان مسئولیت شغلی، امکان پیشرفت شغلی و امکان بروز ابتکار و خلاقیت در عامل اول قرار می گیرند و متغیرهای ساعات کاری در هفته، حقوق دریافتی و میزان مرخصی سالیانه در عامل دوم قرار می گیرند.

جهت نام گذاری عامل ها، از محتوا و معنای متغیرها استفاده می کنیم. با توجه به محتوا و مفهوم متغیرهای قرار گرفته در عامل های اول و دوم، می توان عامل اول را عامل محتوای

کار و عامل دوم را عامل شرایط کاری نامید. همان طور که محتوای متغیرها نشان می‌دهد متغیرهای مسئولیت شغلی، روابط بین همکاران، پیشرفت شغلی و بروز ابتکار و خلاقیت بر محتوای کار تاکید دارند و متغیرهای ساعات کاری، مرخصی سالیانه و حقوق دریافتی بر شرایط کار تاکید دارند.

مقادیر بارهای عاملی هر متغیر قابل قبول بوده و برای تمامی متغیرها بیشتر از مقدار ۰/۷۰ است که مقدار مناسب و بالایی به شمار می‌آید و نشان از این دارد که نتیجه تحلیل عاملی مطلوب بوده و عامل‌های مناسبی استخراج شده اند. یکی از مهم‌ترین نشان‌های مطلوب بودن تحلیل عاملی، تعداد کم عامل‌های استخراج شده، برخورداری هر عامل از حداقل ۴ متغیر و مقادیر بالای بارهای عاملی است. ضمن این که به این موارد باید اضافه کرد که زمانی که هر متغیر تنها با یک عامل دارای همبستگی بالا (بار عاملی) باشد و با عامل‌های دیگر دارای همبستگی ضعیف و ناچیزی (کمتر از ۰/۳۰) باشد، نشان از مناسب بودن عامل‌های استخراج شده دارد.

Rotated Component Matrix^a

	Component	
	1	2
روابط بین همکاران	.906	
مسئولیت شغلی	.905	
پیشرفت شغلی	.857	
بروز ابتکار	.836	
ساعات کاری		.842
حقوق		.768
مرخصی سالیانه		.752

گزارش:

داده‌ها به وسیله تحلیل عناصر اصلی، با چرخش واریماکس، تحلیل شدند. شاخص‌های توانایی عامل‌شدن مناسب بودند و نتیجه آزمون KMO و بارتلت تایید می‌کنند که راه‌حل، راه‌حل خوبی بوده است (مقدار KMO برابر با ۰/۷۹۳، و آزمون بارتلت در سطح معنی‌داری کمتر از ۰/۰۵، قرار دارد). دو عامل با مقدار ویژه بیشتر از ۱ یافت شدند و نمودار سنگریزه هم دو عامل را تایید کرد. دو عامل استخراج شده در مجموع حدود ۷۳ درصد واریانس تمامی متغیرها یا شاخص‌ها را تبیین می‌کنند. عوامل به دست آمده را می‌توان عوامل شرایط کار و محتوای کار نامید. جدول زیر عوامل و متغیرهایی که بر روی آن‌ها بار شده‌اند را نشان می‌دهد.

عامل‌های نهایی استخراج شده و شاخص‌ها یا معرف‌های هر عامل	
عامل دوم: شرایط کار	عامل اول: محتوای کار
ساعات کاری	مسئولیت شغلی
مرخصی سالیانه	روابط بین همکاران
حقوق دریافتی	پیشرفت شغلی
	بروز ابتکار و خلاقیت

در گزارش نتایج، مقدار KMO، درصد واریانس تبیین‌شده برای هر گویه، مقدار ویژه و درصد واریانس استخراج شده به وسیله هر عامل و بارهای عاملی بعد از چرخش را در جداول مناسب ارائه دهید.

چند نکته:

الف) نکته‌ای که در تحلیل عاملی وجود دارد این است که قطع‌نظر از متغیرهای مورد استفاده، مجموعه‌ای از عوامل به دست می‌آیند که چه بسا بی‌معنا باشند. از آن‌جا که حاصل کار (استخراج عوامل) مبتنی بر همبستگی بین متغیرهاست چه بسا عوامل به دست آمده فاقد هرگونه ربط منطقی و مفهومی باشند. ممکن است دریابیم متغیر تحصیلات، سن و درآمد دارای همبستگی تجربی‌اند و در نتیجه آن‌ها را عاملی در تحلیل عاملی به شمار آوریم. اما این بدان معنا نیست که واقعا یک بعد نگرشی بنیادی را می‌سنجد. ممکن است متغیرها به جای این که بازتاب عاملی بنیادی باشند خود دارای رابطه علی باشند. در انتخاب متغیرهایی که در تحلیل عاملی وارد می‌شوند باید فرض ما

این باشد که همبستگی بین متغیرها غیرعلّی است. در واقع همبستگی بین متغیرها باید محصول عامل دیگری، یعنی عامل مشترک سوّمی باشد. در عمل این بدان معناست که هنگام انتخاب متغیرهای مورد تحلیل باید از انتخاب متغیرهایی که احتمالا علّت متغیرهای دیگر تحلیل هستند، اجتناب کرد.

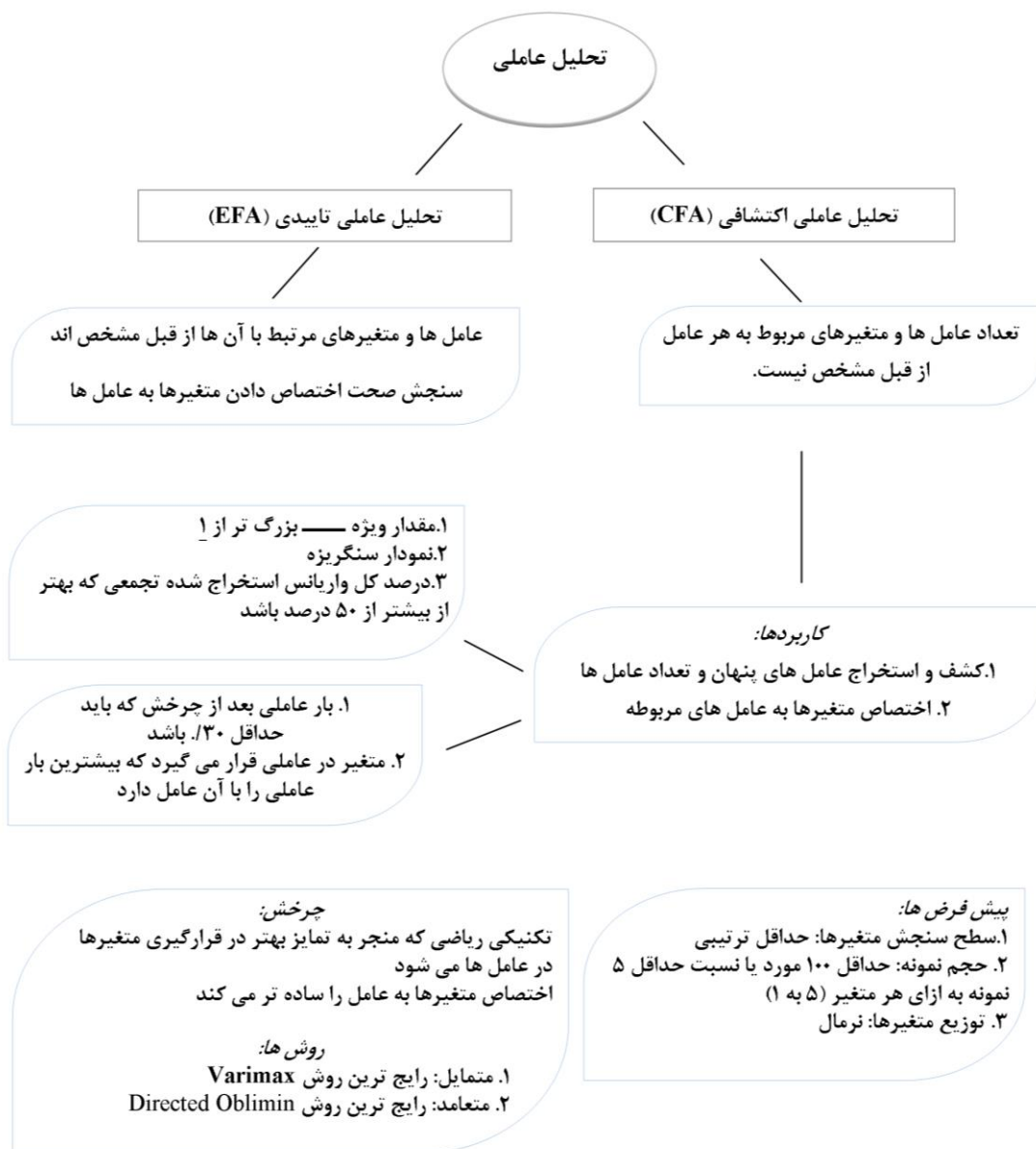
ب) بهترین مدل تحلیل عاملی، مدلی است که نتیجه آن کمترین تعداد عامل (استخراج شده) با بیشترین درصد واریانس تبیین شده تجمعی باشد.

ج) گاهی مقدار ویژه تعداد زیادی عامل بیشتر از ۱ می‌شود که در این حالت باید تعدادی از آن‌ها را حذف کرد. یکی از راه‌های کاهش عامل‌ها، بررسی مقدار کل واریانس تبیین شده مجموعه تمام متغیرها به وسیله شماری از عامل‌هاست. در اینجا باید نقطه-ای را که بعد از آن افزایش عامل‌ها منجر به افزایش زیادی در مقدار کل واریانس تبیین شده نمی‌شود تبیین کرد و فقط عامل‌های بالای این نقطه را انتخاب کرد.

د) اگر متغیری بار اندکی بر تمام عامل‌ها داشته باشد نشان می‌دهد که اشتراک متغیر پایین است و باید آن متغیر را از تحلیل حذف کنیم.

و) بار منفی: علامت بار متغیرها به معنای شدّت رابطه متغیر و عامل نیست. علامت منفی بارهای عاملی فقط در ارتباط با علامت متغیرهای دیگر معنا پیدا می‌کند: علامت‌های مختلف فقط نشان دهنده جهت متضاد رابطه متغیرها با عامل‌ها است و به همین دلیل قبل از تحلیل باید جهت کدگذاری متغیرها را یکسان کرد.

ه) لازم است هر عامل به وسیله تعداد معقولی از متغیرها نشان داده شود. این امر دست کم دو دلیل دارد: (۱) ما می‌خواهیم ماهیت عامل را به روشنی بدانیم و هر اندازه متغیرهای بیشتری روی عامل بار داشته باشند، تفسیر ما بر اطلاعات بیشتری استوار خواهد بود و (۲) اگر قصد داریم که خرده‌مقیاس‌هایی بر پایه ساختار عاملی ایجاد کنیم، به متغیرهای کافی نیاز داریم که با آن عامل همبستگی بالا داشته باشند تا خرده‌مقیاس‌های ما از اعتبار قابل قبول برخوردار باشد. اگر چه فراهم آوردن تعداد دقیق متغیرهایی که این ملاک‌ها را داشته باشند دشوار است، به طور کلی می‌توان گفت که هر عامل نباید کمتر از ۴ یا ۵ متغیر داشته باشند (میزر و دیگران، ۱۳۹۱: ۶۰۰).



بخش دوم:

پایایی (آلفای کرونباخ)

تفاوت‌های فرهنگی یا تغییر در زبان به دلیل مرور زمان، یا تفاوت‌های مربوط به نمونه‌ها و موقعیت‌ها ممکن است بر یک مقیاس تأثیر بگذارد. بنابراین، چه وقتی که شما یک مقیاس می‌سازید و چه هنگامی که از یک مقیاس موجود استفاده می‌کنید و چه هنگامی که یک مقیاس را ترجمه و بومی می‌کنید تحلیل داده‌ها از نظر پایایی^{۲۱} کار مناسبی است.

یک ابزار اندازه‌گیری زمانی پایایی دارد که به نتایج یکسانی در موقعیت‌های تکراری منجر شود (استراب و دیگران، ۲۰۰۴). پایایی یا قابلیت اعتماد به معنای آن است که آیا روش انتخاب شده، موضوع مورد نظر را به طور دقیق می‌سنجد یا خیر. در واقع، ابزار اندازه‌گیری تا چه حد پایایی (قابلیت تکرار) دارد و اگر با همان واحد تحلیل به طور مکرر به کار رود نتایج یکسانی به دست می‌آید یا خیر. به عبارت دیگر پایایی به میزان ثبات و انسجام درونی اجزای یک مفهوم و این که در صورت تکرار آزمون با یک ابزار و در شرایط مشابه؛ نتایج حاصله تا چه میزان مشابه‌اند اطلاق می‌شود (صفری شالی، ۱۳۸۶: ۷۹). روش‌های متعددی برای سنجش پایایی وجود دارد که سنجش پایایی به روش همسازی درونی یا سازگاری اجزاء از رایج‌ترین آن‌هاست.

همسازی درونی یا سازگاری اجزاء

از رایج‌ترین روش‌های ارزیابی پایایی مقیاس (ابزار اندازه‌گیری)، همسازی درونی است. همسازی درونی یا سازگاری اجزاء، آزمونی است برای بررسی سازگاری پاسخ‌های فرد با همه عناصر ابزار اندازه‌گیری. توضیح این که تا وقتی اجزاء آزمون یا ابزار به طور مستقل مفهوم یکسانی را اندازه بگیرند با یکدیگر همبستگی دارند. مشهورترین آزمون برای

²¹ Reliability

پایایی سازگاری اجزاء عبارت است از ضریب آلفای کرونباخ^{۲۲} (برای پرسش‌های ترتیبی یا فاصله‌ای) و فرمول کودر-ریچاردسون (برای پرسش‌های دو وجهی).

ضریب همبستگی آلفای کرونباخ نوعی اندازه پایایی ثبات درونی^{۲۳} است. پایایی ثبات درونی، همانند تحلیل عاملی روشی برای مشاهده شدت همبستگی بین گویه‌هاست، اما در محاسبه پایایی، در یک زمان تنها با یک‌سری از گویه‌ها مواجهیم و تحلیل تنها بر روی یک سری از گویه‌ها انجام می‌شود. محاسبات پایایی برای شناسایی گویه‌های ضعیفی که ممکن است در تحلیل بعدی حذف شوند، نیز مفید است. در هر صورت گویه‌های که در تحلیل عاملی یک عامل قوی را تشکیل می‌دهند، هنگامی که در مقیاس قرار می‌گیرند معمولاً دارای ضریب آلفای کرونباخ قابل قبول هستند و لذا فراهم کننده شواهدی برای پایایی ثبات درونی هستند (مونرو، ۱۳۸۹: ۳۸۳).

استفاده از ضریب آلفای کرونباخ، معمول‌ترین روش محاسبه پایایی است و از بقیه روش‌ها بیشتر گزارش شده است. قاعده کلی این است که مقدار آلفای کرونباخ یک مقیاس باید حداقل ۰/۷۰ باشد (بریس و همکاران، ۱۳۹۱: ۴۶۶). به بیان دیگر چنانچه مقیاسی دارای میزان آلفای کرونباخ ۰/۷ و بالاتر باشد، می‌توان گفت که آن مقیاس دارای پایایی است. البته برای سازه‌های اکتشافی می‌توان مقدار ۰/۶ را معیار قرار داد (نونالی، ۱۹۷۸).

اما باید توجه کرد که حتی وقتی که یک مقیاس دارای مقدار آلفای کرونباخ بالایی است، باید تک‌تک سوالات مورد بررسی قرار بگیرد. چرا که ممکن است یک مقیاس دارای آلفای کرونباخ بالایی باشد، اما برخی از سوالات آن همبستگی کمی با سایرین داشته باشند و نیازمند اصلاح و حذف باشند (حبیب‌پور و صفری، ۱۳۸۸: ۳۳۶). بر اساس یک قاعده سرانگشتی، مقادیر آلفا را در پنج دسته تقسیم می‌کنند (جدول ۴-۳).

در مجموع می‌توان گفت روش آلفای کرونباخ برای سنجش سازگاری درونی گویه‌های یک مقیاس به کار می‌رود و عمدتاً برای پرسشنامه‌هایی به کار می‌رود که گویه‌ها یا سوالات آن به صورت طیف لیکرت (و نیز فاصله‌ای/نسبی) طراحی شده باشد.

²² Cronbach's Alpha

²³ Internal Consistency Reliability

جدول ۴-۳- مقادیر آلفای کرونباخ و شیوه تفسیر آن

۹. $\geq$ عالی

۸. $\geq$ خوب

۷. $\geq$ قابل قبول

۶. $\geq$ سوال برانگیز و قابل تردید

۵. $\geq$ ضعیف

۵. $\leq$ غیرقابل قبول

☑ نکته: تنها پایایی سوالاتی را می‌توانیم محاسبه کنیم که دارای گزینه‌های همسان و مساوی هستند. به عنوان مثال پایایی سوالات ۵ گزینه‌ای را باید با هم سنجید و پایایی سوالات ۳ گزینه‌ای را با هم.

☑ نکته: گاهی برخی پژوهش‌گران به جای محاسبه پایایی هر مقیاس یا متغیر اقدام به محاسبه پایایی کل پرسشنامه می‌کنند؛ یعنی پایایی پرسشنامه‌ای که حاوی چندین مقیاس یا مفهوم متفاوت است را اندازه می‌گیرند. این امر اشتباه بوده و تنها باید در هر مرحله به محاسبه پایایی گویه‌های یک مقیاس یا متغیر واحد پرداخت. البته می‌توان پایایی کل را برای متغیری که دارای چندین بعد یا مولفه است نیز سنجید، اما برای متغیرهایی که مفهوم متفاوتی دارند نمی‌توان پایایی کل را حساب کرد، مثلاً نمی‌توان پایایی دو متغیر اعتماد اجتماعی و افسردگی را باهم سنجید.

☑ نکته: قبل از محاسبه پایایی باید تمامی سوالات را هم‌جهت کرد (کدگذاری مجدد). یعنی باید سوالات مثبت و منفی (معکوس) را هم‌جهت کرده و جهت کدگذاری تمامی سوالات را یکسان کنیم.

عوامل مؤثر بر میزان پایایی

ضریب آلفای کرونباخ تحت تأثیر عوامل متعددی است که مهم‌ترین آن‌ها عبارتند از:

۱- تعداد سوالات: هر چقدر تعداد سوالات یک مقیاس بیشتر باشد، ضریب آلفای

آن مقیاس بالاتر می‌رود.

- ۲- حجم نمونه: هر چه حجم نمونه بزرگتر باشد، ضریب آلفا بالاتر می‌رود. همچنین، برای حجم نمونه کوچکتر از ۳۰، مقدار آلفا مورد تردید است.
- ۳- سطح سنجش: هر چه متغیرهای مورد سنجش از مقیاس بالاتری (فاصله‌ای/نسبی) برخوردار باشند، میزان آلفای کرونباخ آن بالاتر می‌رود (سطح سنجش از بالا به پایین به ترتیب بدین صورت است: نسبی/فاصله‌ای، ترتیبی و اسمی).

مثال

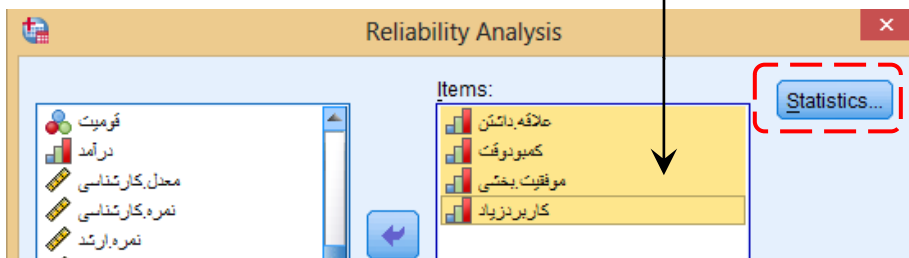
در این اینجا می‌خواهیم میزان پایایی مقیاس "تمایل به آموختن SPSS" را بسنجیم. ما مفهومی به نام تمایل به آموختن برنامه SPSS داریم و قصد داریم با مقیاسی که ساخته‌ایم، میزان تمایل به آموختن SPSS را بسنجیم. قبل از این که اقدام به انجام نمونه‌گیری اصلی کنیم باید از پایایی مقیاس مورد نظر اطمینان حاصل کنیم. در نتیجه پرسشنامه اولیه را بین حدود ۳۰ نفر از پاسخگویان توزیع کرده و بعد از جمع‌آوری پرسشنامه‌ها و وارد کردن اطلاعات در برنامه، اقدام به محاسبه پایایی می‌نماییم.

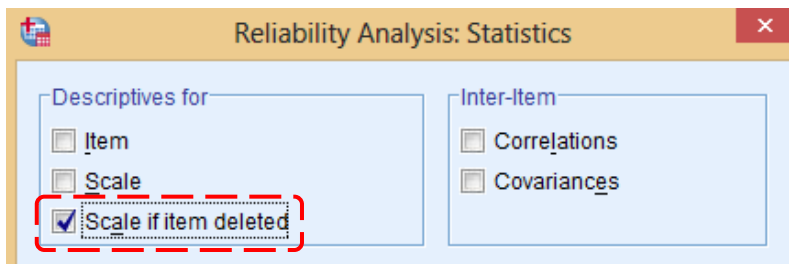
اجرا:

مسیر زیر را دنبال می‌کنیم.

Analyze ---> Scale ---> Reliability Analysis

از کادر سمت چپ سوالات مربوط به مقیاس "تمایل به آموختن SPSS" را انتخاب کرده و وارد کادر گویه‌ها (Items) می‌کنیم. بر روی گزینه Statistics کلیک می‌کنیم.





در پنجره Statistics و در کادر Descriptive for گزینه "پایایی مقیاس اگر گویه حذف شود" (Scale if items deleted) را انتخاب می‌کنیم.

نتایج:

مقدار آلفای کرونباخ در جدول زیر نشان داده شده است. همان‌طور که مشاهده می‌شود مقدار آلفای کرونباخ برابر با ۰/۷۳. به دست آمده است. مقدار آلفای به دست آمده بیشتر از مقدار ۰/۷۰ است و بدین معناست که مقیاس تمایل به آموختن SPSS از پایایی قابل قبولی برخوردار است.

Reliability Statistics

Cronbach's Alpha	N of Items
ضریب آلفای کرونباخ	تعداد سوالات یا گویه‌ها
.731	4

با وجود این که مقیاس مدنظر ما دارای پایایی قابل قبولی است، اما حتما باید به ارزیابی تک‌تک سوالات مقیاس بپردازیم. شاید با وجود پایابودن کل مقیاس، یک یا چند سوال از مقیاس دارای همبستگی ضعیف با کل مقیاس باشند و در نتیجه بتوانیم با اصلاح یا حذف برخی سوالات، پایایی را به اندازه خوب (بالاتر از ۰/۸) یا عالی (بالاتر از ۰/۹) برسانیم.

ستون آخر جدول بعد، به این معناست که اگر ما هر کدام از سوالات را از مقیاس حذف کنیم و دوباره پایایی را محاسبه کنیم، میزان پایایی به‌چه میزان می‌رسد. در واقع مقادیر آلفای ستون آخر نشان می‌دهد که حذف هر سوال چه تأثیری بر مقدار پایایی دارد.

مقدار پایایی کل برابر با ۷۳٪ است که پایایی قابل قبولی است. چنانچه با اصلاح یا حذف سوالی بتوان مقدار آلفا را به بالاتر از ۸٪ یا ۹٪ رساند می‌توان اصلاح یا حذف آن سوال را بررسی کرد. در جدول زیر مشاهده می‌شود که تنها با حذف سوال کمبود وقت می‌توان میزان پایایی را افزایش داد. اگر سوال مدنظر را از مقیاس حذف کنیم، میزان آلفای کرونباخ مقیاس از ۷۳٪ به ۸۱٪ می‌رسد و افزایش قابل توجهی را نشان می‌دهد. در نتیجه با توجه به جوانب دیگر (مهم بودن سوال مورد نظر یا تعداد کل سوالات مقیاس) می‌توان به اصلاح یا حذف سوال مدنظر اقدام کرد.

همبستگی پایین بین هر گویه با کل مقیاس نشان می‌دهد که گویه‌ها، سازه یا مفهوم یکسانی را نشان نمی‌دهند و اندازه نمی‌گیرند. میزان همبستگی کمتر از ۳٪ نشان می‌دهد که گویه مورد نظر همبستگی خوبی با مقیاس کلی ندارد و در نتیجه باید از تحلیل کنار گذاشته شود (هیر و دیگران، ۲۰۱۰). البته مقادیر همبستگی بسیار بالا (بیشتر از ۹۵٪) گمانی است بر این امر که بین گویه‌ها هم‌خطی چندگانه وجود دارد یا این که احتمال دارد پاسخگویان به طور واقعی پاسخ نداده باشند (تاباچینک و فیدل، ۲۰۰۱). در نتیجه میزان همبستگی هر گویه با کل مقیاس باید بین ۳٪ تا ۹۵٪ باشد.

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha if Item Deleted
	میانگین مقیاس با حذف هر گویه	واریانس مقیاس با حذف هر گویه	همبستگی گویه با مقیاس (جزء با کل)	مقدار آلفای کرونباخ با حذف هر گویه
علاقه داشتن	9.37	10.59	.491	.693
کمبود وقت	9.44	11.05	.336	.811
موفقیت بخشی	9.16	8.86	.622	.612
کاربرد زیاد	8.81	10.43	.653	.615

مقدار آلفای منفی

گاهی اوقات پس از اجرای دستور محاسبه آلفای کرونباخ در نرم افزار SPSS، مقدار آلفا برای برخی گویه‌ها منفی به دست می‌آید. به عبارتی، با وجود آن که مقدار آلفا به لحاظ نظری بین (۰) تا (۱) نوسان دارد، اما در عمل مقدار آن ممکن است بین مقادیر $-\infty$ (منفی بی‌نهایت) تا ۱ نوسان داشته باشد. این موضوع موردسوال بسیاری از پژوهش‌گران می‌باشد که در پاسخ باید گفت مقدار منفی آلفا در سه حالت اتفاق می‌افتد:

- ۱- زمانی که جهت کدگذاری گویه‌ها مخالف هم باشد. مثلاً زمانی که برخی گویه‌ها از (۱) تا (۵) و برخی دیگر از (۵) تا (۱) کدگذاری شوند.
- ۲- زمانی که مجموع واریانس گویه‌ها، بزرگتر از واریانس مقیاس باشد.
- ۳- زمانی که میانگین کوواریانس بین گویه‌ها، کوچکتر از صفر و به عبارتی منفی باشد (حبیب‌پور و صفری‌شالی، ۱۳۸۸: ۳۶۶).



تعریف

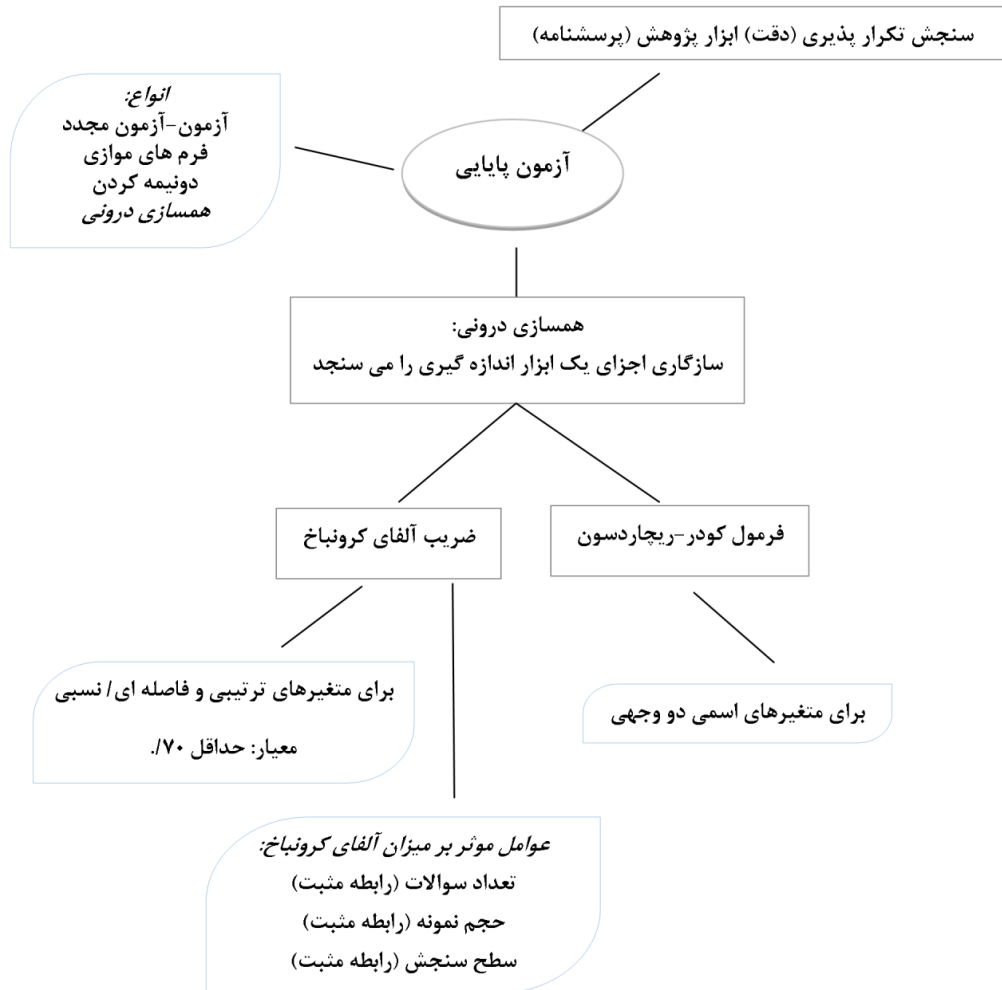
در پژوهشی به بررسی میزان اعتماد اجتماعی در بین شهروندان تهرانی پرداخته شد. مقیاس اعتماد اجتماعی با چهار شاخص سنجیده شد و در آن از پاسخگویان پرسیده شد که «به هرکدام از گروه‌های زیر چقدر اعتماد دارند: اعضای خانواده، بستگان نزدیک، همسایگان و همکاران». شرکت‌کنندگان به سوالات در قالب طیف لیکرت پنج گزینه‌ای از خیلی کم (کد ۱)، کم (کد ۲)، متوسط (کد ۳)، زیاد (کد ۴) و خیلی زیاد (کد ۵) پاسخ دادند، نتایج در جدول بعد نشان داده شده است. اطلاعات را وارد برنامه کنید.

شماره پرسشنامه	جنس	خانواده	تعداد پاسخ	همسایگان	همکاران	شماره پرسشنامه	جنس	خانواده	تعداد پاسخ	همسایگان	همکاران
۱	۱	۴	۴	۲	۴	۱۶	۲	۵	۴	۳	۲
۲	۱	۵	۴	۲	۳	۱۷	۲	۵	۳	۲	۴
۳	۱	۳	۳	۱	۴	۱۸	۲	۴	۴	۳	۳
۴	۱	۵	۵	۵	۴	۱۹	۲	۳	۲	۲	۳
۵	۱	۴	۳	۲	۲	۲۰	۲	۵	۵	۵	۵
۶	۱	۳	۲	۳	۲	۲۱	۲	۱	۲	۲	۱
۷	۱	۵	۵	۴	۳	۲۲	۲	۴	۳	۴	۲
۸	۱	۴	۵	۱	۳	۲۳	۲	۴	۱	۴	۲
۹	۱	۲	۳	۳	۵	۲۴	۲	۵	۴	۳	۴
۱۰	۱	۴	۵	۳	۵	۲۵	۲	۵	۴	۲	۲
۱۱	۱	۲	۱	۱	۲	۲۶	۲	۴	۵	۵	۴
۱۲	۱	۲	۲	۲	۲	۲۷	۲	۵	۳	۳	۳
۱۳	۱	۴	۴	۳	۳	۲۸	۲	۴	۲	۲	۳
۱۴	۱	۳	۲	۳	۵	۲۹	۲	۵	۳	۲	۲
۱۵	۱	۴	۲	۵	۳	۳۰	۲	۴	۴	۳	۳

۱- آیا مقیاس اعتماد اجتماعی پایایی دارد؟ در صورت حذف کدام یک از چهار گویه مقیاس اعتماد اجتماعی، میزان پایایی افزایش می‌یابد؟ آیا حذف سوالی از مقیاس را پیشنهاد می‌دهید.

واژه نامه فصل چهارم

Construct validity	اعتبار سازه
Cronbach's Alpha	آلفای کرونباخ
Reliability	پایایی (قابلیت اعتماد)
Internal Consistency Reliability	پایایی ثبات درونی
Exploratory Factor Analysis (EFA)	تحلیل عاملی اکتشافی



تحلیل داده‌ها (آمار استنباطی)

محتوای فصل

- «بخش اول: ضرایب همبستگی و پیوستگی»
- «بخش دوم: آزمون‌های مقایسه‌ای یا تفاوتی»
- «بخش سوم: آزمون‌های چندمتغیره»
- «بخش چهارم: آزمون‌های ناپارامتریک»

بخش اول:

ضرایب همبستگی و پیوستگی

تحلیل آماری رابطه دو متغیر، اولین گام در بررسی روابط علی بین متغیرهاست. بررسی رابطه علی بین پدیده‌ها از اهداف اصلی علم است. در علوم اجتماعی می‌خواهیم بدانیم علت تفاوت‌های اجتماعی چیست. برای مثال، چرا برخی از مردم نگرش سنتی دارند و برخی دیگر غیر سنتی؟ چرا بازده کار در برخی از جوامع بالاست و در برخی جوامع پایین؟ چرا برخی از خانواده‌ها بزرگ هستند و برخی کوچک؟ چرا در جایی اعتماد اجتماعی بالاست و در جای دیگر پایین؟ چرا مشارکت اجتماعی زنان در جایی گسترده‌است و جای دیگر محدود؟

پایه و اساس تحلیل دو متغیره، بررسی رابطه (پیوستگی) بین دو متغیر است. زمانی دو متغیر را مرتبط یا پیوسته می‌خوانیم که توزیع مقادیر یکی از متغیرها برحسب مقادیر مختلف متغیر دیگر تغییر کند. مثلاً اگر آراء انتخاباتی مردم بر حسب طبقه اجتماعی آن‌ها متفاوت باشد، آراء انتخاباتی و طبقه اجتماعی پیوسته‌اند.

دانشمندان با استدلال نظری فرض می‌کنند که پدیده‌ای متأثر از پدیده‌ای دیگر است. قائل بودن به پیوند و رابطه علی بین پدیده‌ها مبتنی بر دو اصل است: **اصل توأمی و اصل هم‌تغییری**. اصل توأمی به معنای آمدن پدیده‌ای به دنبال پدیده دیگر است و اصل هم‌تغییری به معنای تغییر کردن پدیده‌ای به دنبال تغییر کردن پدیده‌ای دیگر. روش‌های آماری دومتغیره بر اساس داده‌های گردآوری شده، نشان می‌دهد که آیا این توأمی یا هم‌تغییری دو متغیر در واقعیت هم وجود دارد یا نه؟ وجود رابطه آماری بین دو متغیر بیان‌گر توأمی یا هم‌تغییری تجربی دو متغیر است.

البته این امر به‌خودی‌خود بیان‌گر رابطه علی بین دو متغیر نیست. به عبارت دیگر، وجود رابطه آماری بین دو متغیر لزوماً به معنای رابطه علی بین دو متغیر نیست. با این وجود، رابطه آماری می‌تواند یک مشاهده تجربی در تایید فرضیه‌ای باشد که بر اساس استدلال نظری قائل به وجود رابطه علی بین دو متغیر است. وجود رابطه آماری بین دو متغیر بیان‌گر توأمی یا هم‌تغییری دو متغیر در واقعیت تجربی معینی است، ولی علی

بودن یا نبودن این رابطه متکی بر استدلال نظری است (نایی، ۱۳۸۸: ۱۳۶). در هیچ رابطه دو متغیره‌ای نمی‌توان علّیت بین دو متغیر را مفروض گرفت و رابطه دو متغیر را علّی دانست. چرا که ممکن است متغیرهای اندازه‌گیری شده یا اندازه‌گیری نشده دیگری وجود داشته باشند که در نتایج مؤثر باشند (برنز^{۲۴} و دیگران، ۲۰۰۸: ۸۹).

☑ نکته: اصطلاح اندازه‌های پیوستگی برای اشاره به رابطه بین متغیرهایی که حداقل یک طرف متغیرها اسمی است به کار می‌رود. چون در مورد رابطه بین متغیرهایی که حداقل یک طرف متغیرها اسمی است، خطّی یا غیرخطّی بودن رابطه بین متغیرها مطرح نیست و متغیرها جهت ندارند.

روش‌های مختلفی برای سنجش رابطه بین دو متغیر وجود دارد. این روش‌ها را می‌توان برحسب نوع و سطح سنجش متغیرها به دو دسته کلی تقسیم کرد. دسته اول این روش‌ها شامل استفاده از جدول تقاطعی و نمودار پراکندگی و دسته دوم این روش‌ها شامل استفاده از ضرایب همبستگی است. ضریب همبستگی رابطه دو متغیر را به شیوه عددی نشان می‌دهد. معنی‌دار بودن آماری رابطه دو متغیر را فقط از طریق آزمون‌های همبستگی می‌توان مشخص کرد.

ساده‌ترین و گویاترین روش آماری برای تحلیل رابطه دو متغیر اسمی (یا متغیر اسمی با متغیر ترتیبی) استفاده از جدول تقاطعی است و در مورد متغیرهای فاصله‌ای استفاده از نمودار پراکندگی. نمودار پراکندگی در فصل سوم - بخش دوم آموزش داده شده است. جدول تقاطعی در ادامه و به همراه ضریب فی و وی کرامر توضیح داده شده است.

ضریب همبستگی

بهتر است اطلاعات جدول تقاطعی و یا نمودار پراکندگی را در یک عدد خلاصه کنیم. بدین ترتیب می‌توان به برداشتی دقیق و ساده از رابطه دو متغیر رسید. همه این آماره‌های خلاصه شده را ضرایب همبستگی یا اندازه‌های پیوستگی می‌خوانند. اساساً ضریب همبستگی فقط یک شاخص است که توصیف خلاصه‌ای از ویژگی رابطه دو متغیر ارائه می‌دهد. ضرایب همبستگی انواع مختلفی دارد و هر نوع ضریب همبستگی کاربرد معینی دارد.

²⁴ Burns

آماره‌های مناسب برای سطوح سنجش مختلف

ضرایب همبستگی متعددی وجود دارند که برحسب سطح سنجش متغیرها تقسیم‌بندی شده و بکار می‌روند. برخی از مهم‌ترین ضرایب همبستگی در ادامه گزارش شده است.

- متغیرهای اسمی: مناسب‌ترین ضرایب پیوستگی زمانی که هر دو متغیر اسمی هستند، ضرایب فی و وی کرامر و ضریب مجذور کای استقلال هستند.
- متغیرهای ترتیبی: رایج‌ترین آماره در جایی که هر دو متغیر ترتیبی هستند ضریب اسپیرمن است.
- متغیرهای فاصله‌ای/نسبی: زمانی که هر دو متغیر در سطح سنجش فاصله‌ای/نسبی هستند، باید از ضریب همبستگی پیرسون استفاده کنیم.

سایر حالت‌ها:

- زمانی که می‌خواهیم رابطه بین یک متغیر اسمی و یک متغیر فاصله‌ای را بسنجیم از ضریب χ^2 استفاده می‌کنیم. البته برخی پژوهش‌گران ترجیح می‌دهند در این حالت به جای استفاده از آزمون‌های همبستگی، از آزمون t مستقل که یک آزمون مقایسه میانگین به حساب می‌آید استفاده کنند.
- زمانی که می‌خواهیم رابطه بین یک متغیر اسمی و یک متغیر ترتیبی را بسنجیم از مقادیر پیوستگی فی و وی کرامر استفاده می‌کنیم. فی برای جداول 2×2 و وی کرامر برای جداول بزرگتر از 2×2 استفاده می‌شود.
- زمانی که می‌خواهیم رابطه بین یک متغیر ترتیبی با یک متغیر فاصله‌ای را بسنجیم از ضریب همبستگی اسپیرمن استفاده می‌کنیم.

☑ نکته: ضریب همبستگی اسپیرمن یک آزمون ناپارامتریک است، در نتیجه هنگامی که داده‌ها از مفروضات آزمون‌های پارامتریک انحراف دارند، می‌تواند مورد استفاده قرار بگیرد. بدین معنا که چنانچه متغیرهای مدنظر در سطح فاصله‌ای/نسبی باشند اما توزیع متغیرها نرمال نباشد، به جای آزمون پیرسون، از آزمون اسپیرمن استفاده می‌کنیم.

☑ نکته: همواره می‌توان سطوح سنجش بالاتر را به سطوح سنجش پایین‌تر تقلیل داد (اما عکس این عمل امکان‌پذیر نیست). در نتیجه سطح سنجش متغیر پایین‌تر تعیین‌کننده نوع

آزمون همبستگی است. مثلاً اگر یک متغیر فاصله‌ای است و دیگری ترتیبی، هر دو را ترتیبی می‌گیریم و از ضریب همبستگی مناسب استفاده می‌کنیم.

خصوصیات ضرایب همبستگی

چنانچه مشخص شد بین دو متغیر رابطه وجود دارد باید ویژگی‌های رابطه یعنی نوع، جهت و شدت آن را مشخص کرد:

۱- جهت: وقتی متغیرها ترتیبی یا فاصله‌ای هستند باید جهت رابطه را مشخص کرد: رابطه یا مثبت است یا منفی. رابطه‌ای مثبت است که در آن احتمال این که افرادی که در متغیری نمره بالایی دارند در متغیر دیگر هم نمره بالا داشته باشند بیش از دیگران است و احتمال این که افرادی که در متغیری نمره پایینی دارند در متغیر دیگر هم نمره پایینی داشته باشند بیش از دیگران است. در رابطه مثبت بین «جایگاه افراد» در دو متغیر همسازی و هم‌خوانی وجود دارد. به طور ساده، رابطه مثبت رابطه‌ای است که تغییرات دو متغیر در یک راستا است: با افزایش مقدار متغیر A ، مقدار متغیر B هم افزایش می‌یابد و با کاهش مقدار متغیر A ، مقدار متغیر B کاهش می‌یابد.

۲- نوع: رابطه متغیرهای ترتیبی و فاصله‌ای یا خطی است یا منحنی شکل. رابطه خطی رابطه‌ای است به شکل «خط راست». گاه ممکن است در جداول پرتبقة، برخی از درصدها اندکی از جهت مستقیم منحرف شده باشند اما اگر روند اصلی در یک جهت باشد به احتمال زیاد رابطه دو متغیر خطی است. آماره‌هایی که در ادامه مطرح می‌شوند (ضریب همبستگی پیرسون و اسپیرمن) برای سنجش رابطه خطی مناسب هستند.

۳- شدت: اگر تفاوت زیادی بین خرده‌گروه‌ها وجود داشته باشد، رابطه محکمی بین دو متغیر وجود دارد. یعنی اگر خرده‌گروه‌ها بر حسب خصوصیات‌شان در متغیر وابسته تفاوت زیادی با هم داشته باشند، رابطه محکمی بین دو متغیر وجود دارد و اگر تفاوت اندک باشد رابطه ضعیف است. شدت رابطه نشان‌دهنده میزان و بزرگی رابطه بین دو متغیر است.

ملاحظات ضرایب همبستگی

- ۱- مقدار ضریب همبستگی همواره بین ۰ تا ۱ است. هر چه مقدار آن بالاتر باشد رابطه قوی تر است. صفر به معنای عدم رابطه و یک به معنای رابطه کامل است.
- ۲- ممکن است ضرایب های همبستگی داده های ترتیبی و فاصله ای دارای علامت منفی باشند. در این صورت رابطه منفی است (با افزایش مقادیر یک متغیر، مقادیر متغیر دیگر کاهش می یابد). فقدان علامت به معنای رابطه مثبت است.
- ۳- برخی ضرایب های همبستگی (مانند پیرسون و اسپیرمن) فقط برای رابطه خطی مناسب اند.
- ۴- این که کدام متغیر، متغیر مستقل باشد، بر گروهی از ضرایب های همبستگی اثر دارد. این گونه ضرایب همبستگی، اندازه های نامتقارن^{۲۶} خوانده می شوند (مانند ضرایب فی و وی کرامر). آن دسته از ضرایب های همبستگی که از این که کدام متغیر، متغیر مستقل باشد تأثیر نمی گیرند، اندازه های متقارن خوانده می شوند (مانند پیرسون و اسپیرمن).
- ۵- وجود رابطه بین دو متغیر تایید کننده رابطه علی بین آنها نیست.

²⁶ Asymmetric measures

آزمون فی و وی کرامر

زمانی که متغیرها در سطح سنجش کیفی (اسمی و ترتیبی) باشند برای بررسی رابطه آن‌ها باید از آزمون‌هایی مانند فی و وی کرامر استفاده کرد. به طور مشخص زمانی که یک متغیر اسمی و دیگری ترتیبی و یا زمانی که هر دو متغیر اسمی باشند، سنجش رابطه (پیوستگی) دو متغیر با استفاده از ضرایب فی و وی کرامر انجام می‌شود. فقط زمانی که دو متغیر هر دو دارای دو طبقه باشند (اصطلاحاً 2×2 باشند)، از ضریب فی کرامر استفاده می‌کنیم و در سایر حالت‌ها از ضریب وی کرامر استفاده می‌کنیم. مثلاً اگر بخواهیم رابطه بین جنس (زن یا مرد) و وضعیت اشتغال (دو طبقه: شاغل یا غیرشاغل) را بسنجیم چون هر دو متغیر دارای دو گروه یا طبقه هستند، از آزمون فی کرامر استفاده می‌کنیم.

مثال

رابطه بین جنس و علاقه‌مندی به درس آمار را با آزمون فی کرامر می‌سنجیم. هر دو متغیر در سطح سنجش اسمی هستند. جنس پاسخگویان شامل دو طبقه زن و مرد است و علاقه‌مندی به درس آمار شامل دو گزینه بلی و خیر است.

بدین جهت که هر دو متغیر اسمی هستند از آزمون فی و وی کرامر^{۲۷} استفاده می‌کنیم. هر دو متغیر دارای دو طبقه هستند از ضریب فی کرامر استفاده کنیم (در حالت 2×2 از ضریب فی کرامر و در سایر حالت‌ها از ضریب وی کرامر استفاده می‌کنیم). مراحل آزمون پیوستگی (یا رابطه) بین جنس پاسخگویان و علاقه‌مندی آنان به درس آمار با استفاده از آزمون فی کرامر در ادامه آمده است.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Descriptive Statistics ---> Crosstabs

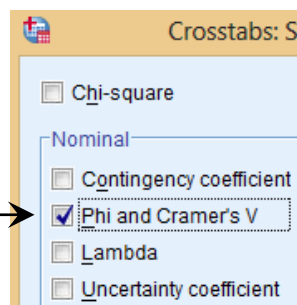
²⁷ Phi and Cramer's V

- ۱- در کادر Row، متغیر وابسته (علاقه مندی به آمار) را وارد می کنیم.
در کادر Column، متغیر مستقل (جنس پاسخگویان) را وارد می کنیم.

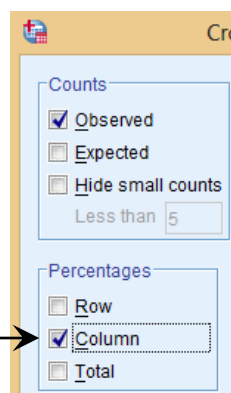


- ۲- گزینه های Statistics و Cells را انتخاب می کنیم و مطابق شکل های بعد دستورات را فعال می کنیم.

- در کادر Nominal بر روی گزینه Phi and Cramer's V کلیک می کنیم.



- با انتخاب گزینه Cells این پنجره ایجاد می شود. در این پنجره گزینه ستون (Column) را انتخاب می کنیم.
انتخاب این گزینه موجب می شود که در جدول تقاطعی ارائه شده در خروجی برنامه، به جای مقادیر فراوانی از درصد فراوانی طبقات استفاده شود و فهم و تفسیر نتایج آسان تر شود.
در نهایت OK Continue را انتخاب می کنیم.



نتایج:

با توجه به جدول بعد می توان گفت که سطح معنی داری به دست آمده (Approx. Sig) برابر با ۰/۱۶، به دست آمده است که از مقدار سطح معنی داری مفروض (۰/۰۵) کمتر

است ($P < .05$). در نتیجه می‌توان گفت که از جنبه آماری بین دو متغیر جنس و علاقه به آمار پیوند (رابطه) وجود دارد. شدت این رابطه برابر با $-.24$ - به دست آمده است.

جهت رابطه به دست آمده منفی است اما چون متغیرها اسمی هستند، تفسیر جهت رابطه باید با توجه به جدول تقاطعی انجام شود.

Symmetric Measures

	Value	Approx. Sig.
	مقدار یا شدت ضریب	سطح معنی‌داری تقریبی
Phi	-.242	.016
Nominal by Nominal Cramer's V	.242	.016
N of Valid Cases (تعداد مورد ها)	100	

با استفاده از جدول تقاطعی می‌توان به ماهیت و چگونگی پیوند (رابطه) دو متغیر پی برد. همان‌طور که درصد فراوانی‌ها نشان می‌دهد، ۶۸ درصد زنان بیان کرده‌اند که به درس آمار علاقه‌مند هستند ولی ۴۴ درصد مردان به درس آمار علاقه دارند. در نتیجه می‌توان گفت بین جنس و علاقه به درس آمار رابطه وجود دارد. نتایج نشان می‌دهد که علاقه‌مندی زنان به درس آمار بیشتر از مردان است.

Crosstabulation جنس * علاقه به آمار

	جنس		Total
	زن	مرد	
Count تعداد خیر	16	28	44
% within Jans علاقه به آمار	32%	56%	44%
Count بله	34	22	56
% within Jans	68%	44%	56%
Count Total	50	50	100
% within Jans	100%	100%	100%

☑ نکته: وقتی نتایج را گزارش می‌کنیم، معمول است مقادیر اعشاری آمار توصیفی (مانند درصد فراوانی، میانگین و انحراف استاندارد) را گرد کنیم، اما وقتی نتایج آمار استنباطی (مانند ضرایب همبستگی، سطح معنی‌داری) را گزارش می‌کنیم باید تمام ارقام اعشاری را گزارش کنیم و از گرد کردن خودداری کنیم.

گزارش:

در گزارش نتایج می‌نویسیم:

بین جنس پاسخگویان و علاقه به درس آمار رابطه معنی‌داری وجود دارد ($P < .05$) و $n=100$ و $r = -.242$ = ضریب فی کرامر). شدت رابطه ضعیف است. جدول تقاطعی نشان داد که درصد زنان علاقه‌مند به درس آمار (۶۸ درصد) بیشتر از مردان (۴۴ درصد) است که نشان از علاقه بیشتر زنان نسبت به درس آمار در مقایسه با مردان دارد.

در تکمیل گزارش نتایج، جداول و نتایجی که در خروجی برنامه به دست می‌آید را در قالب جداول مناسب گزارش کنید.

آزمون مجذور کای استقلال

در این بخش به معرفی آزمون مجذور کای استقلال^{۲۸} می‌پردازیم. آزمون مجذور کای استقلال (χ^2) یک آزمون معتبر آماری است که به وسیله آن می‌توان پی برد که آیا بین دو متغیر رابطه معنی‌دار وجود دارد یا خیر. آزمون مجذور کای معمولاً برای رابطه‌هایی به کار می‌رود که هر دو متغیر ناپارامتری باشند (اسمی با اسمی و یا اسمی با ترتیبی).

هرگاه در یک نمونه مورد مطالعه، هیچ رابطه منظمی بین دو متغیر وجود نداشته باشد، می‌توان نتیجه گرفت که دو متغیر از یکدیگر مستقل هستند که اصطلاحاً به آن استقلال آماری می‌گویند.

اساس آزمون مجذور کای، بررسی فراوانی‌های مشاهده شده با فراوانی‌های مورد انتظار است. یعنی می‌خواهیم بدانیم که آیا بین فراوانی‌های مشاهده شده و فراوانی‌های مورد انتظار تفاوت وجود دارد یا خیر. به عبارتی، در این آزمون، با استفاده از تفاوت بین این دو فراوانی، به وجود یا عدم وجود رابطه بین دو متغیر پی می‌بریم. بین مقدار مجذور کای و معنی‌داری رابطه بین دو متغیر، ارتباط مستقیمی وجود دارد. یعنی هرچه مقدار مجذور کای بزرگتر باشد، احتمال وجود رابطه بین دو متغیر بیشتر است.

به طور کلی، آزمون مجذور کای فقط به تشخیص این امر کمک می‌کند که آیا متغیرها مستقل از یکدیگرند یا با هم رابطه دارند، اما هیچ‌گاه چگونگی و شدت رابطه را توضیح نمی‌دهد. بنابراین، پس از محاسبه مقدار مجذور کای، در صورت وجود رابطه بین متغیرها، باید با استفاده از شاخص‌های پیوند (مانند فی و وی کرامر)، جهت و شدت رابطه را تعیین کنیم (حبیب پور و صفری، ۱۳۸۸: ۴۲۱).

پیش‌فرض‌ها

از آزمون مجذور کای استقلال زمانی می‌توانیم استفاده کنیم که:

۱- حجم نمونه آماری بیشتر از ۴۰ باشد.

²⁸ Independence χ^2

۲- درجه آزادی معادل ۱ نباشد. اگر یک بود، فراوانی مورد انتظار در هر یک از طبقات متغیرها نباید کمتر از ۱۰ باشد (ساعی، ۱۳۸۸: ۹۵).

مثال

آزمون مجذور کای برای بررسی وجود رابطه بین دو متغیر کیفی (اسمی/ترتیبی) استفاده می‌شود. مثلاً زمانی که بخواهیم بررسی کنیم که آیا بین جنس افراد (زن یا مرد) با علاقه به کشیدن سیگار (بلی یا خیر) و یا بین نوع دین افراد (اسلام، مسیحیت و یهودیت) و تمایل به اهداء خون (کم، متوسط، زیاد) رابطه‌ای وجود دارد یا خیر، از آزمون مجذور کای استفاده کنیم.

در مثال کتاب، می‌خواهیم بررسی کنیم که آیا بین جنس دانشجویان با علاقه به درس آمار رابطه وجود دارد یا خیر. جنس شامل دو طبقه زن و مرد و علاقه به درس آمار شامل دو طبقه خیر و بلی می‌شود. در نتیجه فرضیه پژوهش به صورت زیر است:

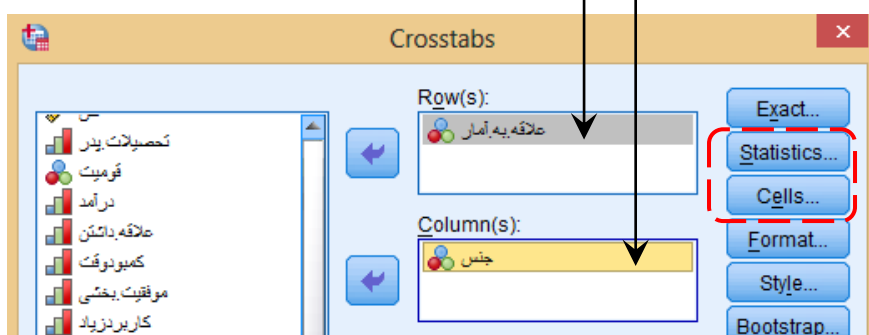
فرضیه: بین جنس دانشجویان و علاقه به درس آمار رابطه وجود دارد.

اجرا:

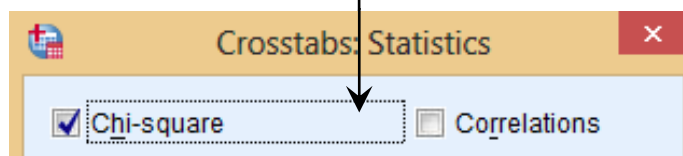
مسیر زیر را دنبال می‌کنیم:

Analyze ---> Descriptive Statistics ---> Crosstabs

متغیر جنس را وارد کادر Column می‌کنیم. در این کادر باید متغیر مستقل یا متغیری که می‌خواهیم بر اساس آن متغیر دیگر را گروه‌بندی کنیم وارد کنیم. متغیر وابسته (علاقه به آمار) را وارد کادر Row می‌کنیم. مانند آزمون فی و وی کرامر بر روی گزینه Cells کلیک می‌کنیم و در پنجره باز شده و در کادر Percentage گزینه Column را انتخاب می‌کنیم. با این انتخاب، در جدول تقاطعی که برنامه ارائه می‌دهد، درصد فراوانی هر کدام از طبقات گزارش می‌شود. در انتها گزینه Statistics را انتخاب می‌کنیم.



در پنجره Statistics آزمون خی دو (Chi-Square) را فعال می‌کنیم. سپس گزینه Continue و سپس OK را انتخاب می‌کنیم.



نتایج:

بعد از اجرای دستور مجذور کای، در خروجی برنامه ابتدا جدول تقاطعی ارائه می‌شود. بررسی جدول تقاطعی می‌تواند به طور تقریبی ما را از وجود رابطه بین متغیرها و شدت آن آگاه کند. جدول تقاطعی بیانگر فراوانی و درصد فراوانی هر کدام از گزینه‌های علاقه به آمار و عدم علاقه به آمار در بین دانشجویان دختر و پسر است که به جهت آسانی، برای مقایسه گروه‌ها از درصد فراوانی به جای فراوانی استفاده می‌کنیم.

مقایسه درصد فراوانی دختران و پسران نشان می دهد که ۳۲ درصد از دختران بیان کرده اند که علاقه ای به درس آمار ندارند، این مقدار در بین پسران ۵۶ درصد است. به طور معکوس، ۶۸ درصد دختران به آمار علاقه مند هستند اما در مقام مقایسه ۴۴ درصد پسران به آمار علاقه دارند. بررسی درصد فراوانی ها نشان می دهد تعداد و نسبت بیشتری از دختران در مقایسه با پسران به درس آمار علاقه دارند (داده ها واقعی نیست).

Crosstabulation جنس * علاقه به آمار

	جنس		Total
	زن	مرد	
Count	16	28	44
% within Jans	32%	56%	44%
Count	34	22	56
% within Jans	68%	44%	56%

جهت بررسی معنی دار بودن رابطه به دست آمده به نتایج آزمون مجذور کای مراجعه می کنیم. درجه آزادی (df) به دست آمده برابر با ۱ است اما چون فراوانی تمام طبقات جدول تقاطعی بیشتر از ۱۰ است در نتیجه می توانیم نتایج آزمون مجذور کای را گزارش کنیم. عموماً نتایج ردیف اول (Pearson Chi-Square) گزارش می شود اما اگر درجه آزادی جدول برابر با یک باشد نتایج ردیف دوم (Continuity Correction) گزارش می شود. چون درجه آزادی (df) برابر با ۱ است از نتایج ردیف دوم بهره می گیریم.

همان طور که مشاهده می شود مقدار مجذور کای برابر با ۴.۹۱۱ به دست آمده است که این مقدار مجذور کای در سطح معنی داری کمتر از ۰.۰۵ قرار دارد. به عبارتی چون مجذور کای به دست آمده در سطح معنی داری کمتر از ۰.۰۵ ($P < 0.05$) معنی دار است، می توان گفت که بین جنس دانشجویان و علاقه به درس آمار رابطه معنی دار وجود دارد.

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	5.84^a	1	.016
Continuity Correction^b	4.91	1	.027
Likelihood Ratio	5.91	1	.015
Linear-by-Linear Association	5.79	1	.016
N of Valid Cases	100		

☑ نکته: این آزمون را در عینا با آزمون فی کرامر هم انجام دادیم و همین نتایج حاصل شد (مقدار معنی داری برابر است). اما تفاوت آزمون مجذور کای و وی کرامر در این است که آزمون مجذور کای در مورد شدت و جهت رابطه اطلاعاتی به ما نمی دهد ولی آزمون فی کرامر نتایج کاملی به ما می دهد که شامل معنی داری، و جهت و شدت رابطه است.

گزارش:

در گزارش نتایج می نویسیم:

آزمون مجذور کای استقلال نشان داد که بین جنس دانشجویان و علاقه به درس آمار رابطه وجود دارد ($P < .05$ و $df=1$ و $n=100$ و $4.91 =$ مقدار مجذور کای یا χ^2). جدول تقاطعی نشان داد که درصد زنان علاقه مند به درس آمار (۶۸ درصد) بیشتر از مردان (۴۴ درصد) است که بدین معناست که زنان در مقایسه با مردان، علاقه بیشتری به درس آمار دارند.

آزمون همبستگی اسپیرمن

اگر یک (یا هر دو) متغیر موجود در طرح پژوهش غیرپارامتری باشند نمی‌توان از آزمون‌های پارامتری مانند پیرسون برای تعیین همبستگی استفاده کرد. در چنین موقعیت‌هایی باید اندازه‌گیری غیرپارامتری همبستگی را به کار برد. در چنین مواقعی می‌توان از ضریب رو اسپیرمن استفاده کرد (بریس، کمپوسلنگار، ۱۳۹۱: ۱۶۷). آزمون همبستگی اسپیرمن (رو اسپیرمن) از آزمون‌های پرکاربرد است که برای سنجش رابطه بین دو متغیر ترتیبی، یا بین یک متغیر ترتیبی با فاصله‌ای/نسبی و یا بین دو متغیر فاصله‌ای/نسبی که غیرنرمال هستند استفاده می‌شود.

مثال

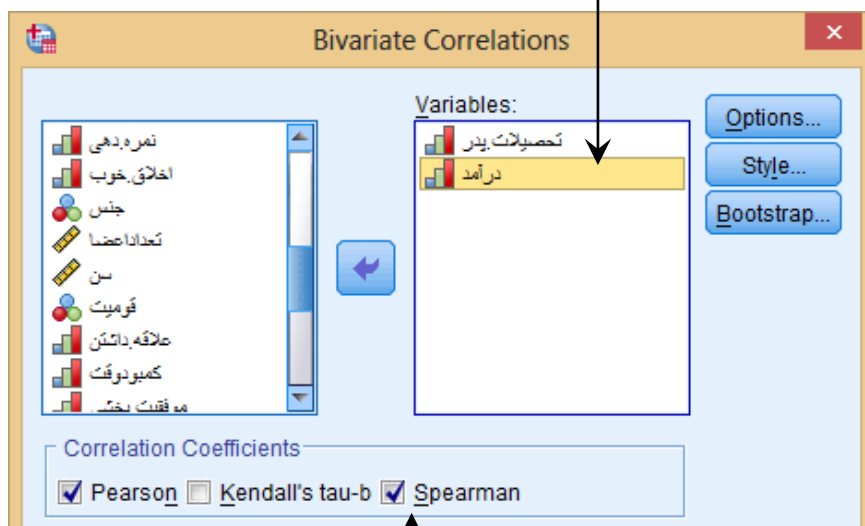
قصد داریم رابطه بین میزان تحصیلات پدر و میزان درآمد خانواده را بررسی کنیم. متغیرهای میزان تحصیلات و میزان درآمد در سطح سنجش ترتیبی قرار دارند در نتیجه جهت بررسی رابطه بین دو متغیر از ضریب همبستگی اسپیرمن^{۲۹} استفاده می‌شود (زمانی که سطح سنجش یک متغیر ترتیبی و متغیر دیگر فاصله‌ای/نسبی باشد، باز هم از آزمون همبستگی اسپیرمن استفاده می‌کنیم). متغیر تحصیلات پدر شامل پنج طبقه می‌شود: سیکل و پایین‌تر، دیپلم، فوق دیپلم، لیسانس و طبقه فوق لیسانس و بالاتر می‌شود. متغیر درآمد خانواده شامل چهار طبقه می‌شود: کمتر از ۵۰۰ هزار تومان، ۵۰۰ هزار یا ۱ میلیون، ۱ تا ۱.۵ میلیون و طبقه بیشتر از ۱.۵ میلیون.

دستور آزمون همبستگی اسپیرمن (و پیرسون) علاوه بر بخش Crosstabs در بخش دیگری از برنامه هم وجود دارد. به جهت آسان‌تر بودن مسیر دوم، دستور آزمون اسپیرمن (و نیز پیرسون) را از مسیر زیر دنبال می‌کنیم.

Analyze ---> Correlate ---> Bivariate

²⁹ Spearman's rho

در کادر متغیرها (Variables) دو متغیر میزان تحصیلات و درآمد افراد را وارد می‌کنیم.



در بخش ضرایب همبستگی (Correlation Coefficients) آزمون مدنظر (آزمون اسپیرمن) را فعال می‌کنیم. در انتها گزینه OK را انتخاب می‌کنیم.

نتایج:

نتایج آزمون همبستگی اسپیرمن در جدول زیر ارائه شده است. نخست به معنی‌داری رابطه و سپس به میزان و جهت رابطه دقت می‌کنیم. سطح معنی‌داری به دست آمده یا Sig. (2-tailed) برابر با ۰/۰۰۲ است. به دست آمده است و کمتر از سطح معنی‌داری ۰/۰۵ است. در نتیجه می‌توان گفت که بین دو متغیر رابطه معنی‌دار وجود دارد. میزان یا شدت رابطه برابر با ۰/۳۱- و جهت رابطه منفی است. جهت منفی نشان می‌دهد که با افزایش میزان تحصیلات پاسخگویان، میزان درآمد آنان کاهش می‌یابد و بالعکس! (مقادیر دو ردیف ارائه شده در جدول همبستگی پیرسون و اسپیرمن یکسان است و بررسی نتایج یک ردیف کافی است).

☑ نکته: علامت دو ستاره (**) در بالای ضریب همبستگی نشان می‌دهد که همبستگی مشاهده‌شده در سطح معنی‌داری کمتر از ۰/۰۱ معنی‌دار است. یعنی سطح معنی‌داری به دست

آمده کمتر از ۰.۱/ شده است ($P < ۰.۱$). چنانچه یک ستاره وجود داشته باشد نشان می‌دهد که رابطه در سطح معنی‌داری کمتر از ۰.۵/ قرار دارد ($P < ۰.۵$).

Correlations

		تحصیلات	درآمد
تحصیلات	Correlation Coefficient	1	-.310**
	Sig. (2-tailed)	.	.002
	N	100	100
	Spearman's rho		
درآمد	ضریب همبستگی	1	-.310**
	Sig. (2-tailed)	.	.002
	سطح معنی‌داری در وضعیت دو دامنه		
	N	100	100

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

علامت دو ستاره، وجود همبستگی را در سطح معنی‌داری ۰.۱/ نشان می‌دهد

گزارش:

در گزارش نتایج می‌نویسیم:

جهت بررسی رابطه میزان تحصیلات پدر و میزان درآمد خانواده از ضریب همبستگی اسپیرمن استفاده شد. یافته‌ها نشان می‌دهد بین میزان تحصیلات و میزان درآمد همبستگی منفی وجود دارد ($P < ۰.۱$) و $n = ۱۰۰$ و $r_s = -.۳۱۰$. جهت منفی رابطه نشان می‌دهد که با افزایش میزان تحصیلات، میزان درآمد کاهش می‌یابد.

آزمون همبستگی پیرسون

احتمالا، گسترده‌ترین کاربرد شاخص آماری همبستگی دو متغیری، ضریب همبستگی گشتاوری پیرسون است که به‌طور معمول همبستگی پیرسون نامیده می‌شود. علامت اختصاری آن r است. ضریب پیرسون نشان می‌دهد که تا چه اندازه بین متغیرهای کمی رابطه خطی وجود دارد (میزر، گامست و گارینو، ۱۳۹۱: ۱۵۲).

کاربرد اصلی ضریب پیرسون زمانی است که متغیرها از نوع پارامتری باشند؛ بدین معنا که توزیع نرمال داشته باشند و در سطح فاصله‌ای/نسبی باشند. البته زمانی که متغیرها از نوع شبه فاصله‌ای باشند (یعنی هر متغیر ترکیبی از چند متغیر ترتیبی باشد که اصطلاحاً به آن مقیاس‌های تراکمی می‌گویند)، برخی از پژوهش‌گران از ضریب پیرسون استفاده می‌کنند. برخی از نویسندگان استفاده از ضریب پیرسون برای یک متغیر دو ارزشی و یک متغیر فاصله‌ای/نسبی را هم مجاز شمرده‌اند. تفسیر همبستگی پیرسون زمانی که یکی از متغیرها دوازشی (فقط شامل دو سطح) اما متغیر دیگر کمی است نیز می‌تواند منطقی باشد (میزر، گامست و گارینو، ۱۳۹۱: ۱۶۴).

تفسیر شدت رابطه در همبستگی پیرسون

بعد از تعیین معنی‌داری و جهت رابطه، باید شدت رابطه ارزیابی شود. برای تفسیر شدت رابطه دومتغیر، تقسیم‌بندی‌های گوناگونی ارائه شده‌است. تقسیم‌بندی زیر یکی آن‌هاست.

جدول ۵-۱- شیوه تفسیر شدت رابطه در همبستگی پیرسون

شدت رابطه	تفسیر
۰/۸ تا ۱	رابطه بسیار قوی
۰/۶ تا ۰/۸	رابطه قوی
۰/۴ تا ۰/۶	رابطه متوسط
۰/۲ تا ۰/۴	رابطه کم (یا ضعیف)
صفر تا ۰/۲	فقدان رابطه یا رابطه ناچیز

(منبع: میلر، ۱۳۸۰: ۲۹۹)

مثال

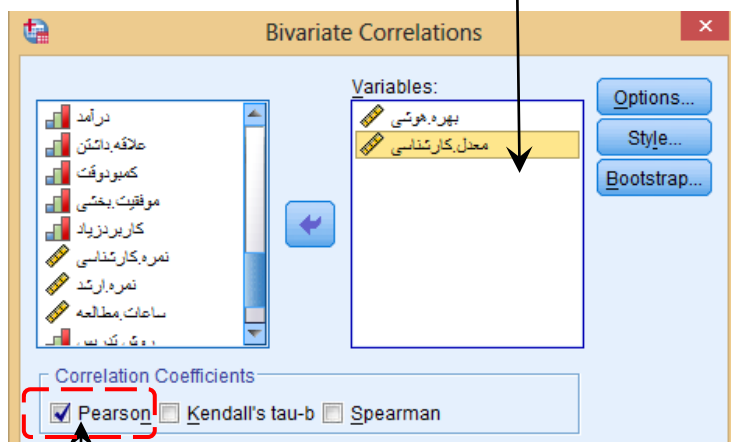
در این بخش به بررسی همبستگی بین دو متغیر بهره‌های هوشی و معدل مقطع کارشناسی می‌پردازیم. انتظار داریم که دو متغیر با یکدیگر همبسته باشند، به نحوی که با افزایش بهره‌های هوشی، معدل افزایش بیابد. به عبارت دیگر انتظار داریم افرادی که بهره‌های هوشی بالاتری دارند، معدل بالاتری هم داشته باشند. هر دو متغیر کمی بوده و در سطح سنجش فاصله‌ای/نسبی قرار دارند. با توجه به این که هر دو متغیر در سطح سنجش فاصله‌ای/نسبی هستند از همبستگی پیرسون^{۳۰} استفاده می‌کنیم. از پیش‌فرض‌های آزمون همبستگی پیرسون نرمال بودن توزیع متغیر در جمعیت آماری است، در این مثال فرض می‌کنیم که این پیش‌فرض برقرار است و توزیع داده‌ها نرمال است.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Correlate ---> Bivariate

دو متغیر بهره‌های هوشی و معدل کارشناسی را از کادر سمت چپ وارد کادر متغیرها (Variables) می‌کنیم.



در بخش ضرایب همبستگی (Correlation Coefficients) آزمون پیرسون را فعال می‌کنیم (این آزمون به طور پیش‌فرض فعال است). در انتها گزینه OK را انتخاب می‌کنیم.

³⁰ Pearson Correlation

نتایج:

در جدول بعد، نتایج آزمون همبستگی پیرسون بین دو متغیر بهره هوشی و معدل کارشناسی نشان داده شده است. نخست به سطح معنی داری به دست آمده نگاه می کنیم. سطح معنی داری به دست آمده برابر با 0.060 است. به دست آمده است که بسیار بیشتر از مقدار مفروض 0.05 است. در نتیجه بین دو متغیر بهره هوشی و معدل مقطع کارشناسی پاسخگویان رابطه معنی داری وجود ندارد. با توجه به این که بین دو متغیر همبستگی وجود ندارد، شدت و جهت رابطه مورد بررسی قرار نمی گیرد.

Correlations

		بهره هوشی	معدل کارشناسی
بهره هوشی	Pearson Correlation	1	-.053
	Sig. (2-tailed)		.600
	N	100	100
معدل کارشناسی	Pearson Correlation	-.053	1
	همبستگی پیرسون Sig. (2-tailed)	.600	
	N	100	100

گزارش:

در گزارش نتایج می نویسیم:

از آزمون همبستگی پیرسون جهت آزمون رابطه دو متغیر بهره هوشی و معدل مقطع کارشناسی استفاده شد. بین میزان بهره هوشی و معدل مقطع کارشناسی همبستگی معنی دار مشاهده نشد ($P = 0.060$ و $n = 100$ و $r = -0.053$). در نتیجه از جنبه آماری دو متغیر بهره هوشی و معدل کارشناسی با یکدیگر رابطه ندارند.

(فرض می کنیم رابطه به دست آمده معنی دار باشد و سطح معنی داری به دست آمده برابر با 0.05 شده است و ضریب پیرسون برابر با 0.45 به دست آمده است. در این صورت به این صورت گزارش می دهیم:

آزمون همبستگی پیرسون نشان داد که بین میزان بهره‌هوشی و معدل مقطع کارشناسی همبستگی وجود دارد ($P = .004$ و $n = 100$ و $r = .45$). جهت رابطه بین بهره‌هوشی و معدل کارشناسی مثبت است. شدت همبستگی به‌دست آمده در حد متوسط است. واریانس توضیح داده شده ۲۰.۳٪ است. نتایج نشان می‌دهد دانشجویانی که بهره‌هوشی بالاتری دارند، معدل کارشناسی بالاتری هم کسب کرده‌اند.

چند نکته:

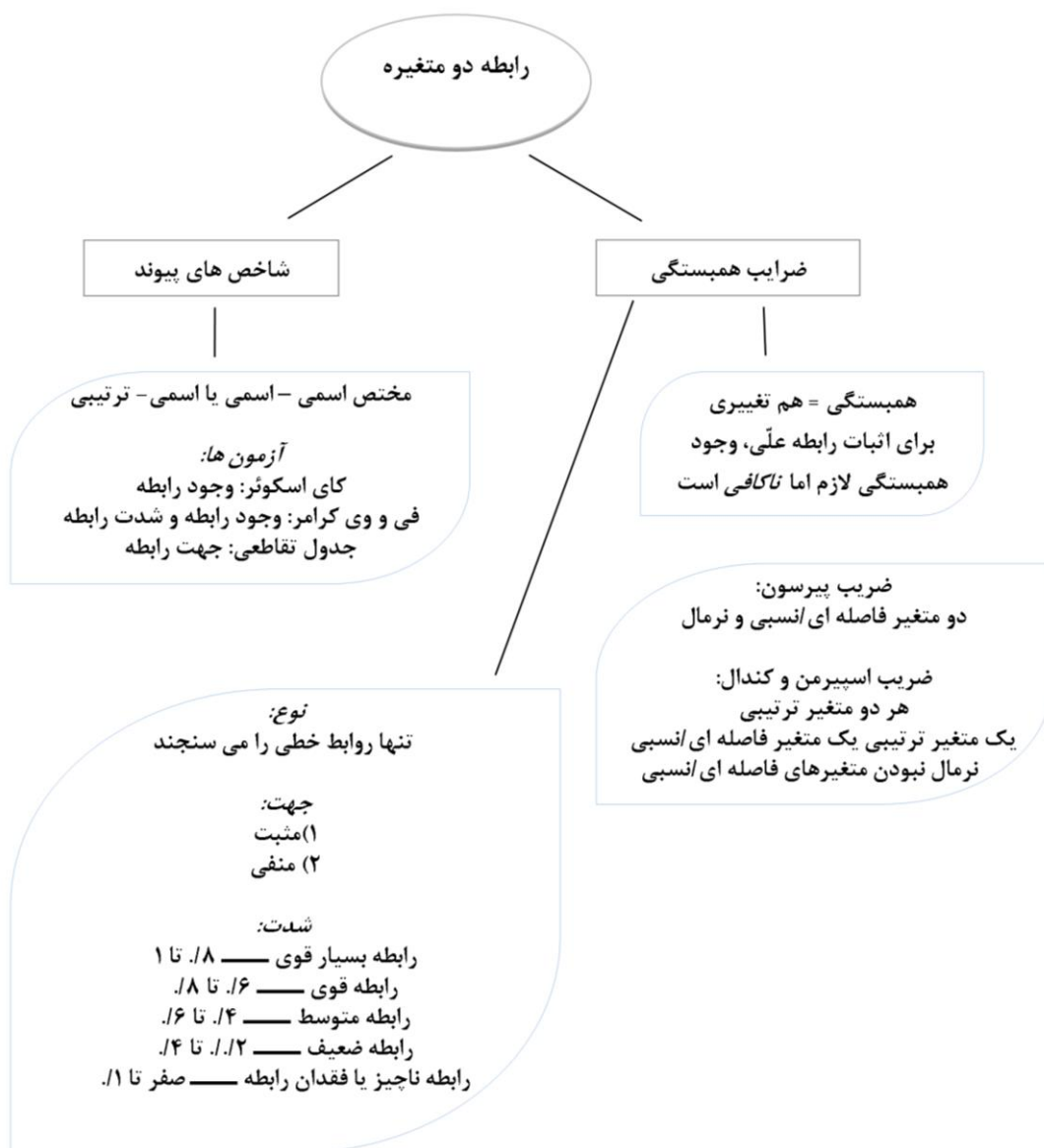
ضریب همبستگی پیرسون را با r نشان می‌دهند. واریانس توضیح داده‌شده همان r^2 است که از به‌توان‌دو رساندن ضریب همبستگی پیرسون (r) به دست می‌آید و نشان‌دهنده واریانس مشترک دو متغیر است. بهتر است نمودار پراکندگی دو متغیر بهره‌هوشی و معدل در گزارش ذکر شود (قبل از نتایج آزمون همبستگی پیرسون).



تعریف

به جدول صفحه ۹۶ رجوع کنید:

- ۱- آیا بین جنس افراد و میزان تحصیلات آنان رابطه وجود دارد؟ نوع و شدت رابطه را به دست بیاورید (ضرایب مجذور کای و وی کرامر).
- ۲- آیا بین میزان تحصیلات و میزان ورزش در افراد همبستگی وجود دارد؟ آیا افزایش تحصیلات افراد با کاهش میزان ورزش در آنان همراه است؟
- ۳- رابطه بین میزان ورزش و وزن افراد را بررسی کنید. جهت و شدت رابطه را مشخص کنید.



آزمون همبستگی تفکیکی (جزئی)

اصولا وقتی همبستگی بین دو متغیر محاسبه می‌شود، تفسیر آن بحث برانگیز و ابهام‌آمیز است. پرسش‌هایی در این زمینه مطرح می‌شود مانند این که همبستگی به دست آمده چه چیز را توجیه می‌کند؟ آیا رابطه بین متغیرها مستقیم یا ناشی از نفوذ متغیر دیگری است؟ تحلیل همبستگی جزئی تلاش می‌کند برخی از ابهامات موجود در تفسیر همبستگی را حل کند.

به منظور آن که بتوانیم رابطه بین زوج معینی از متغیرها را ارزشیابی کنیم، گاه ضرورت دارد اثر متغیر سومی را حذف کنیم. برای مثال، انتظار می‌رود بین گنجینه واژگان و اندازه کفش، در میان کودکان خردسال همبستگی مثبت قابل ملاحظه‌ای وجود داشته باشد. اما تردید وجود دارد که بین این دو متغیر واقعا رابطه علت و معلولی وجود داشته باشد، زیرا پشت این دو متغیر، متغیر سوم یعنی سن وجود دارد. با افزایش سن، هر دو متغیر اندازه کفش و خزانه واژگان افزایش می‌یابد. بنابراین برای تعیین رابطه واقعی بین دو متغیر، سن باید کنترل شود. متغیری که حذف شده یا ثابت نگه داشته می‌شود، متغیر کنترل نامیده می‌شود. اگر بتوانیم تأثیر سن را حذف کنیم، انتظار می‌رود همبستگی بین این دو متغیر عملا صفر باشد. با انتخاب آزمودنی‌های هم سن و سال می‌توان این عامل را به سادگی کنترل کرد. در این صورت واریانس متغیر سن به صفر یا نزدیک به صفر کاهش می‌یابد.

روش همبستگی جزئی، یک روش مهم کنترل (آماري) است. ضریب همبستگی جزئی^{۳۱} بیان‌گر رابطه بین دو متغیر در حالتی است که تأثیر و نفوذ یک یا چند متغیر دیگر حذف شده باشد.

ضریب همبستگی برای داده‌های فاصله‌ای نشان‌دهنده شدت رابطه دو متغیر است و این ضریب به ما می‌گوید اگر نمره افراد را در یک متغیر بدانیم، با چه دقتی می‌توانیم نمره آن‌ها را در متغیر دیگر پیش‌بینی کنیم. ضریب همبستگی جزئی هم نشان‌دهنده همین اطلاعات است، منتهی با از بین بردن اثر متغیرهای مستقل دیگری که حذف شده‌اند. با محاسبه ضریب همبستگی جزئی بین هر متغیر مستقل و متغیر وابسته و مقایسه آن‌ها

³¹ Partial Correlation

معلوم می‌شود که کدام متغیر مستقل، دقیق‌ترین پیش بین متغیر وابسته است (یعنی کدام متغیر مستقل دارای قوی‌ترین رابطه با متغیر وابسته است). ضریب همبستگی جزئی عمدتاً در مورد متغیرهای فاصله‌ای به کار می‌رود و مبتنی بر ضریب همبستگی پیرسون است.

پیش‌فرض‌ها:

الف) تمامی متغیرها باید دارای مقیاس کمی (فاصله‌ای/نسبی) باشند. یعنی هم دو متغیر اصلی و هم متغیر کنترل باید کمی باشند.

ب) متغیرهای مورد آزمون باید رابطه خطی با همدیگر داشته باشند.

❑ نکته: همبستگی جزئی (یا تفکیکی) تنها برای مدل‌های کوچک مفید است. یعنی مدلهایی که ۳ یا حداکثر ۴ متغیر را در بر می‌گیرند. بنابراین، برای مدل‌های بزرگ، زمانی که داده‌ها فاصله‌ای/نسبی و یا نزدیک به آن‌ها باشند، می‌توانیم از تحلیل مسیر یا مدل‌سازی معادلات ساختاری و زمانی که مقیاس متغیرها در سطح سنجش کیفی (اسمی/ترتیبی) باشند، از مدل-سازی لگاریتم خطی استفاده می‌کنیم.

مفاهیم کلیدی

متغیر آزمون: به متغیری که کنترل می‌شود، می‌گویند. در واقع، همان متغیر کنترل است.

رابطه مرتبه صفر: به رابطه اولیه بین دو متغیر، قبل از کنترل آماری اطلاق می‌شود.
رابطه مشروط (تفکیکی/جزئی): به رابطه بین دو متغیر بعد از حذف اثر متغیر کنترل گفته می‌شود. این رابطه مشروط می‌تواند به صورت مرتبه‌های اول، دوم و... باشد:
 ۱- رابطه مشروط مرتبه اول: به رابطه بین دو متغیر بعد از حذف اثر فقط یک متغیر کنترل گفته می‌شد.

۲- رابطه مشروط مرتبه دوم: به رابطه بین دو متغیر بعد از حذف اثر همزمان دو متغیر کنترل گفته می‌شود.

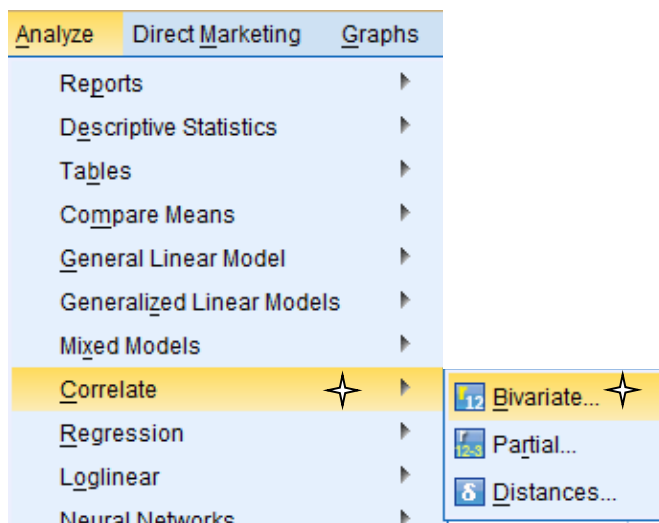
مثال

می‌خواهیم بدانیم که آیا رابطه بین نمره درس SPSS در مقاطع کارشناسی و کارشناسی ارشد دانشجویان واقعی است و یا وجود رابطه بین این دو متغیر با کنترل متغیر تعداد ساعات مطالعه از بین می‌رود. در اینجا می‌خواهیم پی ببریم که آیا با کنترل متغیر سومی به نام تعداد ساعات مطالعه در هفته، رابطه بین متغیرهای نمره درس SPSS در مقطع کارشناسی و کارشناسی ارشد تغییر می‌کند یا خیر. به عبارتی می‌خواهیم بررسی کنیم که نمره درس SPSS در مقطع کارشناسی و در مقطع کارشناسی ارشد واقعی بوده یا کاذب است و نمرات دانشجویان ناشی از عاملی به نام تعداد ساعات مطالعه در هفته است یا خیر.

اجرا:

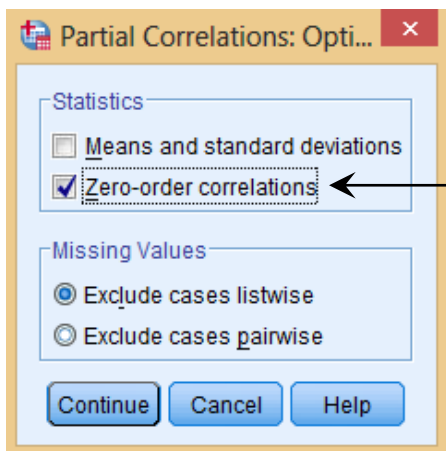
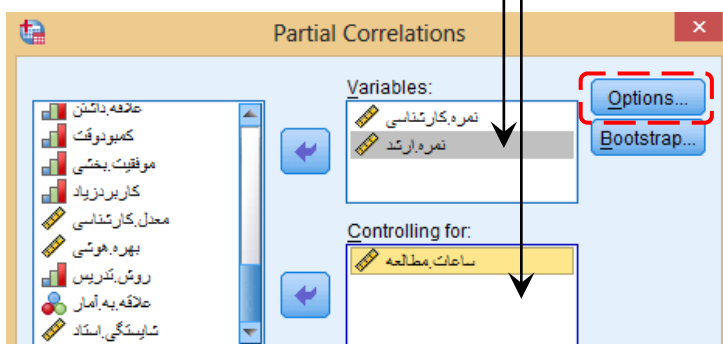
مسیر زیر را دنبال می‌کنیم:

Analyze ---> Correlate ---> Partial



۲- در پنجره‌ای که باز می‌شود، متغیرهای نمره درس SPSS در مقطع کارشناسی و کارشناسی ارشد را وارد کادر Variables می‌کنیم و متغیر ساعات مطالعه که نقش متغیر کنترل را دارد، وارد کادر Controlling for می‌نماییم.

در پنجره‌ای که باز می‌شود، متغیرهای نمره درس SPSS در مقطع کارشناسی و در مقطع کارشناسی ارشد را وارد کادر Variables می‌کنیم. متغیر ساعات مطالعه که نقش متغیر کنترل را دارد، وارد کادر Controlling for می‌نماییم. سپس گزینه Options را انتخاب می‌کنیم.



در کادر اصلی روی Options کلیک کرده و گزینه Zero-order correlation را انتخاب می‌کنیم. در انتها OK و Continue را انتخاب می‌کنیم.

نتایج:

جدول زیر نتایج کنترل متغیر ساعات مطالعه را نشان می‌دهد. در بخش اول نتایج همبستگی پیرسون بین متغیرها ارائه شده است. در این بخش هیچ متغیری کنترل نشده است و تمامی روابط، روابط مرتبه صفر (بدون کنترل آماری) هستند.

در بخش دوم، متغیر ساعات مطالعه کنترل شده است. در نتیجه روابط از نوع مشروط است و چون تنها یک متغیر کنترل داریم رابطه از نوع رابطه مشروط مرتبه اول است. در

بخش دوم جدول چون متغیر ساعات مطالعه کنترل شده است در نتیجه رابطه بین ساعات مطالعه و متغیرهای دیگر نمایش داده نمی‌شود. در این بخش اثر ساعات مطالعه افراد بر رابطه بین نمره کارشناسی و نمره ارشد آنان کنترل شده است و رابطه بین نمره مقطع کارشناسی و نمره مقطع ارشد افراد بعد از حذف اثر متغیر کنترل ساعات مطالعه نمایش داده شده است.

بخش اول: همبستگی متغیرها زمانی

که هیچ متغیری کنترل نشده است

Correlations			نمره کارشناسی	نمره ارشد
Control Variables متغیر کنترل				
متغیر کنترل ↓	نمره ارشد	Correlation (همبستگی)	.411	1
		Significance (2-tailed)	.000	.
		df (درجه آزادی)	90	0
	نمره کارشناسی	Correlation	1	.411
		Significance (2-tailed)	.	.000
		df	0	90
	ساعات مطالعه	Correlation	.601	.664
		Significance (2-tailed)	.000	.000
		df	90	90
ساعات مطالعه ↑	نمره ارشد	Correlation	.020	1
		Significance (2-tailed)	.853	.
		df	89	0
	نمره کارشناسی	Correlation	1	.020
		Significance (2-tailed)	.	.853
		df	0	89

بخش دوم: متغیر ساعات مطالعه در

هفته کنترل شده است.

نحوه تفسیر نتایج آزمون همبستگی تفکیکی

۱- اگر رابطه مرتبه صفر (رابطه اولیه یا بدون کنترل) دو متغیر معنی دار شود، و رابطه مشروط مرتبه اول (با کنترل) هم معنی دار شود، در آن صورت رابطه واقعی است (در چنین حالتی، احتمال وجود رابطه بین متغیر کنترل با هر دو متغیر اصلی کم است).

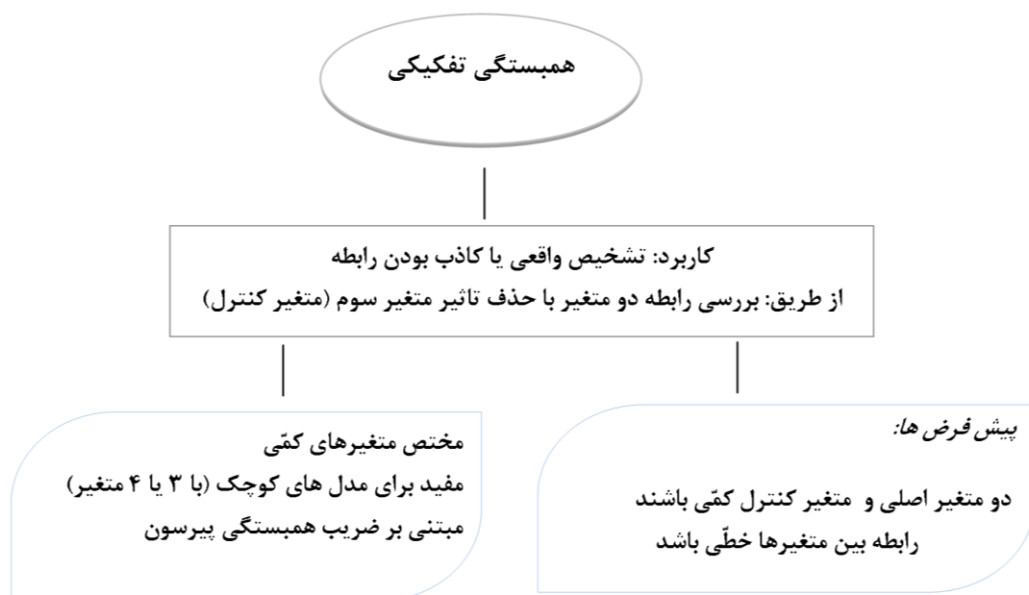
۲- اگر رابطه مرتبه صفر دو متغیر معنی دار شود، اما رابطه آن‌ها با وجود متغیر کنترل معنی دار نشود، در آن صورت رابطه کاذب است. (در چنین حالتی احتمال وجود رابطه بین متغیر کنترل با هر دو متغیر اصلی زیاد است). بنابراین می‌توان چنین تصور کرد که رابطه مشاهده شده بین دو متغیر واقعی نبوده و وجود چنین رابطه‌ای به خاطر وجود رابطه‌ای است که بین هر کدام از این دو متغیر اصلی و متغیر کنترل برقرار است.

۳- اگر رابطه مرتبه صفر دو متغیر معنی دار شود و در عین حال رابطه آن‌ها با وجود متغیر کنترل، باز هم معنی دار شود اما شدت همبستگی بین دو متغیر بعد از ورود متغیر کنترل کمتر شود (همبستگی مرتبه صفر بیشتر از همبستگی مرتبه اول، اما هر دو از نظر آماری معنی‌دار)، به این معناست که بخشی از همبستگی اولیه ناشی از همبستگی هر دو متغیر با متغیر سوم است. به عبارت دیگر، متغیر سوم بخشی از همبستگی اولیه را تبیین یا تفسیر می‌کند.

گزارش:

نتایج آزمون همبستگی تفکیکی با کنترل اثر متغیر ساعات مطالعه در هفته نشان می‌دهد که قبل از کنترل متغیر ساعات مطالعه، رابطه بین دو متغیر نمره کارشناسی و نمره ارشد افراد معنی‌دار است ($P < .001$) و شدت رابطه بین دو متغیر برابر با $.411$ است که به این معناست که بین دو متغیر رابطه وجود دارد. اما بعد از کنترل متغیر سوم (ساعات مطالعه)، مشاهده می‌کنیم که رابطه بین دو متغیر نمره مقطع کارشناسی و ارشد معنی‌دار نیست ($P = .853$).

به عبارتی با کنترل متغیر ساعات مطالعه، رابطه اولیه بین نمره مقطع کارشناسی و نمره مقطع ارشد از بین می‌رود و این‌گونه تفسیر می‌شود که رابطه بین نمره کارشناسی و نمره ارشد کاذب است و تغییرات دو متغیر را متغیر سوم به نام ساعات مطالعه تبیین می‌کند.



تفسیر نتایج		
نتیجه	رابطه مشروط مرتبه اول دارد؟	رابطه مرتبه صفر دارد؟
رابطه اولیه می تواند واقعی باشد	بله : بدون تغییر	بله
بخشی از همبستگی اولیه دو متغیر، ناشی از متغیر سوم است	بله : با کاهش میزان رابطه	بله
رابطه اولیه کاذب است	خیر : رابطه غیرمعنی دار	بله

بخش دوم:

آزمون‌های مقایسه‌ای یا تفاوتی

آزمون‌های تفاوتی آزمون‌هایی هستند که به تفاوت مقدار میانگین یا میانه در بین یک یا چند گروه می‌پردازند. هنگام تحلیل داده‌ها در نرم‌افزار SPSS، برای آن دسته از فرضیه‌های تفاوتی که تفاوت مقدار میانگین در یک یا چند گروه را مطرح می‌کنند، از آزمون‌های پارامتری استفاده می‌شود. آزمون‌های پارامتری t و F را می‌توان برای این-گونه فرضیه‌ها به کار برد. اما برای آن دسته از فرضیه‌های تفاوتی که به بررسی تفاوت مقدار میانه در یک یا چند گروه می‌پردازند، از آزمون‌های ناپارامتری استفاده می‌شود. آزمون‌هایی همچون آزمون دوجمله‌ای، مجذور کای تک‌متغیره، مک‌نمار، کوکران، مان-ویتنی، کروسکال والیس و ویل کاکسون از آزمون‌های ناپارامتری به شمار می‌روند.

پیش‌فرض‌ها

در آزمون فرضیه‌های تفاوتی، زمانی از آزمون‌های پارامتری استفاده می‌شود که شرایط زیر برقرار باشد:

- الف) واریانس نمونه‌ها برابر یا تقریباً برابر باشد.
- ب) داده‌ها در سطح سنجش فاصله‌ای و نسبی باشند.
- ج) توزیع داده‌ها در جامعه، نرمال و یا نزدیک به نرمال باشد.

☑ نکته: اگر گروه‌های مورد بررسی اندازه یکسانی داشته باشند (تعداد اعضای گروه‌های مورد مقایسه برابر باشد) در این صورت فرض برابری واریانس نمونه‌ها چندان مهم نیست. در ضمن در نمونه بزرگ، حتی اگر واریانس یک گروه دو برابر دیگری باشد، باز هم می‌توان از آزمون‌های پارامتری استفاده کرد.

☑ نکته: اگر اطلاعات جمع‌آوری شده این سه شرط را نداشته باشند، می‌توان داده‌ها را از وضعیت پارامتری به ناپارامتری تبدیل کرد و از روش‌های آماری ناپارامتری استفاده کرد. روش

عمده برای تبدیل داده‌های پارامتری به ناپارامتری، رتبه‌بندی کردن (تبدیل داده‌های فاصله‌ای-نسبی به داده‌های ترتیبی) آن‌هاست.

آزمون‌های پارامتری و ناپارامتری براساس مستقل یا وابسته بودن گروه‌ها و تعداد متغیرها به چندین دسته تقسیم می‌شوند که در جدول ۵-۲ نشان داده شده است.

انتخاب آزمون برای مقایسه میانگین‌ها

آزمون‌های t و تحلیل واریانس عمده‌ترین آزمون‌های آماری برای مقایسه میانگین گروه‌ها می‌باشند. از آن‌جا که گروه‌های مورد بررسی ممکن است مستقل یا همبسته باشند، بنابراین هریک از آزمون‌های فوق به دوبرخش مستقل و همبسته تقسیم می‌شوند. تصمیم‌گیری در خصوص این که از هر آزمونی در چه زمانی باید استفاده کرد، مهم‌ترین مسأله در تحلیل داده‌های کمی است. روش انتخاب آزمون مناسب برحسب تعداد گروه‌ها و مستقل یا وابسته بودن گروه‌ها در جدول زیر آمده است.

جدول ۵-۲- انواع آزمون‌های مقایسه‌ای از نوع پارامتریک

نوع متغیر	وضعیت گروه‌ها	نام آزمون
پارامتریک	یک گروه	t تک نمونه‌ای
	دو گروه	t مستقل
		t همبسته (زوجی)
	سه گروه و بیشتر	تحلیل واریانس (ANOVA)

آزمون t تک نمونه‌ای

آزمون‌های آماری پارامتری برای یک گروه زمانی به کار می‌روند که قصد داشته باشیم میانگین یک نمونه را با یک میانگین مفروض و نظری مقایسه کنیم. این میانگین مفروض یا نظری می‌توانید یک مقدار معمول یا رایج، یک مقدار استاندارد و یا یک مقدار مورد انتظار باشد. به عبارتی، زمانی که قصد داشته باشیم میانگین یک متغیر در پژوهش را با یک میانگین تعیین شده مقایسه کنیم از آزمون t تک نمونه‌ای^{۳۲} بهره می‌گیریم.

به عنوان مثال از این آزمون می‌توان در موارد زیر استفاده کرد: مقایسه میانگین درآمد مردم یک شهر با میانگین درآمد کل کشور، مقایسه میزان اعتماد اجتماعی مردم یک شهر با مقدار متوسط، مقایسه میانگین معدل دانشجویان ارشد یک رشته با میانگین معدل دانشجویان ارشد کل یک دانشگاه.

پیش‌فرض‌ها

برای استفاده از آزمون تی تک نمونه‌ای شرط‌ها و پیش‌فرض‌های زیر باید رعایت شود:

- ۱- متغیر مورد نظر باید در سطح سنجش فاصله‌ای/نسبی باشد.
- ۲- توزیع متغیر در جامعه و در نتیجه در نمونه نرمال باشد.

☑ نکته: در بسیاری از پایان‌نامه‌ها و مقالات (داخلی و خارجی)، زمانی که مقیاس متغیر از نوع پاسخ تراکمی (فصل اول - سطح سنجش ترتیبی) بوده است از آزمون تی تک نمونه‌ای استفاده شده است. این امر نشان می‌دهد که پژوهش‌گران بسیاری، برای متغیرهایی که مقیاس تراکمی دارند، با کمی تسامح و تساهل از آزمون‌های پارامتری استفاده می‌کنند.

مثال

می‌خواهیم میانگین بهره‌هوشی دانشجویان ایران را با میانگین بهره‌هوشی کل کشور (کل مردم ایران) مقایسه کنیم. بر اساس استدلال و یافته‌های شخصی ادعا می‌کنیم که بهره‌هوشی دانشجویان بیشتر از سایر مردم است. فرض می‌کنیم طبق آمارها و یافته‌های

³² One-Sample T test

دستگاه‌های ذیربط، میزان بهره هوشی مردم ایران برابر با ۱۰۰ است. در نتیجه ما باید میانگین بهره هوشی دانشجویان را با عدد ۱۰۰ مقایسه کنیم. بنابراین فرضیه‌ای تدوین می‌کنیم و آن را آزمون می‌کنیم:

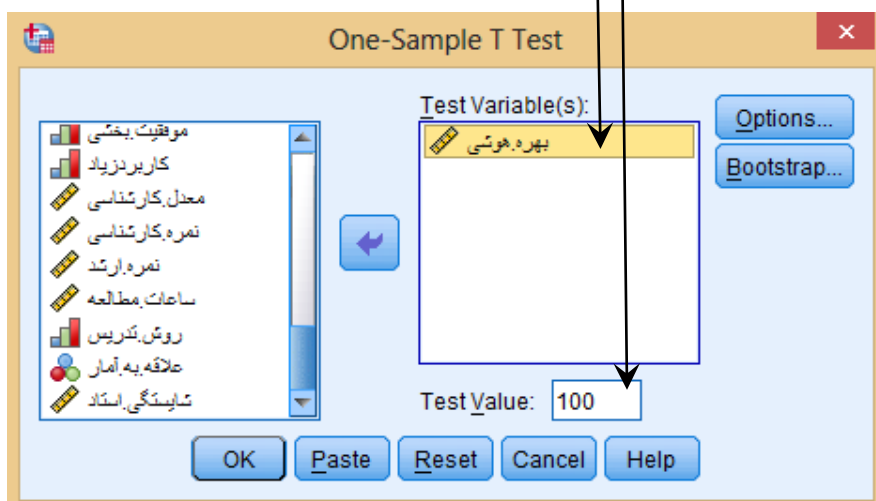
فرضیه: بهره هوشی دانشجویان به طور معنی‌داری از بهره هوشی کل کشور بالاتر است.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Compare Means ---> One-Sample T Test

در کادر متغیرهای آزمون (Test Variables) بهره هوشی دانشجویان را وارد می‌کنیم.
در کادر مقدار آزمون (Test Value) مقدار مفروض یا استاندارد (۱۰۰) را وارد می‌کنیم.
در انتها گزینه OK را انتخاب می‌کنیم.



نتایج:

جدول زیر نشان می‌دهد که تعداد دانشجویانی (N) که ما میانگین بهره هوشی آنان را به دست آورده ایم صد نفر است. میانگین بهره هوشی این نمونه از دانشجویان برابر با ۱۱۶.۶ است که بیشتر از میانگین بهره هوشی کشوری است. اما جهت بررسی معنی‌دار بودن آماری تفاوت بهره هوشی بین دانشجویان و کل مردم باید به جدول بعد و سطح معنی‌داری نگاه کنیم.

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
	فراوانی	میانگین	انحراف استاندارد	خطای استاندارد میانگین
بهره هوشی	100	116.61	15.08	1.51

جدول آزمون تی تک‌نمونه‌ای نشان می‌دهد که بین میانگین بهره هوشی دانشجویان و کل مردم ایران تفاوت معنی‌داری وجود دارد. این تفاوت با بررسی سطح معنی‌داری (Sig) حاصل شده که کمتر از مقدار 0.05 به دست آمده است آشکار می‌گردد. بررسی میانگین‌ها نشان داد که میانگین دانشجویان (116.61) بیشتر از میانگین کل کشور (100) است و این تفاوت از جنبه آماری تأیید می‌شود ($P < 0.05$). در نتیجه فرضیه پژوهش مورد تأیید قرار می‌گیرد و نشان می‌دهد بهره هوشی دانشجویان بیشتر از میانگین بهره هوشی کل کشور است (داده‌ها واقعی نیست و صرفاً برای آموزش است).

One-Sample Test

	Test Value = 100					
	t	df	Sig. (2-tailed)	Mean Difference تفاوت میانگین گروه‌ها	95% Confidence Interval of the Difference	
					Lower	Upper
IQ	11.02	99	.000	16.61	13.62	19.60

گزارش:

میانگین بهره هوشی دانشجویان تفاوت معنی‌داری با میانگین بهره هوشی کشور دارد ($t = 11.02$ و $df = 99$ و $P < 0.001$). میانگین بهره هوشی دانشجویان ایرانی 116.61 به دست آمده است که به طور معنی‌داری بیشتر از میانگین بهره هوشی کل مردم ایران (100) است.

آزمون t مستقل

آزمون t گروه‌های مستقل^{۳۳} میانگین دو گروه از پاسخگویان را با هم مقایسه می‌کند. به عبارتی، در این آزمون، میانگین‌های به دست آمده از نمونه‌های تصادفی مود قضاوت قرار می‌گیرند. بدین معنی که از دو جامعه مختلف، نمونه‌هایی اعم از این که تعداد نمونه مساوی یا غیرمساوی باشند، به طور تصادفی انتخاب کرده و میانگین‌های آن دو جامعه را با هم مقایسه می‌کنیم (منصورفر، ۱۳۸۴: ۲۰۱).

به عنوان مثال افراد برحسب جنس به دو گروه زن و مرد تقسیم می‌شوند و می‌توانیم دو گروه زنان و مردان را در متغیرهای مختلف با هم مقایسه کنیم. مثلاً می‌توانیم زنان و مردان را در متغیرهای میزان درآمد (به تومان)، بهره‌هوشی، میزان اضطراب و ... مقایسه کنیم. می‌توانیم طول عمر مردم شهرنشین را با مردم روستانشین مقایسه کنیم و یا می‌توانیم میزان امید به زندگی را در بین مردان شاغل و مردان بیکار باهم مقایسه کنیم.

پیش‌فرض‌ها

برای استفاده از آزمون t مستقل، پیش‌فرض‌های زیر باید برقرار باشد:

- ۱- مقیاس متغیر وابسته باید کمی و در سطح فاصله‌ای/نسبی باشد.
- ۲- مقیاس متغیر مستقل باید کیفی و در سطح اسمی باشد.
- ۳- مقادیر دو متغیر باید مستقل و از دو جامعه باشند.
- ۴- توزیع داده‌های متغیر وابسته باید نرمال باشد (حبیب‌پور و صفری، ۱۳۸۸: ۵۴۶).

مثال

در مثال کتاب، معدل درس SPSS (در مقطع کارشناسی) برای دانشجویان دختر و پسر ارائه شده است. می‌خواهیم میانگین نمره دانشجویان دختر و پسر را در درس SPSS با یکدیگر مقایسه کنیم. توزیع متغیر نمره درس SPSS را هم نرمال در نظر می‌گیریم. در ارتباط با برقراری پیش‌فرض‌ها توجه شود که بهره‌هوشی یک متغیر کمی است و جنسیت متغیری کیفی (اسمی) است. فرضیه‌ای در این زمینه تدوین می‌کنیم و آن را می‌آزماییم (به جمله‌بندی و شیوه طرح فرضیه مقایسه‌ای توجه شود):

³³ Independent-Samples T test

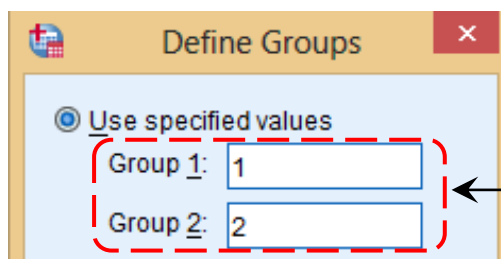
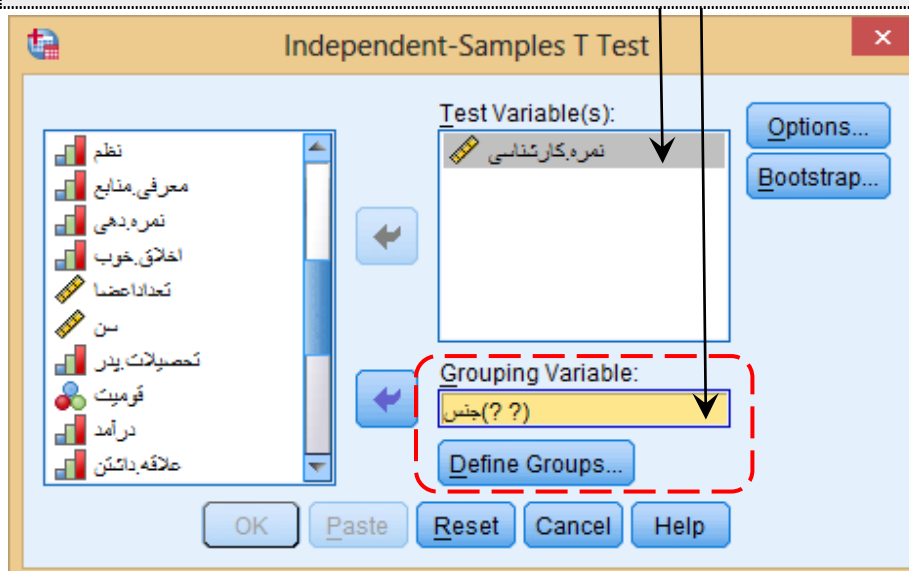
فرضیه: میانگین نمره درس SPSS در دانشجویان دختر و پسر متفاوت است.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Compare Means ---> Independent – Samples T Test

در کادر متغیرهای آزمون (Test Variables) نمره کارشناسی را وارد می‌کنیم. در این کادر می‌توانیم دو یا چند متغیر را وارد کنیم و دو گروه دانشجویان را در چند متغیر مقایسه کنیم. متغیر گروه‌بندی، یعنی جنس را وارد کادر متغیر گروه‌بندی Grouping Variable می‌کنیم. سپس گزینه پایین آن یعنی معرفی گروه‌ها (Define Groups) را انتخاب می‌کنیم.



در پنجره Define Groups باید کدهایی را که به دو طبقه دختر و پسر هنگام تعریف متغیر جنسیت داده‌ایم را وارد کنیم. ما به زن که ۱ و به مرد که ۲ داده‌ایم و این دو عدد را وارد کادرهای Group می‌کنیم.

نتیج:

جدول زیر توصیفی از متغیر نمره کارشناسی افراد برحسب جنس است. مقادیر فراوانی هرگروه، میانگین، انحراف استاندارد و خطای استاندارد میانگین در این جدول ارائه شده است. همان‌طور که مشاهده می‌گردد میانگین نمره دختران در درس SPSS برابر با ۱۷.۱۲ و میانگین پسران برابر با ۱۶.۶۵ به‌دست آمده است. دختران میانگین نمره بالاتری در درس SPSS دارند، اما معنی‌دار بودن این تفاوت میانگین را باید با بررسی سطح معنی‌داری مشخص کرد.

Group Statistics

جنس	N	Mean	Std. Deviation	Std. Error Mean
زن	50	17.12	2.46	.348
نمره کارشناسی مرد	50	16.65	1.61	.228

در جدول بعد نتایج آزمون t گروه‌های مستقل گزارش شده است. جدول زیر دارای دو بخش است. بخش اول مربوط به آزمون لَوْن یا لَوْن^{۳۴} است. از آزمون لَوْن جهت بررسی تفاوت واریانس دو گروه استفاده می‌شود. کاربرد این آزمون بدین‌صورت است که چنانچه سطح معنی‌داری آزمون لَوْن بیشتر از ۰.۰۵ باشد از نتایج ردیف بالای آزمون تی (t-test) استفاده می‌کنیم و اگر مقدار سطح معنی‌داری آزمون لَوْن کمتر از ۰.۰۵ باشد از نتایج ردیف پایین استفاده می‌کنیم.

مقدار سطح معنی‌داری آزمون لَوْن (Sig) برابر با ۰.۰۰۲ به دست آمده است که کمتر از مقدار ۰.۰۵ است، در نتیجه ما باید به نتایج ردیف پایین آزمون تی استناد کنیم. با توجه به سطح معنی‌داری (Sig. (2-tailed)) آزمون تی مستقل که برابر با ۰.۲۵۷ به‌دست آمده است و بیشتر از مقدار مفروض ۰.۰۵ است نتیجه می‌گیریم که بین میانگین نمره درس SPSS در پسران و دختران تفاوت معنی‌داری وجود ندارد. در نتیجه فرضیه صفر تایید می‌شود و فرضیه پژوهش رد می‌شود.

³⁴ Levene's test

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means			
		آزمون لون: سنجش برابری واریانس ها		آزمون t: سنجش برابری میانگین ها			
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference
نمره کارشناسی	Equal variances assumed	10.36	.002	1.18	98	.254	.475
	Equal variances not assumed			1.14	84.57	.257	.475

گزارش:

آزمون تی مستقل نشان داد که تفاوت معنی داری بین نمره کارشناسی درس SPSS در بین دانشجویان دختر و پسر وجود ندارد ($P = .254$ و $df = 84.57$ و $t = 1.18$). میانگین نمره درس SPSS پسران ۱۶.۶۵ و دختران ۱۷.۱۲ است و این تفاوت از جنبه آماری معنی دار نیست.

آزمون t همبسته

آزمون تی گروه‌های همبسته^{۳۵} (t همبسته یا جفتی) معمولاً در پژوهش‌های آزمایشی و تجربی و در شرایطی بکار می‌رود که قصد داریم تأثیر نوعی مداخله را بر روی یک گروه و در دو زمان متفاوت مورد بررسی قرار دهیم (مثلاً، مقایسه وضعیت بیمار در زمان‌های قبل و بعد از استفاده از یک دارو، و یا بررسی اثر اجرای یک دوره آموزشی شغلی بر روی عملکرد شغلی افراد). در آزمون t همبسته، هر فرد یا شیء دوبار در دو وضعیت متفاوت (معمولاً قبل و بعد) مورد مشاهده قرار می‌گیرد. به عنوان مثال، میزان درد بیماران قبل و بعد از مصرف یک داروی مسکن جدید. فرض صفر در چنین آزمون‌هایی این است که اختلافی بین مقادیر میانگین در دو نمونه از جامعه وجود ندارد.

❑ نکته: زمانی که داده‌ها به صورت جفتی باشند (مانند مقادیر یک اندازه‌گیری در دو مقطع زمانی)، در آن صورت نتایج آزمون t همبسته، در مقایسه با آزمون t مستقل، مقدار واریانس‌ها را کمتر نشان داده و در نتیجه از توان بالاتری برای کشف تفاوت‌ها برخوردار است.

پیش‌فرض‌ها

زمانی می‌توانیم از آزمون t همبسته استفاده کنیم که پیش‌فرض‌های زیر برقرار باشد:

- ۱- مقیاس متغیر وابسته باید کمی و در سطح فاصله‌ای/نسبی باشد.
- ۲- مقادیر دو متغیر باید جفت (وابسته) و از یک جامعه باشند.
- ۳- توزیع داده‌های متغیر وابسته باید نرمال باشد (حبیب‌پور و صفری، ۱۳۸۸: ۵۴۶).

مثال

در مثال کتاب، هر دانشجوی دو نمره در درس SPSS کسب کرده است: نمره در مقطع کارشناسی و نمره در مقطع ارشد. در این جا ما با یک گروه از پاسخگویان مواجهیم که در یک درس مشخص دارای دو نمره در دو مقطع زمانی (مقطع کارشناسی و کارشناسی ارشد) هستند. به بیان دیگر ما می‌خواهیم ببینیم که آیا بین نمره درس SPSS دانشجویان، زمانی که در مقطع کارشناسی تحصیل می‌کرده‌اند و زمانی که در مقطع کارشناسی ارشد تحصیل می‌کرده‌اند تفاوتی وجود دارد؟ و دانشجویان در کدام مقطع

³⁵ Paired-Samples T test

تحصیلی میانگین بالاتری در درس SPSS کسب کرده‌اند؟ در نتیجه فرضیه مناسبی طرح می‌کنیم:

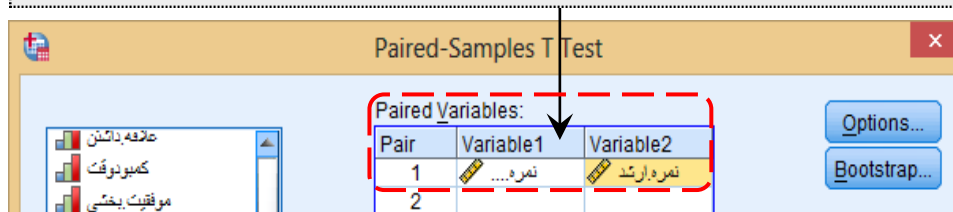
فرضیه: میانگین نمره درس SPSS دانشجویان در مقطع کارشناسی و کارشناسی ارشد متفاوت است.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Compare Means ---> Paired – Samples T Test

در ابتدا متغیر نمره کارشناسی و سپس نمره ارشد پاسخگویان را از کادر سمت چپ وارد کادر متغیرهای همبسته (Paired Variables) می‌کنیم و در انتها گزینه OK را انتخاب می‌کنیم.



نتایج:

جدول زیر توصیف نمره درس SPSS دانشجویان در دو مقطع کارشناسی و کارشناسی ارشد است. همان طور که مشاهده می‌شود میانگین نمره در مقطع کارشناسی بیشتر از مقطع ارشد است. میانگین نمره پاسخگویان در مقطع کارشناسی برابر با ۱۶.۸۹ و در مقطع ارشد برابر با ۱۵.۹۸ است و اختلاف میانگین دو مقطع تحصیلی برابر با ۰.۹۰ است.

Paired Samples Statistics

	Mean	N	Std. Deviation	Std. Error Mean
Pair 1 نمره کارشناسی	16.89	100	2.086	.208
نمره ارشد	15.98	100	2.197	.219

نتایج آزمون تی جفتی در جدول بعد گزارش شده است. نتایج نشان می‌دهد که بین نمره پاسخگویان در دو مقطع کارشناسی و کارشناسی ارشد تفاوت معنی‌دار وجود دارد ($P < .05$). سطح معنی‌داری به دست آمده (Sig) کمتر از مقدار مفروض $.05$ به دست آمده است و بدین معناست که بین میانگین نمره درس SPSS دانشجویان در دو مقطع کارشناسی و کارشناسی ارشد تفاوت معنی‌دار وجود دارد و دانشجویان کارشناسی میانگین بالاتری در درس SPSS کسب کرده‌اند.

Paired Samples Test

	Paired Differences			t	df	Sig. (2-tailed)
	Mean	Std. Deviation	Std. Error Mean			
Pair 1 نمره ارشد - نمره کارشناسی	.905	2.35	.23	3.84	99	.000

گزارش:

آزمون تی همبسته نشان داد که تفاوت معنی‌داری بین میانگین نمره دانشجویان در دو مقطع کارشناسی و کارشناسی ارشد وجود دارد ($P < .001$ و $df = 99$ و $t = 3.844$). میانگین نمره درس SPSS دانشجویان هنگامی که در مقطع کارشناسی مشغول به تحصیل شده‌اند برابر با 16.89 و هنگامی که در مقطع کارشناسی ارشد مشغول به تحصیل بوده‌اند برابر با 15.98 به دست آمده است. اختلاف میانگین دو مقطع تحصیلی $.905$ است. دانشجویان در مقطع کارشناسی در مقایسه با مقطع کارشناسی ارشد، به طور معنی‌داری میانگین بالاتری در درس SPSS کسب کرده‌اند.

تعریف

۱- جدول زیر تعداد ضربان قلب ۱۵ مرد مبتلا به تپش قلب را در دو مقطع زمانی نشان می‌دهد. در ابتدا تعداد ضربان قلب اندازه‌گیری شد و سپس به آزمودنی‌ها (بیماران) یک داروی جدید کاهش ضربان قلب داده شد. یک ساعت بعد مجدداً تعداد ضربان قلب آنان اندازه‌گیری شد. آیا استفاده از داروی جدید منجر به کاهش تعداد ضربان قلب آزمودنی‌ها شده است؟ به بیان دیگر آیا تعداد ضربان قلب افراد بعد از مصرف داروی جدید تفاوت معنی‌داری کرده است؟

آزمودنی‌ها	۱	۲	۳	۴	۵	۶	۷	۸	۹	۱۰	۱۱	۱۲	۱۳	۱۴	۱۵
تعداد ضربان قلب اولیه (قبل از مصرف دارو)	۷۵	۸۲	۷۴	۷۶	۹۰	۸۴	۸۴	۷۲	۷۶	۷۰	۷۴	۷۰	۷۱	۷۸	۸۶
تعداد ضربان قلب ثانویه (بعد از مصرف دارو)	۷۵	۷۴	۷۲	۷۱	۸۰	۷۰	۶۹	۶۸	۷۵	۷۱	۷۳	۶۶	۶۵	۷۰	۸۵

آزمون تحلیل واریانس یک راهه

فیشر اولین کسی بود که مفهوم «واریانس» را به کار برد و با ارائه نظریه‌ای در رابطه با تحلیل واریانس، ارزش و اهمیت آن را در تحلیل پدیده‌های واقعی توجیه و تبیین کرد. تحلیل واریانس (یک‌راهه)^{۳۶} که به آن ANOVA یا آزمون F می‌گویند، یکی از تکنیک‌های آماری مؤثر و پرکاربرد در پژوهش‌های اقتصادی، اجتماعی، علوم تربیتی، روان‌شناسی، مدیریت و حتی کشاورزی، بیولوژی و غیره است. همان‌طور که قبلاً نیز بیان شد، زمانی که بخواهیم میانگین‌های دو جامعه (یا نمونه) را با همدیگر مقایسه کنیم و معنی‌داری تفاوت بین آن‌ها را بررسی نماییم، از آزمون های t استفاده می‌کنیم. اما زمانی که پژوهش‌گر بخواهد به بررسی تفاوت‌های میانگین بیش از دو جامعه (یا نمونه) بپردازد، به کارگیری آزمون‌هایی چون t امکان‌پذیر نخواهد بود. برای این منظور در این‌گونه پژوهش‌ها از روش تحلیل واریانس یا آزمون F استفاده می‌گردد. به عنوان مثال اگر بخواهیم تفاوت درآمد بین سه گروه کارمند، کشاورز و کارگر را بررسی کنیم، از آزمون F یا تحلیل واریانس استفاده می‌کنیم. این روش تفاوت معنی‌دار بین درآمد گروه‌های شغلی سه‌گانه را از طریق مقایسه میانگین درآمدهای آنان بررسی می‌کند.

آزمون تحلیل واریانس یک راهه یا آزمون F، گسترش یافته و تعمیم یافته آزمون t مستقل است و زمانی به کار می‌رود که ما بخواهیم بیش از دو (سه یا بیشتر) گروه یا وضعیت را با هم مقایسه کنیم. زمانی که تعداد گروه‌ها سه یا بیشتر باشد، پژوهش‌گران و آمارشناسان معمولاً به جای انجام چند آزمون تی مستقل، از آزمون تحلیل واریانس استفاده می‌کنند. مثلاً زمانی که ما بخواهیم میزان درآمد مردم سه شهر تهران، کرج و مشهد را با هم مقایسه کنیم و یا بخواهیم تعداد ساعات مطالعه دانشجویان چهار مقطع فوق دیپلم، کارشناسی، کارشناسی ارشد و دکترا را با هم مقایسه کنیم از این آزمون استفاده می‌کنیم.

تحلیل واریانس (آنووا) در واقع روشی برای آزمایش تفاوت بین گروه‌های مختلف داده‌ها یا نمونه‌هاست. این روش کل واریانس موجود در یک مجموعه از داده‌ها را به دو بخش تقسیم می‌کند. بخشی از این واریانس ممکن است بخاطر شانس و تصادف حادث شده

³⁶ One-Way ANOVA

باشد و بخش دیگر ممکن است ناشی از دلایل یا عوامل خاصی باشد. یعنی واریانس موجود ممکن است ناشی از تفاوت بین گروه‌های مورد مطالعه و یا به خاطر تفاوت موجود در درون نمونه‌ها حادث شده باشد. بنابراین تحلیل واریانس به عنوان یک روش تحلیل یا بررسی مجموع این تفاوت‌ها به تبیین پدیده موردنظر می‌پردازد.

به طور خلاصه در تحلیل واریانس دو نوع واریانس باید برآورد گردد که یکی از آن‌ها واریانس بین گروه‌ها و دیگری واریانس درون گروه‌هاست. واریانس بین گروه‌ها عبارت از مجموع مجذورات انحراف میانگین هر گروه از میانگین کل، و واریانس درون گروه‌ها عبارت از مجموع مجذورات انحراف مقادیر هر فرد از میانگین گروه خود می‌باشد.

اگر در نتایج به دست آمده از آزمون تحلیل واریانس در نرم‌افزار SPSS، سطح معنی‌داری به دست آمده برای F ، کوچکتر از 0.05 باشد ($P < 0.05$)، می‌توان قضاوت کرد که بین میانگین گروه‌های مورد مطالعه در سطح 95 تفاوت درصد وجود دارد (یعنی بین میانگین حداقل دو گروه تفاوت وجود دارد) و در صورتی که سطح معنی‌داری کوچکتر از 0.01 باشد، می‌توان مدعی شد که تفاوت بین گروه‌ها در سطح 99 درصد معنی‌دار است.

پیش‌فرض‌ها

پیش‌فرض‌های زیر باید برقرار باشد تا بتوانیم از آزمون تحلیل واریانس استفاده کنیم:

۱. متغیر وابسته باید در سطح سنجش فاصله‌ای/نسبی باشد.
۲. متغیر مستقل باید کیفی و در سطح سنجش اسمی یا ترتیبی باشد.
۳. توزیع متغیر در جامعه و در نتیجه در نمونه نرمال باشد.
۴. اعضای هر گروه باید یک نمونه تصادفی مستقل از جامعه باشند.

آزمون تعقیبی

نتیجه آزمون تحلیل واریانس تفاوت بین گروه‌های مورد مقایسه را به طور کلی نشان می‌دهد اما نشان نمی‌دهد که کدام یک از این گروه‌ها با هم تفاوت دارند. یعنی ممکن است نتیجه آزمون تحلیل واریانس در مقایسه میزان درآمد پزشکان، کارگران و کشاورزان این باشد میزان درآمد این سه گروه تفاوت دارد، اما نشان نمی‌دهد که کدام گروه‌ها درآمدشان متفاوت است. اگر نتیجه آزمون تحلیل واریانس معنی‌دار باشد و نشان دهد که

درآمد این سه گروه با یکدیگر متفاوت است، نمی‌توان ادعا کرد که درآمد همه گروه‌های مورد مقایسه با هم متفاوت است. ممکن است میزان درآمد پزشکان بیشتر از درآمد کشاورزان و کارگران باشد اما درآمد کشاورزان و کارگران تقریباً یکسان باشد و تفاوتی با هم نداشته باشند. آزمون تحلیل واریانس فقط وجود تفاوت معنی‌دار بین حداقل دو گروه از گروه‌های مورد مقایسه را نشان می‌دهد، اما برای این که بدانیم کدام گروه‌ها در متغیر موردنظر با یکدیگر تفاوت دارند باید از آزمون‌های تعقیبی استفاده کنیم. آزمون‌های تعقیبی به مقایسه جفتی و دوبه‌دوی گروه‌ها می‌پردازد.

آزمون‌های تعقیبی انواع مختلفی دارند ولی می‌توان گفت که سه مورد از آن‌ها رایج‌تر است: آزمون توکی^{۳۷}، آزمون شِفِه^{۳۸} و آزمون دانکن^{۳۹}. از آزمون توکی زمانی که حجم گروه‌ها برابر (یا تقریباً برابر) باشد استفاده می‌کنیم و از آزمون شِفِه زمانی که حجم گروه‌ها نابرابر باشد استفاده می‌کنیم. در مثال کتاب، چون حجم گروه‌ها متفاوت است (ترک ۱۸ نفر، فارس ۴۵، کرد ۱۵ و سایر قومیت‌ها ۲۲) از آزمون تعقیبی شِفِه استفاده می‌کنیم. جدول زیر انواع آزمون‌های تعقیبی (مقایسه‌ای) چندگانه را نشان می‌دهد.

☑ نکته: خنثی بودن حجم نمونه برای آزمون به این معناست که آن آزمون به برابری بودن حجم نمونه گروه‌های مورد مقایسه، حساس نیست.

☑ نکته: زمانی که پیش‌فرض برابری واریانس‌ها برقرار باشد (واریانس‌ها برابر یا همگن باشد)، از آزمون‌های بخش اول و زمانی که واریانس بین گروه‌ها نابرابر باشد از آزمون‌های بخش دوم (تی دو تمهنه و جیمز-هوئل) استفاده می‌کنیم.

☑ نکته: محافظه کار بودن یک آزمون به این معناست که آن آزمون برای این که تفاوت بین میانگین‌ها را معنادار نشان دهد، باید میزان بالایی از این تفاوت بین گروه‌ها وجود داشته باشد. لیبرال بودن به این معناست که آن آزمون تفاوت گروه‌ها از هم را زیاد نشان می‌دهد. روش‌های میانه رو بین این دو قرار دارند.

³⁷ Tukey

³⁸ Scheffe

³⁹ Duncan

جدول ۵-۳- انواع آزمون‌های مقایسه‌ای چندگانه (آزمون‌های تعقیبی)

نوع فرض درباره واریانس	آزمون	حجم نمونه	خصوصیات
بخش اول: برابری واریانس‌ها	LSD فیشر	خنثی	لیبرال ترین روش است کم تر مورد استفاده قرار می گیرد
	Bonferroni (بونفرونِی)	خنثی	محافظه کارترین روش است توان این آزمون در تمایز گروه‌ها از همدیگر کم است برای آزمون تعداد جفت‌های کم تر، از این آزمون استفاده شود تا توکی
	Scheffe (شِفِه)	خنثی	به مقایسه جفتی گروه‌ها می پردازد محافظه کار است عدم حساسیت به انحراف از نرمال بودن توزیع داده‌ها و همگونی واریانس‌ها
	Tukey (توکی)	برابر	یک روش میانه رو است توان این آزمون در تمایز گروه‌ها از همدیگر کم است قوی تر از شفه است مناسب برای آزمون تعداد زیادی جفت میانگین
	Duncan (دانکن)	خنثی	
بخش دوم: عدم برابری واریانس‌ها	Tamhane's T2 (تی دو تمهنه)	نابرابر	محافظه کار است
	Games-Howell (جیمز-هوئل)	بسیار متفاوت و نابرابر	عملکرد بهتری نسبت به بقیه دارد اغلب لیبرال است

(منبع: حبیب پور و صفری، ۱۳۸۸: ۵۶۰-۵۶۵)

مثال

می‌خواهیم نمره درس SPSS دانشجویان قومیت‌های مختلف را مقایسه کنیم. یعنی می‌خواهیم بدانیم که دانشجویان کدام یک از قومیت‌ها (ترک، فارس، کرد و سایر قومیت‌ها) در درس SPSS میانگین نمره بالاتری کسب کرده‌اند. در نتیجه فرضیه زیر را جهت آزمون تدوین می‌کنیم:

فرضیه: میانگین نمره درس SPSS دانشجویان قومیت‌های گوناگون بایکدیگر متفاوت است.

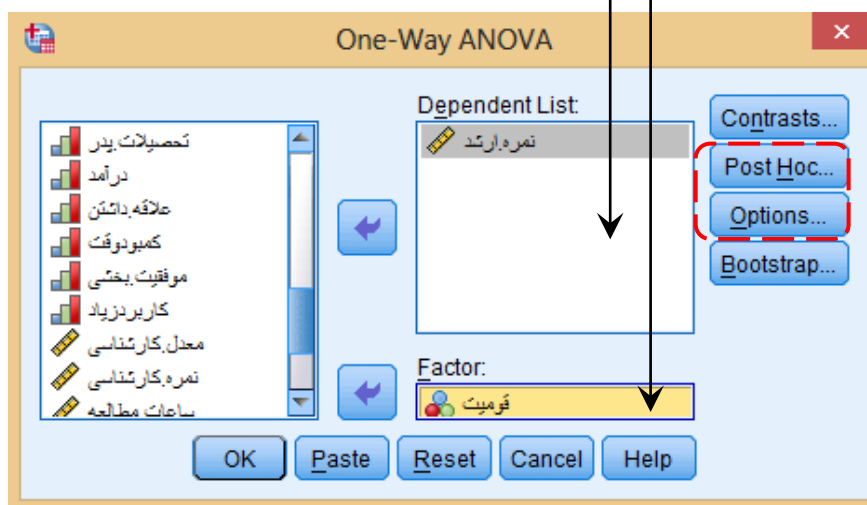
اجرا:

مسیر زیر را دنبال می‌کنیم:

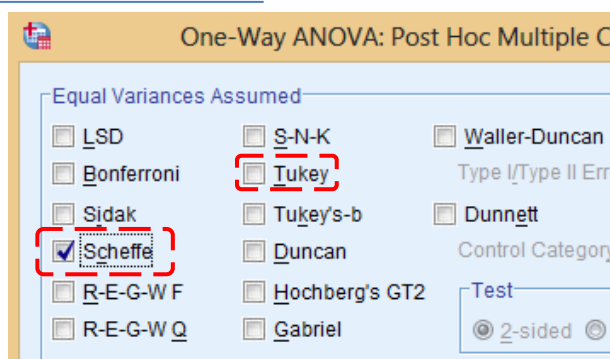
Analyze ---> Compare Means ---> One – Way ANOVA

در کادر فهرست متغیرهای وابسته (Dependent List) متغیر وابسته یا متغیر مورد مقایسه را وارد می‌کنیم، یعنی نمره درس SPSS دانشجویان.

در کادر عامل (Factor)، متغیر مستقل را وارد می‌کنیم، یعنی قومیت دانشجویان. سپس گزینه‌های Options و Post Hoc (آزمون‌های تعقیبی) را انتخاب می‌کنیم.

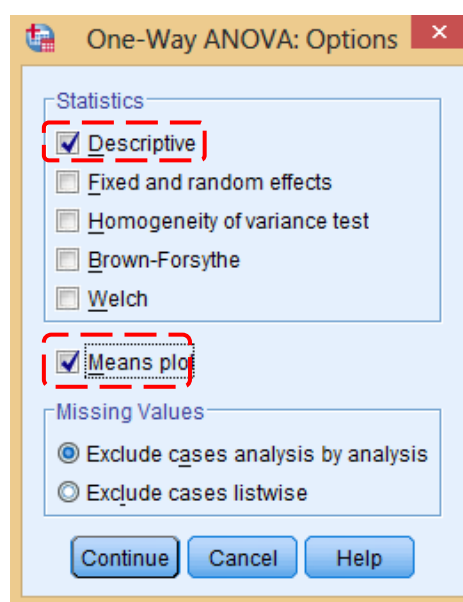


در پنجره Post Hoc گزینه آزمون تعقیبی شِفِه (Scheffe) را انتخاب می‌کنیم. چون حجم گروه‌ها با یکدیگر متفاوت است از آزمون شفه استفاده می‌کنیم. اگر تعداد اعضای گروه‌ها برابر بود از آزمون تعقیبی توکی (Tukey) استفاده می‌کنیم.



در پنجره Options، گزینه های آمار توصیفی (Descriptive) و نمودار میانگین ها (Means plot) را انتخاب می کنیم.

در انتها Continue و OK را انتخاب می کنیم.



نتایج:

در جدول زیر آمار توصیفی هر کدام از قومیت ها ارائه شده است. همان طور که مشاهده می شود فراوانی (N) دانشجویان قومیت های مختلف با یکدیگر متفاوت است و بدین جهت در آزمون های تعقیبی از آزمون تعقیبی شفه استفاده شد. مقایسه میانگین قومیت ها نشان می دهد که دانشجویان با قومیت ترک بالاترین میانگین نمره (۱۸.۳۹) را کسب کرده اند و بعد از آن، دانشجویان کرد با میانگین ۱۶.۰۷ بالاترین میانگین را دارند.

Descriptives

نمره ارشد

	N	Mean	Std. Deviation	Std. Error
ترک	18	18.39	1.06	.251
فارس	45	15.58	2.50	.372
کرد	15	16.07	1.61	.416
سایر قومیت‌ها	22	14.77	.43	.091
Total	100	15.98	2.19	.219

جدول زیر جدول آنووا یا تحلیل واریانس است و مهم‌ترین جدول آزمون تحلیل واریانس محسوب می‌شود. با مراجعه به سطح معنی‌داری (Sig) جدول می‌توان در مورد تأیید یا رد شدن فرضیه و تفاوت میانگین‌ها از جنبه آماری مختلف اظهار نظر کرد.

سطح معنی‌داری به دست آمده کمتر از مقدار تعیین شده 0.05 است ($P < 0.05$) که نشان می‌دهد بین میانگین چهار قومیت در درس SPSS تفاوت معنی‌داری وجود دارد و حداقل دو قومیت دارای میانگین متفاوتی در این درس هستند. از روی نتایج جدول زیر (ANOVA) و میانگین هرکدام از قومیت‌ها نمی‌توان در مورد این که نمره دانشجویان کدام قومیت‌ها با یکدیگر تفاوت دارد به طور قطعی اظهار نظر کرد. جهت مقایسه جفتی دانشجویان با قومیت‌های مختلف، از آزمون‌های تعقیبی استفاده می‌کنیم.

ANOVA

نمره ارشد

	Sum of Squares مجموع مجذورات	df	Mean Square میانگین مجذورات	F آماره F	Sig. سطح معنی‌داری
Between Groups تفاوت بین گروهی	143.71	3	47.90	13.75	.000
Within Groups تفاوت درون گروهی	334.32	96	3.48		
Total	478.03	99			

نتیجه آزمون تعقیبی شفه جهت مقایسه دوبه‌دوی قومیت‌ها در جدول زیر ارائه شده است. با بررسی جفتی قومیت‌ها متوجه می‌شویم که قومیت ترک با همه قومیت‌ها

دارای تفاوت معنی‌داری است ($P < .05$)، ولی بین قومیت‌های دیگر با هم تفاوتی دیده نمی‌شود. دانشجویانی که قومیت ترک دارند دارای میانگین نمره متفاوتی در مقایسه با سایر قومیت‌ها هستند و این تفاوت میانگین از جنبه آماری معنی‌دار است. با توجه به مقادیر میانگین نمره دانشجویان ترک که بیشتر از دانشجویان سایر قومیت‌ها شده است می‌توانیم بگوییم که میانگین نمره درس SPSS در دانشجویان ترک به طور معنی‌داری بالاتر از میانگین نمره دانشجویان قومیت‌های دیگر است. ولی بین میانگین نمره درس SPSS دانشجویان قومیت‌های دیگر با یکدیگر تفاوتی وجود ندارد.

Multiple Comparisons

Dependent Variable: نمره ارشد

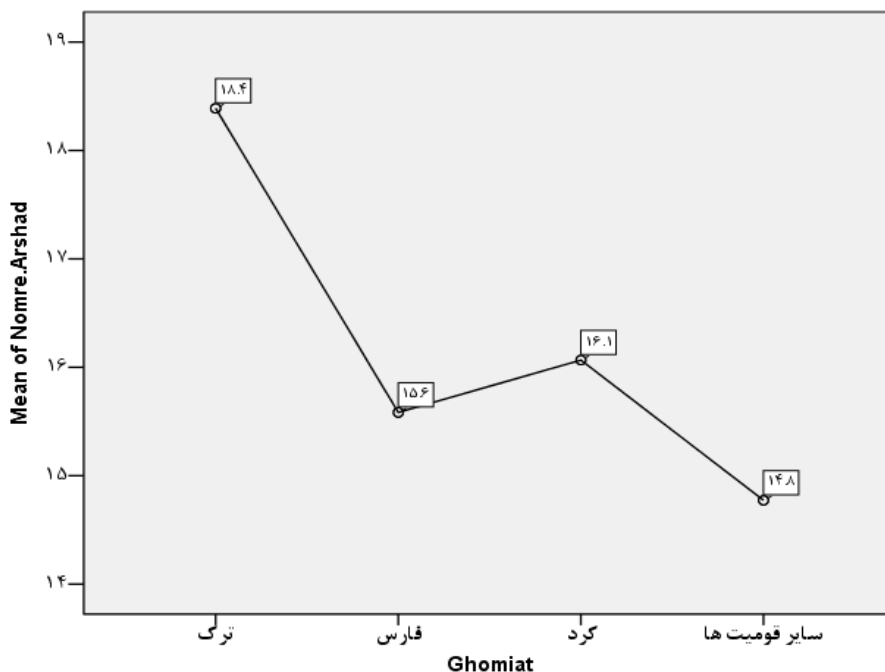
آزمون تعقیبی شفه Scheffe

قومیت (I)	قومیت (J)	Mean Difference (I-J) تفاوت میانگین دو گروه	Std. Error خطای استاندارد	Sig.
ترک	فارس	2.80*	.520	.000
	کرد	2.3*	.652	.008
	سایر قومیت‌ها	3.62*	.593	.000
فارس	ترک	-2.80*	.520	.000
	کرد	-.48	.556	.860
	سایر قومیت‌ها	.81	.485	.430
کرد	ترک	-2.32*	.652	.008
	فارس	.48	.556	.860
	سایر قومیت‌ها	1.29	.624	.239
سایر قومیت‌ها	ترک	-3.62*	.593	.000
	فارس	-.81	.485	.430
	کرد	-1.29	.624	.239

*. The mean difference is significant at the 0.05 level.

علامت (*) تفاوت میانگین گروه‌ها را در سطح معنی‌داری کمتر از $.05$ نشان می‌دهد

نمودار میانگین دانشجویان قومیت‌های مختلف در درس SPSS در مقطع ارشد ارائه شده است. همان‌طور که مشاهده می‌گردد دانشجویان ترک بالاترین میانگین را دارند.



شکل ۵-۱ نمودار خطی میانگین نمره ارشد در دانشجویان با قومیت‌های مختلف

گزارش:

در گزارش می‌نویسیم: دانشجویان قومیت‌های مختلف دارای میانگین نمره متفاوتی در درس SPSS در مقطع ارشد هستند ($P < .001$ و $F_{(3,96)} = 13.755$). آزمون تعقیبی شفه نشان داد که دانشجویان ترک میانگین متفاوتی نسبت به تمامی قومیت‌های دیگر دارند ولی بین قومیت‌های دیگر تفاوتی در میانگین نمره SPSS دیده نمی‌شود. نمودار میانگین قومیت‌ها نشان می‌دهد که دانشجویان ترک بالاترین میانگین ($M=18.39$) را کسب کرده‌اند.

☑ نکته: اندیس‌های F نشان‌دهنده درجه آزادی بین‌گروهی و درون گروهی است (درجه آزادی در جدول تحلیل واریانس ارائه شده است).

تعریف

به جدول صفحه ۱۸۸ رجوع کنید:

۱- آیا میزان اعتماد اجتماعی شهروندان تهرانی بالاست؟ نمره کل مقیاس اعتماد اجتماعی را با عدد ۳ که مقدار متوسط است (دامنه نمرات از ۱ تا ۵ است) مقایسه کنید.

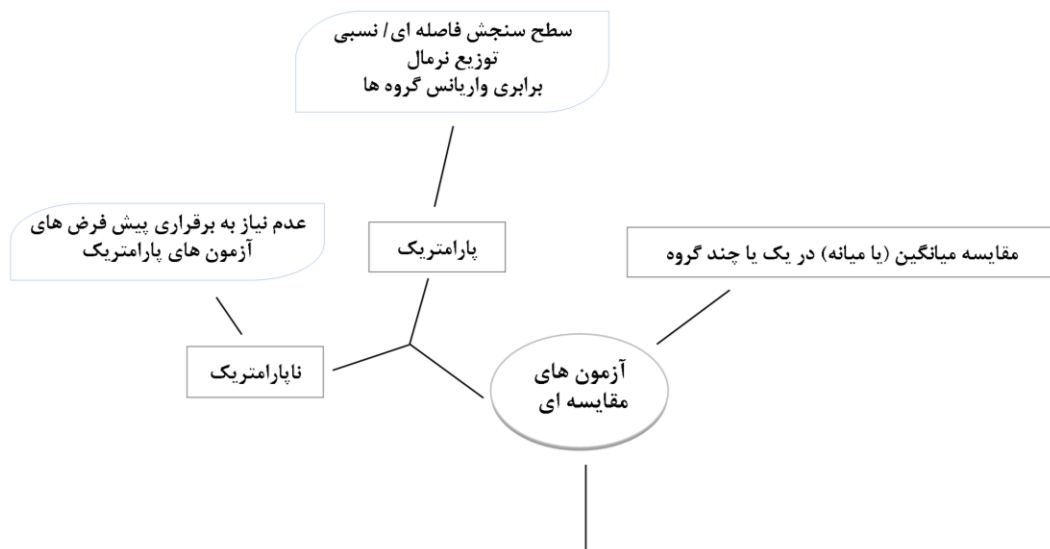
۲- آیا بین میزان اعتماد اجتماعی زنان و مردان تفاوت وجود دارد؟ اعتماد اجتماعی زنان بالاتر است یا مردان؟

۳- داده‌های مربوط به وزن ۱۵ نفر در ادامه آورده شده است. در این پژوهش ۱۵ نفر از افراد از رژیم غذایی کم‌کالری برای کاهش وزن خود استفاده کردند. وزن افراد در دو مقطع زمانی سنجیده شد: قبل از شروع رژیم غذایی و یک ماه بعد از شروع رژیم غذایی کم‌کالری. آیا رژیم غذایی کم‌کالری منجر به کاهش معنی‌دار وزن افراد شده است؟

شماره	۱۵	۱۴	۱۳	۱۲	۱۱	۱۰	۹	۸	۷	۶	۵	۴	۳	۲	۱
وزن اولیه	۸۰	۷۸	۸۸	۹۹	۷۰	۸۹	۶۸	۶۵	۷۵	۸۹	۸۵	۷۵	۸۷	۹۱	۹۴
وزن یک ماه بعد از شروع رژیم غذایی	۷۷	۷۵	۸۰	۸۹	۵۹	۸۸	۵۹	۶۶	۶۷	۸۸	۸۳	۷۵	۸۵	۹۲	۸۹

به جدول صفحه ۹۶ رجوع کنید:

۴- آیا بین میزان ورزش افراد با تحصیلات گوناگون تفاوت وجود دارد؟ افراد با چه میزان تحصیلات، بیشتر از سایرین ورزش می‌کنند؟



آزمون های ناپارامتریک	خصوصیت	آزمون های پارامتریک
	یک نمونه یا یک متغیر	
کای اسکوئر تک متغیره		t تک نمونه ای
	دو گروه	
ویلکاکسون	وابسته	t گروه های همبسته
مان-ویتنی	مستقل	t گروه های مستقل
	سه گروه	
فریدمن	وابسته	تحلیل اندازه های مکرر
کروسکال والیس	مستقل	تحلیل واریانس (آنووا)
نکته: در آزمون های مربوط به گروه های مستقل، لزومی به برابری تعداد اعضای گروه های مورد مقایسه نیست		

بخش سوم:

آزمون‌های چندمتغیره

به آزمون‌هایی چندمتغیره می‌گویند که نسبت به آزمون‌های تک متغیره و چندمتغیره، متغیرهای بیشتری داشته باشند، یعنی حداقل سه متغیر داشته باشند. علاوه بر این، برخی از زیر مجموعه‌های این متغیرها نیز باید با یکدیگر مورد تجزیه و تحلیل قرار بگیرند (که به نحوی با هم ترکیب می‌شوند). آزمون رگرسیون چندمتغیره، تحلیل مسیر، رگرسیون لجستیک و تحلیل عاملی اکتشافی را می‌توان از دسته آزمون‌های چندمتغیره به حساب آورد، چرا که عموماً حداقل سه متغیر در این تحلیل‌ها وجود دارد.

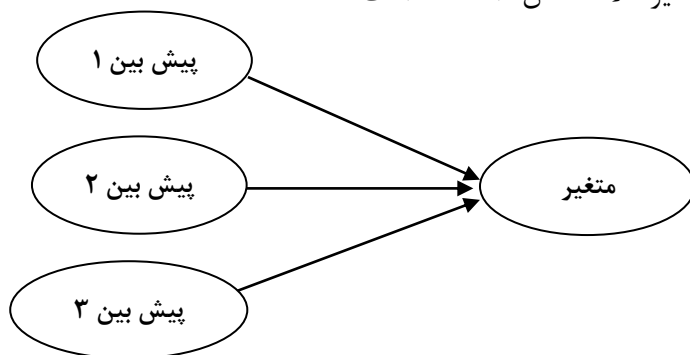
● رگرسیون چند متغیره

همبستگی‌ها می‌توانند ابزار پژوهشی بسیار مفیدی باشند، اما آن‌ها درباره قدرت پیش‌بینی متغیرها چیزی به ما نمی‌گویند. در تحلیل رگرسیون^{۴۰} یک مدل پیش‌بینی کننده را با داده‌های خودمان برازش می‌دهیم و از مدل و نتایج به دست آمده برای پیش‌بینی مقادیر یک متغیر وابسته از روی یک یا چند متغیر وابسته استفاده می‌کنیم (فیلد^{۴۱}، ۲۰۰۰: ۱۰۳).

تحلیل رگرسیون مرحله‌ای است بعد از همبستگی؛ تحلیل رگرسیون هنگامی استفاده می‌شود که بخواهیم مقادیر یک متغیر را از روی مقادیر متغیر (یا متغیرهای) دیگر پیش‌بینی کنیم. در این حالت، متغیری که ما از بهره می‌گیریم تا مقدار متغیر دیگر را پیش‌بینی کنیم متغیر پیش‌بین (مستقل) نام دارد و متغیری که می‌خواهیم پیش‌بینی کنیم را متغیر ملاک (وابسته) نام دارد. از طریق تحلیل رگرسیون به تبیین و تخمین تغییرات متغیر وابسته از طریق متغیر یا متغیرهای مستقل می‌پردازیم. در تحلیل رگرسیون، با اطلاع از نمره فرد در متغیرهای پیش‌بین، می‌توانیم نمره فرد در متغیر

^{۴۰} Regression Analysis^{۴۱} Field

ملاک را پیش‌بینی کرد. در تحلیل رگرسیون، چند متغیر پیش‌بین (حداقل دو متغیر)، مقادیر یک متغیر ملاک را پیش‌بینی کرده و بر متغیر ملاک اثر می‌گذارند. تحلیل رگرسیون چندمتغیری برای پیش‌بینی یک متغیر اندازه‌گیری شده به روش کمی که متغیر وابسته یا ملاک نامیده می‌شود، از روی مجموعه‌ای از متغیرهای مستقل یا پیش‌بین اندازه‌گیری شده به صورت دوارزشی یا کمی به کار می‌رود. این روش گسترش یافته رگرسیون خطی ساده است که فقط یک متغیر پیش‌بین و یک متغیر ملاک دارد. به هر یک از متغیرهای مستقل این مجموعه در مقایسه با پیش‌بینی کننده‌های دیگر وزن داده می‌شود تا یک متغیر مرکب خطی به دست آید و توان پیش‌بینی متغیر ملاک را به حداکثر رساند. اندازه متغیر محاسبه شده برابر با پیش‌بینی اندازه متغیر وابسته است و ترکیب وزنی را می‌توان به عنوان یک مدل ویژه برای پیش‌بینی متغیر ملاک تلقی کرد (میزر و دیگران، ۱۳۹۱: ۳۳-۳۲). در شکل ۵-۲ جایگاه و نحوه ارتباط متغیرهای پیش‌بین با متغیر ملاک نشان داده شده است.



شکل ۵-۲- نحوه ارتباط متغیرهای پیش‌بین با متغیر ملاک در رگرسیون چندمتغیره

پیش‌فرض‌ها

هنگامی می‌توانیم از آزمون رگرسیون استفاده کنیم که پیش‌شرط‌های زیر برقرار باشد:

۱. بین متغیرهای پیش‌بین و ملاک رابطه خطی وجود داشته باشد (بررسی نمودار پراکندگی دو متغیره).
۲. متغیر ملاک باید در سطح سنجش فاصله‌ای/نسبی باشد.
۳. متغیرهای پیش‌بین در سطح سنجش فاصله‌ای/نسبی و یا ترتیبی اندازه‌گیری شده باشند. متغیر پیش‌بین اسمی فقط هنگامی می‌تواند مورد استفاده قرار گیرد که به صورت دوبرخی باشد، یعنی بیش از دوطبقه نداشته باشد. مثلاً جنس با دوطبقه مذکر و مؤنث.

۴. رگرسیون چند متغیری نیازمند مشاهده‌های فراوانی است. تخمین R (همبستگی چندگانه) هم به تعداد متغیرهای پیش‌بین و هم به تعداد شرکت‌کنندگان (حجم نمونه) وابسته است. قانون کلی این است که تعداد شرکت‌کنندگان حداقل ۱۰ برابر تعداد متغیرهای پیش‌بین باشد (بریس و همکاران، ۱۳۹۱: ۳۳۸-۳۳۹).

اصطلاحات رگرسیون

بتا یا ضریب استاندارد شده رگرسیونی (β): بتا^{۴۲} مقیاسی است برای تعیین مقدار تأثیر متغیرهای پیش‌بین بر متغیر ملاک. بتا بر اساس واحد انحراف استاندارد اندازه‌گیری می‌شود (بریس و همکاران، ۱۳۹۱: ۳۳۹). به عنوان مثال مقدار بتای ۵/ نشان می‌دهد که تغییر یک واحد استاندارد در متغیر پیش‌بین، منجر به تغییر ۵/ انحراف استاندارد در متغیر ملاک می‌شود. بنابراین هرچه مقدار بتا بزرگتر باشد، اثر متغیر پیش‌بین بر متغیر ملاک بیشتر خواهد بود. برای مقایسه قدرت پیش‌بینی متغیرهای پیش‌بین از ضریب بتا استفاده می‌شود. هر متغیری که بتای بزرگتری داشته باشد تأثیر بیشتری بر متغیر ملاک دارد. دامنه این ضریب از (-۱) تا (+۱) است.

ضریب همبستگی چندگانه (R): همبستگی چندگانه^{۴۳}، همبستگی بین مقادیر مشاهده شده‌ی متغیر ملاک و مقادیر پیش‌بینی شده توسط مدل رگرسیونی است. بنابراین مقادیر بزرگ ضریب همبستگی چندگانه اشاره به همبستگی بالا و قوی بین مقادیر مشاهده شده و مقادیر پیش‌بینی شده متغیر ملاک دارد (فیلد^{۴۴}، ۲۰۰۰: ۱۰۳). دامنه این ضریب بین (۰) تا (۱) است. هرچه مقدار این ضریب به (۱) نزدیک‌تر باشد، نشان از همبستگی قوی‌تر و هرچه مقدار آن به (۰) نزدیک‌تر باشد، دلالت بر همبستگی ضعیف‌تر بین مقادیر مشاهده شده و پیش‌بینی شده دارد.

ضریب تعیین^{۴۵} (R^2): به مجذور ضریب همبستگی چندگانه^{۴۶} نیز معروف است. این ضریب، میزان واریانس متغیر وابسته که توسط مجموعه متغیرهای مستقل تبیین می‌شود را نشان می‌دهد. دامنه این ضریب بین (۰) تا (۱) است. هر چه مقدار این

^{۴۲} Beta

^{۴۳} Multiple Correlation

^{۴۴} Field

^{۴۵} Coefficient of Determination

^{۴۶} R Square

ضریب به ۱ نزدیک‌تر باشد، نشان از آن دارد که متغیرهای مستقل توانسته‌اند میزان بالاتری از واریانس متغیر وابسته را تبیین کنند (از ضریب تعیین جهت ارزیابی برازش مدل رگرسیونی استفاده می‌شود). همیشه مدلی مناسب‌تر است که از یک طرف مقدار ضریب تعیین آن بالاتر باشد و از طرف دیگر، متغیرهای زیادی را شامل نشود. چرا که هرچه تعداد متغیرها زیاد باشد، میزان برازش مدل بیش از اندازه بالا خواهد بود و در نتیجه تفسیر چنین مدلی با مشکل همراه است (حبیب‌پور و صفری، ۱۳۸۸: ۴۸۷).

ضریب تعیین تعدیل شده^{۴۷}: اشکال ضریب تعیین این است که میزان موفقیت مدل را بیش‌تر از اندازه برآورد می‌کند و کم‌تر تعداد متغیرهای مستقل و همین‌طور حجم نمونه را در نظر می‌گیرد. همچنین ضریب تعیین تعداد درجات آزادی را به حساب نمی‌آورد. از این رو، برخی آمارشناسان ترجیح می‌دهند تا شاخص دیگری به نام ضریب تعیین تعدیل شده را مورد استفاده قرار دهند. این ضریب، مقدار ضریب تعیین را به منظور انعکاس بیشتر میزان نیکویی برازش مدل تصحیح می‌کند. برای تفسیر ضریب تعیین، معمولاً از این مقدار (تعدیل شده) استفاده می‌شود. چون در این ضریب، مقدار ضریب تعیین با درجات آزادی تعدیل شده است (منصورفر، ۱۳۸۵: ۱۶۶).

علت در تحلیل رگرسیون

این که بتوانیم رابطه علت و معلولی را نتیجه‌گیری کنیم به این وابسته است که متغیرها دستکاری شده‌اند یا فقط اندازه‌گیری شده‌اند. در رگرسیون چند متغیره معمولاً نمراتی را که به‌طور طبیعی بر روی تعدادی متغیر پیش‌بین روی می‌دهند را اندازه می‌گیریم و سعی می‌کنیم تعیین کنیم کدام مجموعه از متغیرهای مشاهده شده بهترین پیش‌بین برای متغیر ملاک هستند. چنین پژوهش‌هایی علت را نمی‌رساند. اما اگر بتوان متغیرهای پیش‌بین یک رگرسیون چند متغیره را دستکاری کرد (روش آزمایشی)، می‌توان نتایج علی گرفت.

مثال

در این بخش می‌خواهیم تأثیر متغیرهای تعداد ساعات مطالعه، علاقه به درس آمار و نگرش SPSS در مقطع کارشناسی را بر نمره درس SPSS دانشجویان مقطع ارشد بسنجیم. در نتیجه ما سه متغیر پیش‌بین (مستقل) و یک متغیر ملاک (وابسته) خواهیم

⁴⁷ Adjusted R Square

داشت. می‌خواهیم بدانیم کدام یک از متغیرهای پیش‌بین می‌توانند به‌طور معنی‌داری متغیر ملاک را پیش‌بینی کنند و کدام یک پیش‌بینی کننده قوی‌تری هستند.

جدول ۵-۴- متغیرهای ملاک و پیش‌بین (مثال کتاب)

متغیرهای پیش‌بین	متغیر ملاک
ساعات مطالعه	
علاقه به آمار	نمره درس SPSS (در مقطع ارشد)
نمره درس SPSS در کارشناسی	

☑ نکته: تنها متغیرهایی را می‌توان به عنوان متغیر پیش‌بین وارد آزمون رگرسیون کرد که دارای رابطه خطی معنی‌دار با متغیر ملاک باشند. یعنی ابتدا باید به کمک نمودار پراکندگی، از وجود رابطه خطی بین دو متغیر پیش‌بین و ملاک اطمینان حاصل کرد و سپس متغیرهایی که دارای رابطه خطی با متغیر ملاک هستند را وارد رگرسیون نمود.

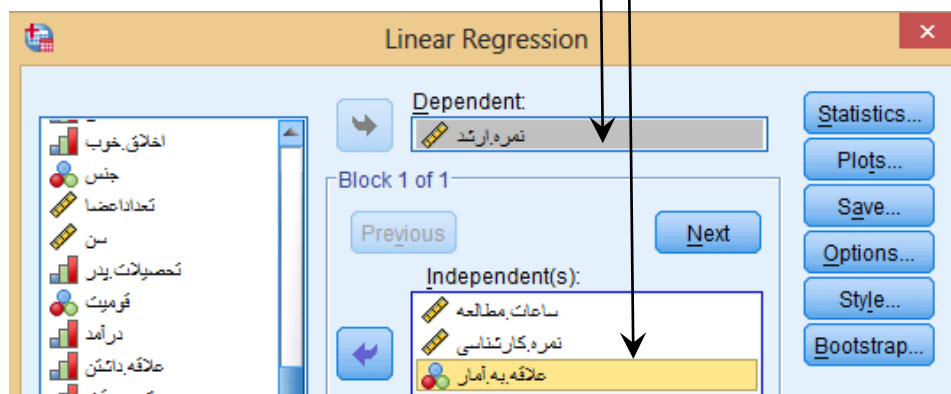
☑ نکته: یکی از مشکلات در اجرای آزمون رگرسیون، وجود متغیرهای اسمی است. جهت حل این مشکل باید متغیرهای اسمی (مانند جنس یا علاقه به آمار) را از حالت اسمی بودن خارج کرده و به صورت تصنّعی به متغیرهای فاصله‌ای تبدیل کنیم. برای این کار، باید هر کدام از طبقات (پاسخ‌های) یک متغیر به عنوان یک متغیر جداگانه با دو طبقه تعریف شود و به طبقه اول کد صفر (۰) و به طبقه دوم کد یک (۱) تعلق بگیرد. به عنوان مثال در تبدیل متغیر جنس (مرد یا زن) به زنان کد ۰ و به مردان کد ۱ می‌دهیم، در مورد متغیر پیش‌بین علاقه به آمار که دارای دو پاسخ بلی و خیر است ما اقدام به کدگذاری مجدد این متغیر کرده و به افرادی که پاسخ خیر داده‌اند کد صفر و به افرادی که پاسخ بله داده‌اند کد یک می‌دهیم و از متغیر علاقه به آمار با کد جدید به عنوان متغیر پیش‌بین استفاده می‌کنیم.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Regression ---> Linear

در کادر Dependent متغیر ملاک یا وابسته را وارد می کنیم، یعنی متغیر نمره ارشد.
در کادر Independent متغیر پیش بین یا مستقل را وارد می کنیم، یعنی متغیرهای ساعات مطالعه، نمره کارشناسی و علاقه به آمار. در انتها گزینه OK را انتخاب می کنیم.



نتایج:

جدول زیر بیان گر مقدار ضریب همبستگی چندگانه (R) و ضریب تعیین (R Square) است. مقدار ضریب همبستگی چندگانه برابر با ۰/۷۴۵ است. به دست آمده است. همان طور که ذکر شد مقدار R بین ۰ تا ۱ است و هرچه مقدار آن به ۱ نزدیک تر باشد نشان دهنده همبستگی قوی تر بین مقادیر واقعی و مقادیر مشاهده شده متغیر ملاک است.

مقدار ضریب تعیین تعدیل شده (Adjusted R Square) نیز برابر با ۰/۵۴ است. به دست آمده است و بدین معناست که متغیرهای پیش بین مدل توانسته اند ۵۴ درصد از واریانس متغیر ملاک را پیش بینی کنند و به عبارت دیگر ۵۴ درصد از سهم واریانس متغیر نمره ارشد ناشی از متغیرهای پیش بین (ساعات مطالعه، نمره کارشناسی و علاقه به آمار) است.

Model Summary

Model	R ضریب همبستگی چندگانه	R Square ضریب تعیین	Adjusted R Square ضریب تعیین تعدیل شده	Std. Error of the Estimate خطای استاندارد تخمین
1	.745 ^a	.556	.540	1.484

a. Predictors: (Constant), علاقه به آمار, ساعات مطالعه, نمره کارشناسی

جدول زیر نتیجه تحلیل واریانس را ارائه می‌دهد که معنی‌داری کل مدل را مورد ارزیابی قرار می‌دهد. از آنجا که سطح معنی‌داری کمتر از ۰.۰۵ است، مدل معنی‌دار است. معنی‌دار بودن آزمون تحلیل واریانس (مقدار F) نشان می‌دهد که متغیرهای پیش‌بین توانسته‌اند به طور معنی‌داری تغییرات متغیر ملاک را پیش‌بینی کنند.

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	242.50	3	80.83	36.67	.000^b
	Residual	194.00	88	2.20		
	Total	436.50	91			

a. Dependent Variable: نمره ارشد (متغیر ملاک)

b. Predictors: (Constant), ساعات مطالعه, علاقه به آمار (متغیرهای پیش‌بین)

جدول زیر جدول ضرایب و مهم‌ترین جدول در رگرسیون است. در این جدول ضریب رگرسیونی غیراستاندارد (B)، ضریب رگرسیونی استاندارد شده ($Beta$)، مقدار t و سطح معنی‌داری (Sig) گزارش شده است. مهم‌ترین بخش این جدول مربوط به مقادیر $Beta$ و سطح معنی‌داری است.

با مشاهده سطح معنی‌داری حاصل شده می‌توان گفت که چون سطح معنی‌داری دو متغیر ساعات مطالعه و علاقه به آمار کمتر از ۰.۰۵ به دست آمده است و تا سه رقم اعشار صفر شده است ($P < ۰.۰۰۱$)؛ در نتیجه می‌توان گفت که این دو متغیر بر متغیر نمره ارشد افراد تأثیرگذار هستند و می‌توانند تغییرات متغیر نمره ارشد افراد را پیش‌بینی کنند. با مقایسه ضریب استاندارد شده دو متغیر ساعات مطالعه (۰.۶۳۳) و علاقه به آمار (۰.۳۴۳) نتیجه می‌گیریم که متغیر ساعات مطالعه دارای تأثیر بیشتری بر نمره ارشد افراد است و پیش‌بینی کننده قوی‌تری است.

مقدار t و مقدار Sig معنی‌دار بودن آماری تأثیر متغیرهای پیش‌بین را نشان می‌دهند. مقدار t (چه مثبت و چه منفی) اگر بزرگتر از ۱.۹۶ باشد و مقدار Sig اگر کوچکتر از ۰.۰۵ باشد نشان می‌دهد که متغیر پیش‌بین بر متغیر ملاک تأثیر معنی‌دار دارد.

Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
(Constant)	13.166	1.369		9.62	.000
1 ساعات مطالعه	.352	.049	.633	7.11	.000
نمره کارشناسی	-.029	.092	-.028	-.31	.756
علاقه به آمار	1.512	.318	.343	4.75	.000

a. Dependent Variable: نمره ارشد

معادله رگرسیونی

ما درباره خط رگرسیون، برحسب یک معادله صحبت می‌کنیم. این معادله به عنوان مدل رگرسیون، یعنی ارائه رابطه پدیدبینی کننده بین دومتغیر عمل می‌کند. با استفاده از معادله رگرسیونی، می‌توانیم نمره هر فرد را در متغیر ملاک پدیدبینی کرده و محاسبه کنیم. این معادله به ما این امکان را می‌دهد که با آگاهی از نمرات فرد در متغیرهای پدیدبین، نمره او را در متغیر ملاک پدیدبینی کنیم و تخمین بزنیم.

معادله رگرسیونی برای نمره‌های خام بدین صورت است:

$$Y = a + b_1 x_1 + b_2 x_2 + \dots + b_n x_n$$

که در این معادله:

Y نمره پدیدبینی شده در مورد متغیر ملاک است.

X ها اندازه متغیرهای پدیدبین است که بر اساس آن‌ها پدیدبینی انجام می‌شود.

b ها ضرایب غیراستاندارد هستند.

a یک مقدار ثابت است که معرف عرض از مبدأ است.

معادله رگرسیونی مثال کتاب بدین صورت است:

$$\text{علاقه به آمار} = ۱.۵۱ + (.۳۵۲ \times \text{ساعات مطالعه فرد}) + ۱۳.۱۶ = \text{نمره ارشد}$$

برای به دست آوردن نمره ارشد هر فرد کافیه که میزان ساعات مطالعه آن فرد و نمره علاقه وی به درس آمار را در معادله قرار داده تا نمره ارشد وی را پدیدبینی کنیم. مثلاً

اگر فردی ۴ ساعت در هفته مطالعه می‌کند و نمره وی در متغیر علاقه به آمار عدد ۱ (بله یا علاقه‌مند به آمار) باشد، پیش‌بینی می‌کنیم که نمره ارشد وی ۱۶.۰۸ باشد.

$$۱۶.۰۸ = ۱.۵۱ (۱) + ۳.۵۲ (۴) + ۱۳.۱۶ = \text{نمره ارشد}$$

گزارش:

در گزارش نتایج می‌نویسیم:

با استفاده از روش Enter (همزمان)، مدل معنی‌داری به دست آمد ($R^2 = .۵۴۰$) ضریب‌تعیین تعدیل شده و $N=۱۰۰$ و $P < .۰۰۱$.

این مدل می‌تواند ۵۴ درصد از واریانس متغیر ملاک (نمره SPSS ارشد) را تبیین کند. دو متغیر ساعات مطالعه و علاقه‌مندی به درس آمار پیش‌بینی کننده معنی‌داری برای متغیر نمره درس SPSS هستند اما نمره کارشناسی پیش‌بینی کننده معنی‌داری نیست (یا تأثیر معنی‌داری ندارد). ضریب رگرسیونی استاندارد شده برای متغیر پیش‌بین ساعات مطالعه برابر با ۰.۶۳۳ و با جهت تأثیر و پیش‌بینی‌کنندگی مثبت و برای متغیر علاقه به آمار ۰.۳۴۳ و مثبت است. معادله رگرسیونی به دست آمده به صورت زیر است:

$$(\text{علاقه به آمار}) + ۱.۵۱ + (\text{ساعات مطالعه فرد}) ۳.۵۲ + ۱۳.۱۶ = \text{نمره ارشد}$$

تعریف

به جدول صفحه ۹۶ رجوع کنید:

۱- آیا سن، میزان ورزش و جنس افراد پیش‌بینی کننده‌های معنی‌داری برای میزان وزن هستند؟ جنس افراد را به صورت تصنعی به کار ببرید.

رگرسیون
چندمتغیره

کاربرد: پیش بینی مقادیر یک متغیر از روی مقادیر دو یا چندمتغیر
مرحله‌ای بعد از همبستگی دو متغیره

پیش فرض های مهم:

۱. وجود رابطه خطی بین متغیرهای پیش بین و متغیر ملاک
۲. سطح سنجش متغیر ملاک باید فاصله ای/نسبی و متغیر پیش بین باید فاصله ای/نسبی، ترتیبی و اسمی دو وجهی (با کد ۰ و ۱) باشد
۳. حجم نمونه حداقل ۱۰ برابر تعداد متغیرهای پیش بین

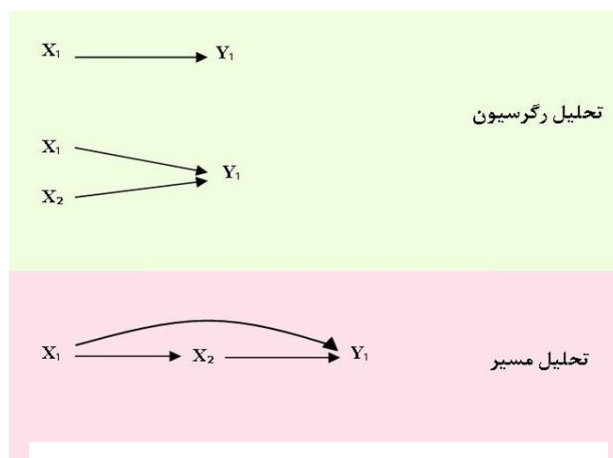
شاخص های مهم:

R^2 یا ضریب تعیین: ارزیابی برازش مدل
Beta یا ضریب استاندارد شده رگرسیونی: مقایسه تاثیر متغیرهای پیش بین بر متغیر ملاک
B یا ضریب غیراستاندارد: تدوین معادله رگرسیونی

● تحلیل مسیر

روش تحلیل مسیر^{۴۸} نخستین بار توسط یک بیولوژیست به نام سول رایت^{۴۹} در سال ۱۹۳۴ در ارتباط با تجزیه کردن مجموع همبستگی بین دو متغیر در یک سیستم علی معرفی گردید. این روش تعمیم یافته روش رگرسیون چند متغیره در ارتباط با تدوین مدل های علی است. تکنیک تحلیل مسیر بر اساس فرض ارتباط بین متغیرهای مستقل و وابسته استوار است. این روش بر استفاده ابتکاری از نمودار تصویری که به دیاگرام یا مدل مسیر معروف است، تاکید خاص دارد.

در روش تحلیل مسیر، علاوه بر تأثیر مستقیم، امکان شناسایی تأثیرات غیرمستقیم هریک از متغیرهای مستقل بر متغیر وابسته وجود دارد. به همین دلیل، در تحلیل مسیر، با چندین معادله خط رگرسیونی استاندارد شده مواجه هستیم، در حالی که در تحلیل رگرسیون، تنها یک معادله خط رگرسیونی استاندارد شده داریم. در واقع تفاوت تحلیل مسیر و رگرسیون در توانایی تحلیل مسیر در سنجش تأثیرات غیرمستقیم (علاوه بر تأثیرات مستقیم) است. این تفاوت در شکل زیر (شکل ۵-۳) نمایش داده شده است.



در تحلیل مسیر می توان اثرات غیرمستقیم را هم سنجید

شکل ۵-۳- مقایسه روش تحلیل مسیر و رگرسیون

⁴⁸ Path Analysis

⁴⁹ Swell Wright

✓ نکته: تحلیل مسیر برای تبیین دقیق‌تر روابط علی متغیرها، به تجزیه همبستگی بین متغیرها می‌پردازد. از طریق این تجزیه، اثرات مستقیم و غیرمستقیم یک متغیر بر متغیر دیگر مشخص می‌گردد. آن چه که ضرایب همبستگی نشان می‌دهند، اثرات مستقیم و غیرمستقیم متغیرهاست. به عبارت دیگر:

$$\text{اثرات غیرمستقیم} + \text{اثرات مستقیم} = \text{همبستگی}$$

پیش‌فرض‌ها

تحلیل مسیر بر مفروضاتی استوار است که مهم‌ترین آن‌ها عبارتند از:

- ۱- پیش‌فرض‌های آزمون رگرسیون خطی در تحلیل مسیر هم صدق می‌کند.
- ۲- روابط بین متغیرهای موجود در مدل، خطی، جمع‌پذیر و علی است و روابط منحنی و تعاملی لحاظ نمی‌شود.
- ۳- متغیرهای باقیمانده با همدیگر و با متغیرهایی که قبل از آن در مدل قرار گرفته‌اند، همبسته نیستند.
- ۴- جریان علیت یک طرفه می‌باشد و علیت متقابل بین متغیرها لحاظ نمی‌شود.

نمودار یا دیاگرام مسیر

مهم‌ترین بخش تحلیل مسیر، طراحی و آزمون نمودار مسیر است که از چند جزء اساسی تشکیل شده است. نمودار مسیر، در واقع، یک مدل ساختاری پیشین یا پیش تجربی با مجموعه معادله ساختاری است که روابط علی ممکن بین متغیرها را توصیف می‌کند. این نمودار مسیر، پس از مرور مبانی نظری و تدوین چارچوب نظری انتخابی پژوهش توسط پژوهش‌گر طراحی می‌شود که در نهایت در تحلیل مسیر مورد آزمون تجربی قرار می‌گیرد.

انواع متغیر در نمودار مسیر

- (۱) متغیرها در تحلیل مسیر شامل متغیرهای وابسته و متغیرهای مستقل است.

الف) متغیر وابسته: در تحلیل مسیر با دو دسته کلی از متغیر وابسته سروکار داریم:

- ۱- متغیر وابسته نهایی: متغیری که در نهایت تمامی تحلیل‌های مبنی بر تأثیرات متغیرهای مستقل باید روی آن انجام شود (در مثال کتاب، میزان تماشای تلویزیون). این متغیر فقط تأثیرپذیر است و بر هیچ متغیر دیگری در مدل اثر نمی‌گذارد.
- ۲- متغیر وابسته میانی: در واقع متغیری مستقل است که از یک طرف متغیرهای مستقل بر آن اثر می‌گذارند و از طرف دیگر، خود بر متغیر وابسته نهایی یا دیگر متغیرهای وابسته میانی اثر می‌گذارد و تأثیر سایر متغیرهای مستقل مدل بر روی آن آزمون می‌شود (در مثال کتاب متغیر درآمد).

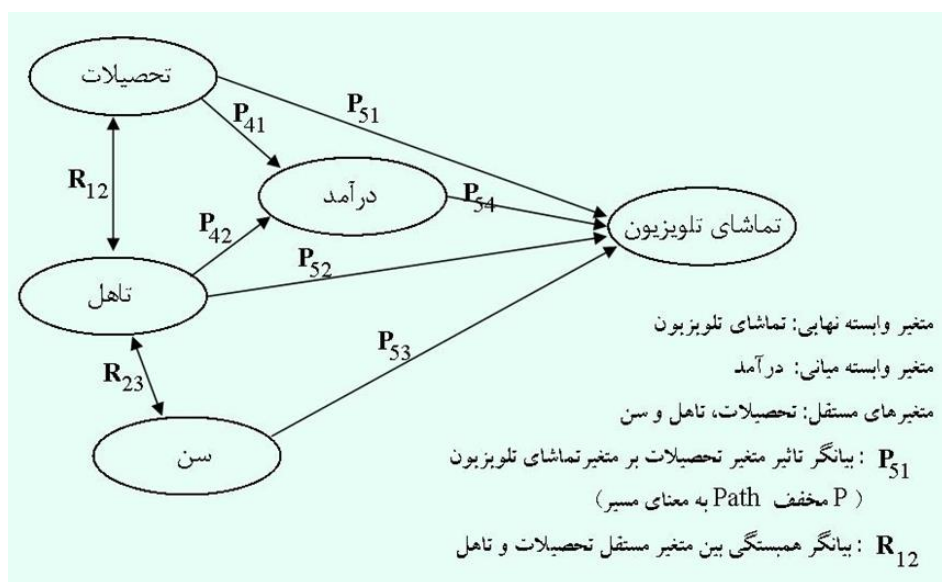
ب) متغیر مستقل: متغیری که بر متغیرهای وابسته (نهایی و میانی) تأثیر می‌گذارد و از هیچ متغیری تأثیر نمی‌پذیرد. در مثال کتاب سایر متغیرها که شامل سن، وضعیت تاهل و تحصیلات می‌شوند متغیر مستقل هستند (شکل ۵-۴).

☑ نکته: در هنگام آزمون مدل نظری، به بررسی میزان همبستگی دوجانبه بین متغیرها می‌پردازیم تا پی‌بیریم آیا متغیرهای مستقل با همدیگر همبستگی دارند یا خیر. این رابطه از نوع همبستگی است و نمی‌توان آن را علی فرض کرد. وجود همبستگی بالا بین متغیرهای مستقل نشان می‌دهد که متغیرهایی که برای این پژوهش انتخاب شده‌اند دارای پشتوانه نظری هستند و افزایش یا کاهش در هر یک از این متغیرها، با افزایش یا کاهش در متغیر مقایسه شده با آن همراه می‌باشد.

☑ نکته: در تحلیل مسیر، برای ارزیابی مدل از آماره R^2 (ضریب تعیین) استفاده می‌شود. این آماره مقدار واریانس متغیر وابسته را نشان می‌دهد که متغیرهای مستقل توانسته‌اند آن را تبیین کنند. در واقع، R^2 نشان‌دهنده این است که مدل تا چه اندازه با مجموعه‌ای از داده‌ها برازش دارد. بنابراین هرچه مقدار R^2 بالاتر باشد، مدل قوی‌تر است و برعکس مقدار پایین R^2 دلالت بر ضعف مدل دارد که باید مدل را اصلاح کرده و مدلی دیگری تدوین کنیم و متغیرهای مهم‌تر و تأثیرگذارتری انتخاب شوند که تبیین‌کننده واریانس بیشتری از متغیر وابسته نهایی باشند.

☑ نکته: متغیر وضعیت تاهل به صورت اسمی (مجرد=۱ و متاهل=۲) تعریف شده است. بنابراین پیش از اجرای رگرسیون باید متغیرهای اسمی را به طور تصنعی به فاصله‌ای تبدیل کرد. در نتیجه در SPSS متغیر دیگری تعریف کرده و جهت تبدیل متغیر اسمی وضعیت تاهل به متغیر فاصله‌ای، به کلیه افراد مجرد کد ۰ و به تمامی افراد متاهل کد ۱ می‌دهیم و این متغیر جدید را وارد تحلیل رگرسیون می‌کنیم. زمانی که می‌خواهیم متغیرهای اسمی دووجهی

را وارد رگرسیون کنیم باید کدگذاری این دسته از متغیرها به شکلی باشد که به یک طبقه کد ۰ (صفر) و به طبقه دیگر کد ۱ (یک) اختصاص یابد (به توضیحات بخش رگرسیون چندمتغیره مراجعه شود).



شکل ۵-۴- مدل مفهومی: مسیرها و نوع متغیرها در تحلیل مسیر (مثال کتاب)

آزمون ضرایب مسیر

در تحلیل مسیر یک مدل نظری را به آزمون می‌گذاریم که در نهایت با اجرای تحلیل، این مدل نظری (که از چارچوب نظری به دست می‌آید) باید به یک مدل تجربی (که از تحلیل مسیر به دست می‌آید) منتهی شود. بنابراین طبیعی است که همواره ساخت روابط علی بین متغیرها در مدل تجربی با مدل نظری فرق دارد. اما چه متغیرهایی باید در مدل باقی بمانند و چه متغیرهایی باید از مدل حذف شوند؟

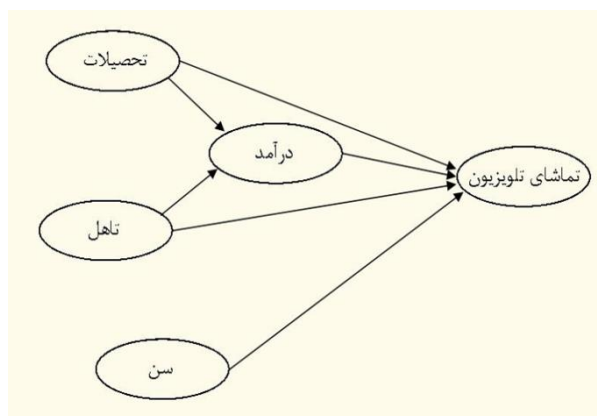
قاعده کلی این است که در اجرای تحلیل مسیر، متغیرهایی که مقدار بتا یا ضریب مسیر آنها در سطح اطمینان ۹۵ درصد معنی‌دار نشود ($P > 0.05$) را از مدل حذف کنیم. بنابراین، مسیر این متغیرها نیز از مدل حذف می‌شود. البته هر چه حجم نمونه بزرگتر باشد، احتمال این که ضرایب مسیر معنی‌دار شوند، بیشتر است.

مثال

پژوهشی با هدف شناخت عوامل مؤثر بر میزان تماشای تلویزیون در شهر کرج انجام شد (پژوهش فرضی است). در این پژوهش به بررسی نقش متغیرهای فردی و جمعیتی شناختی بر میزان تماشای تلویزیون در افراد پرداخته شد.

اجرا:

در فرآیند اجرای این پژوهش، پس از مطالعه مبانی نظری و پیشینه تجربی مدل نظری زیر را تنظیم کرده‌ایم (البته این نمودار مسیر و داده‌های مورد استفاده در آن ساختگی است و تنها برای آموزش مراحل و نحوه اجرای تحلیل مسیر به کار گرفته شده است). این مدل عوامل مؤثر بر میزان تماشای تلویزیون را نشان می‌دهد. عوامل مورد نظر ما عبارت‌اند از: تحصیلات، وضعیت تاهل، سن و متغیر درآمد. مدل نظری و مفهومی به صورت زیر است (شکل ۵-۵).



شکل ۵-۵- مدل نظری عوامل مؤثر بر تماشای تلویزیون (مثال کتاب)

تحلیل مسیر از مجموع چند آزمون رگرسیون خطی ساده و چندمتغیره تشکیل شده است. در روش تحلیل مسیر ما به تعداد متغیرهای وابسته (هم میانی و هم نهایی) باید از روش رگرسیون استفاده کنیم. در مثال کتاب، ما با یک متغیر وابسته میانی (درآمد) و یک متغیر وابسته نهایی (تماشای تلویزیون) روبرو هستیم. در نتیجه باید دو آزمون رگرسیون چندمتغیره اجرا کنیم که در یک آزمون، متغیر تماشای تلویزیون به عنوان متغیر وابسته وارد تحلیل رگرسیون می‌کنیم و تأثیر متغیرهای مستقل را بر آن می‌سنجیم

و سپس یک آزمون رگرسیون چندمتغیره دیگر اجرا می‌کنیم که این بار متغیر درآمد نقش متغیر وابسته را ایفا می‌کند و تأثیر متغیرهای مستقل را بر آن می‌سنجیم. در جدول زیر تعداد آزمون‌های رگرسیون خطی و متغیرهای وابسته و مستقل در هر آزمون نشان داده شده‌اند. جدول زیر از مدل نظری پژوهش استخراج شده است (شکل ۵-۵).

جدول ۵-۵- تعداد آزمون‌های رگرسیون و متغیرهای هر آزمون (مثال کتاب)

تعداد آزمون	متغیر وابسته	متغیرهای مستقل
آزمون ۱	تماشای تلویزیون	درآمد، تحصیلات، وضعیت
آزمون ۲	درآمد	تحصیلات و وضعیت تاهل

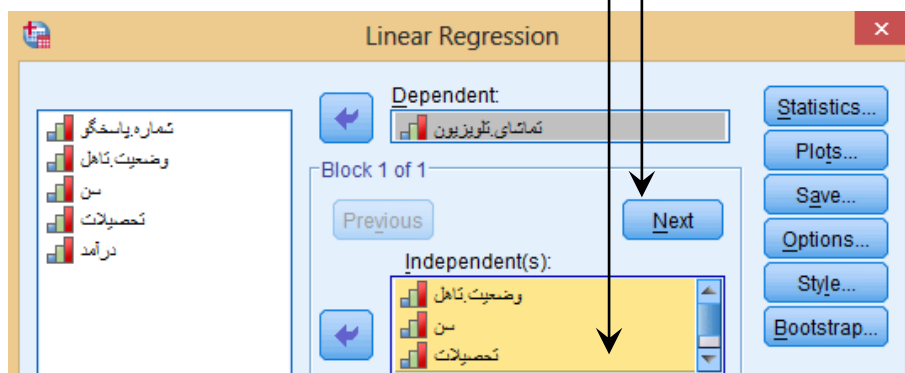
❑ نکته: ترتیب انجام مدل‌های رگرسیونی از متغیر وابسته نهایی به متغیر وابسته میانی است. یعنی ابتدا آزمون رگرسیونی که در آن متغیر وابسته نهایی وجود دارد را اجرا می‌کنیم و سپس به سراغ آزمون رگرسیونی می‌رویم که متغیرهای وابسته میانی در آن قرار دارند.

(۱) آزمون مدل اول (متغیر وابسته: میزان تماشای تلویزیون)

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Regression ---> Linear

متغیر وابسته میزان تماشای تلویزیون را وارد کادر Dependent و تمامی متغیرهای مستقل را وارد کادر Independent می‌کنیم. در انتها OK را انتخاب می‌کنیم.



تا اینجا میزان بتای مسیرهای P_{51} و P_{52} و P_{53} و P_{54} به دست می‌آید اما هنوز بتای سایر مسیرها به دست نیامده است.

نتایج:

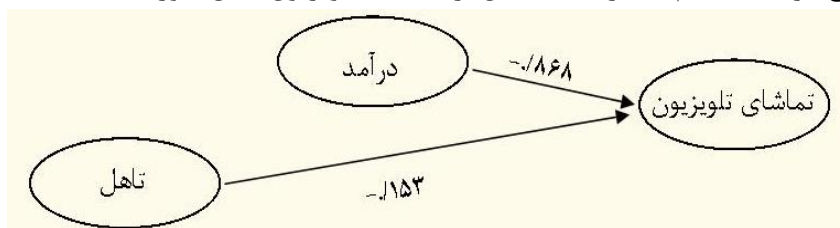
جهت بررسی نتایج به دست آمده به جدول ضرایب (Coefficients) مراجعه می‌کنیم و سطح معنی‌داری هر متغیر را بررسی می‌کنیم. اگر سطح معنی‌داری هر متغیر مستقل یا پیش‌بین کمتر از مقدار ۰/۰۵ باشد نتیجه می‌گیریم که تأثیر آن متغیر بر میزان تماشای تلویزیون تأیید می‌شود و آن متغیر می‌تواند به طور معنی‌داری تغییرات میزان تماشای تلویزیون را پیش‌بینی کند. مطابق نتایج به دست آمده، دو متغیر وضعیت تاهل و درآمد بر میزان تماشای تلویزیون تأثیرگذار هستند ولی دو متغیر سن و تحصیلات بر میزان تماشای تلویزیون مؤثر نیستند.

جدول شماره ۱
Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	4.597	.498		9.22	.000
وضعیت	-.365	.079	-.153	-4.64	.000
سن	.004	.002	.019	.62	.533
تحصیلات	.008	.001	.003	.09	.927
درآمد	-.014	.002	-.868	-24.81	.000

a. Dependent Variable: تماشای تلویزیون

مدل تجربی اثرات مستقیم متغیرهای مستقل بر تماشای تلویزیون بدین صورت است:



شکل ۵-۶- مدل تجربی تأثیرات مستقیم متغیرهای مستقل بر میزان تماشای تلویزیون (مثال)

از میان چهار متغیر مستقل که بر مبنای مدل نظری دارای تأثیرات مستقیم بر میزان تماشای تلویزیون بودند، تنها دو متغیر درآمد و تاهل دارای تأثیر مستقیم معنی‌داری بر میزان تماشای تلویزیون هستند و دو متغیر دیگر پژوهش (تحصیلات و سن) تأثیر

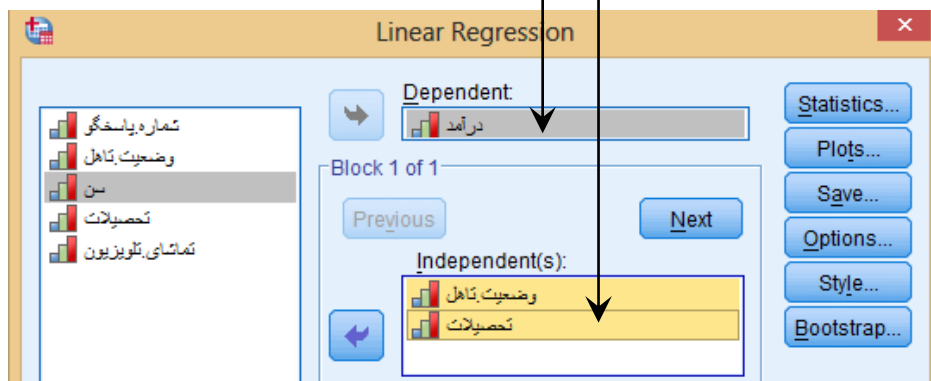
معنی‌داری بر میزان تماشای تلویزیون ندارند ($P > .05$) و پیش‌بینی‌کننده تغییرات آن نیستند.

۲) آزمون مدل دوم (متغیر وابسته: درآمد)

در این مرحله متغیر درآمد را به عنوان متغیر وابسته در نظر می‌گیریم و تأثیر دو متغیر تحصیلات و تاهل را بر روی آن آزمون می‌کنیم (می‌خواهیم بتای مسیرهای P_{41} و P_{42} را به دست بیاوریم).

نکته: متغیر درآمد در مرحله دوم نقش یک متغیر وابسته (میانجی) را دارد و می‌توان تأثیر متغیرهای مستقل را بر روی آن سنجید.

متغیر وابسته یعنی درآمد را وارد کادر Dependent و تمامی متغیرهای مستقل را وارد کادر Independent می‌کنیم. در انتها OK را انتخاب می‌کنیم.



نتایج:

جدول زیر ضرایب مسیر را نشان می‌دهد که معنی‌داری روابط و شدت و جهت ضرایب مسیر در آن ارائه شده است. نتایج نشان می‌دهد که تحصیلات و تاهل پیش‌بینی‌کننده معنی‌داری برای درآمد افراد هستند و جهت تأثیر هر دو متغیر مثبت است.

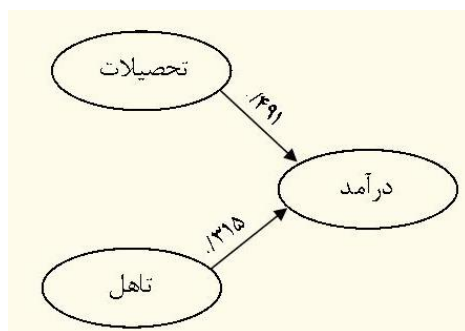
جدول شماره ۲

Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	277.57	38.43		7.22	.000
تحصیلات	28.85	4.32	.491	6.67	.000
وضعیت تاهل	164.68	38.48	.315	4.28	.000

a. Dependent Variable: درآمد

مطابق نتایج به دست آمده از جدول شماره دو، تأثیر هر دو متغیر مستقل بر متغیر وابسته درآمد معنی دار است. در شکل ۵-۷ روابط معنی دار به همراه ضرایب استاندارد شده گزارش شده است.



شکل ۵-۷- مدل تجربی تأثیرات مستقیم متغیرهای مستقل بر درآمد

هدف از تحلیل رگرسیون برآورد ضرایب رگرسیونی (بتا) مسیرهای نمودار مسیر است. در نتیجه ما باید در پی شناسایی ضرایب رگرسیونی باشیم. و مدل نظری خود را با مدل تجربی به دست آمده مقایسه کنیم.

تأثیر متغیرهای مستقل بر متغیر وابسته نهایی به ۳ صورت است:

- (۱) صرفاً به صورت مستقیم
- (۲) صرفاً به صورت غیرمستقیم
- (۳) هم به صورت مستقیم و هم غیرمستقیم

(۱) متغیرهایی که صرفاً به صورت مستقیم بر میزان تماشای تلویزیون تأثیر گذاشته‌اند:

متغیر سن و درآمد متغیرهایی هستند که تنها به صورت مستقیم بر متغیر میزان تماشای تلویزیون تأثیر دارند. نتایج جدول شماره ۱ نشان می‌دهد که متغیر درآمد دارای تأثیر معنی‌داری بر متغیر وابسته است اما متغیر سن دارای تأثیر معنی‌داری بر متغیر وابسته نیست ($P > 0.05$). در نتیجه مسیر P_{53} را از مدل نهایی حذف می‌کنیم.

(۲) متغیرهایی که صرفاً به صورت غیرمستقیم بر میزان تماشای تلویزیون تأثیر گذاشته‌اند:

در مثال کتاب متغیری که تنها به صورت غیرمستقیم بر متغیر وابسته نهایی تأثیر گذاشته باشد و تأثیر مستقیمی بر متغیر وابسته نهایی نداشته باشد، وجود ندارد.

(۳) متغیرهایی که دارای تأثیر دوگانه (هم مستقیم و هم غیر مستقیم) بر متغیر میزان تماشای تلویزیون هستند:

دو متغیر تحصیلات و وضعیت تاهل، متغیرهایی هستند که دارای تأثیر دوگانه بر متغیر وابسته هستند. بر طبق نتایج جدول ۱ از بین این دو متغیر، وضعیت تاهل دارای تأثیر معنی‌داری بر میزان تماشای تلویزیون است و میزان ضریب مسیر این متغیر برابر با 0.153 - می‌باشد. اما متغیر تحصیلات تأثیر معنی‌داری بر متغیر وابسته ندارد ($P > 0.05$). در نتیجه مسیر شماره P_{51} را از مدل حذف می‌کنیم.

در ادامه جهت تعیین مقادیر مسیرهای P_{41} و P_{42} از نتایج جدول ۲ بهره می‌گیریم. بر طبق نتایج، متغیرهای تحصیلات و وضعیت تاهل هر دو دارای تأثیر معنی‌داری بر درآمد هستند ($P < 0.05$) و در نتیجه هر دو مسیر در مدل نهایی باقی می‌مانند.

گزارش:

نحوه محاسبه ضرایب مسیر مستقیم، غیرمستقیم و کل:

بعد از این که روش تحلیل مسیر را طی چندین مرحله اجرا کردیم، باید در قالب جدولی، نتایج مربوط به انواع تأثیرات مستقیم، غیرمستقیم و کل را نشان دهیم. جدولی مانند بعد ترسیم می‌کنیم:

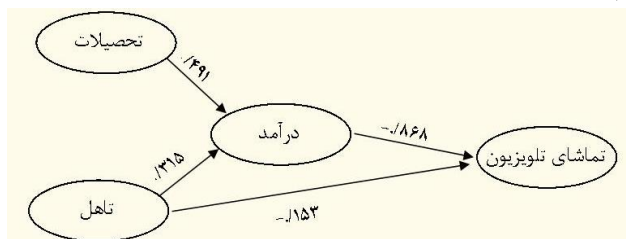
جدول ۵-۶- اثرات مستقیم، غیرمستقیم و کل متغیرهای پیش‌بین بر متغیر ملاک (تماشای تلویزیون)

انواع تأثیر			متغیرهای مستقل
کل (مجموع مستقیم و غیرمستقیم)	غیر مستقیم	مستقیم	
-/۴۳	-/۴۳	—	تحصیلات
-/۴۲	-/۲۷	-/۱۵۳	وضعیت تاهل
—	—	—	سن
-/۸۶۸	—	/۸۶۷	درآمد

اثرات غیرمستقیم هر متغیر برابر با حاصل ضرب مسیرهای آن متغیر است. به عنوان مثال اثر غیرمستقیم متغیر تحصیلات بر تماشای تلویزیون برابر با حاصل ضرب مسیرهای P_{41} و P_{54} است. یعنی حاصل ضرب ضرایب مسیر تحصیلات بر روی درآمد و درآمد بر روی تماشای تلویزیون $(-/۴۳ = -/۸۶۸ \times /۴۹۱)$. همین‌طور تأثیر غیرمستقیم متغیر وضعیت تاهل بر تماشای تلویزیون برابر است با حاصل ضرب ضرایب مسیر تاهل بر روی درآمد و درآمد بر تماشای تلویزیون $(-/۲۷ = -/۸۶۸ \times /۳۱۵)$.

تأثیر کل هر متغیر از جمع کردن اثرات مستقیم و غیرمستقیم آن متغیر به دست می‌آید. در مورد متغیر تاهل اثر کل برابر شده است با $-/۴۲$ که مجموع اثرات مستقیم و غیر مستقیم این متغیر است $(-/۴۲ = -/۲۷ + -/۱۵۳)$.

مدل تجربی به دست آمده از تکنیک تحلیل مسیر در شکل شماره ۵-۸ نشان داده شده است. همان‌طور که ملاحظه می‌کنیم در مدل تجربی به دست آمده دو مسیر P_{51} و P_{53} از مدل (دیاگرام) حذف شده‌اند.

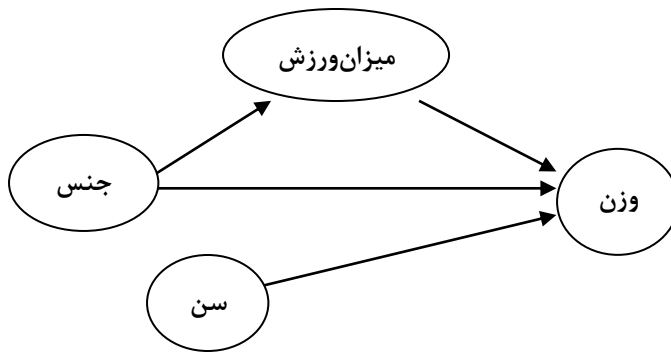


شکل ۵-۸- مدل تجربی تأثیرات مستقیم و غیر مستقیم متغیرهای پیش‌بین بر میزان تماشای تلویزیون

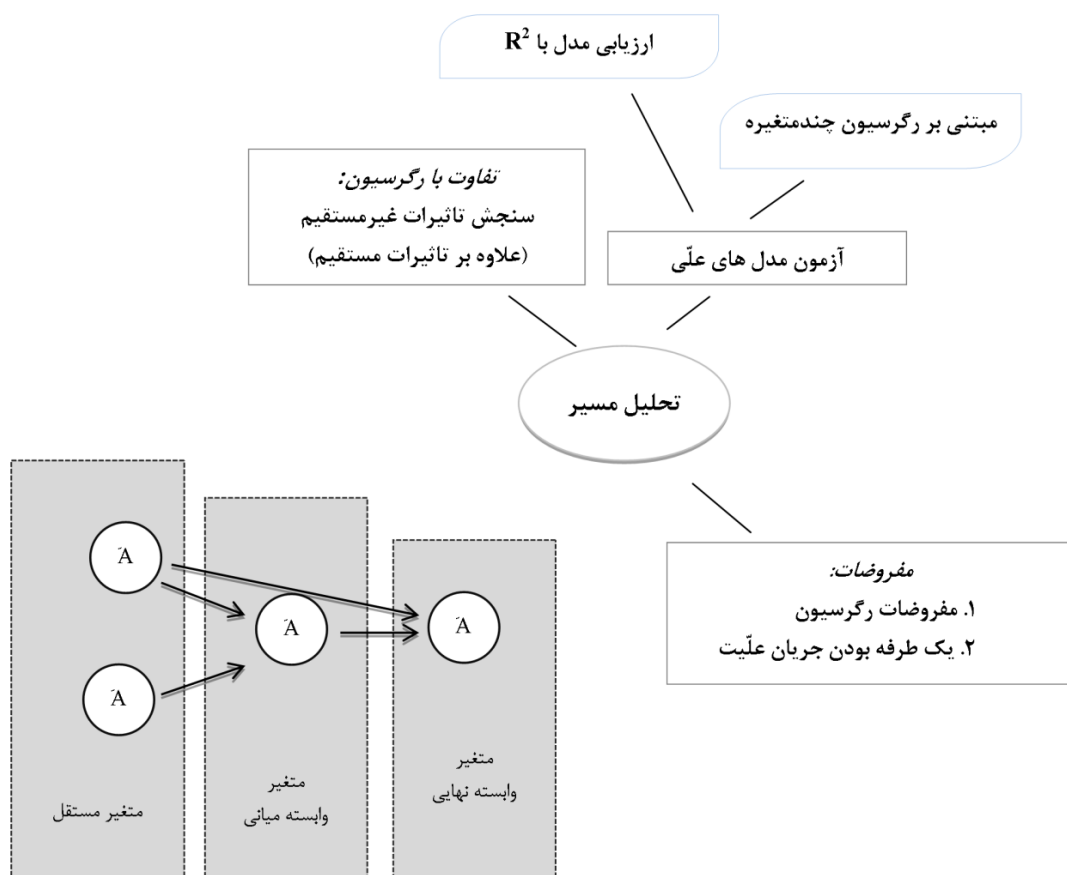
تعریف

به جدول و داده‌های صفحه ۹۶ رجوع کنید:

۱- مدل مفهومی زیر را آزمون کنید. ضریب بتای تمامی چهار مسیر را روی شکل بنویسید. مدل مفهومی را با مدل تجربی مقایسه کنید. تاثیر کدام متغیرها بر وزن افراد صرفاً مستقیم و کدام متغیرها هم مستقیم و هم غیرمستقیم است؟ جدول تاثیرات مسیرهای مستقیم و غیر مستقیم و همچنین اثر کل را ارائه کنید.



☑ نکته: متغیر جنس را به صورت تصنعی به فاصله‌ای تبدیل کنید و به زنان کد ۰ و به مردان کد ۱ بدهید. در حالت کلی بهتر است به گروهی که در متغیر وابسته (در اینجا وزن) میانگین بالاتری دارد کد ۱ و به گروهی که میانگین پایین‌تری دارد کد ۰ تعلق بگیرد (یعنی ابتدا میانگین زنان و مردان را در متغیر وزن به طور جداگانه به دست می‌آوریم و هر گروهی که میانگین وزن بالاتری داشته باشد کد ۱ می‌گیرد. در مثال کتاب میانگین وزن مردان بیشتر از زنان است و در نتیجه زنان کد ۰ و مردان کد ۱ می‌گیرند).



رگرسیون لجستیک

رگرسیون لجستیک^{۵۰} و تحلیل تشخیصی^{۵۱} روندهای آماری هستند که برای پیش‌بینی عضویت گروهی به کار می‌روند نه پیش‌بینی نمره. هر دو روند به ما امکان می‌دهند تا یک متغیر وابسته طبقه‌ای را بر اساس تعدادی متغیر پیش‌بین یا مستقل پیش‌بینی کنیم. این متغیرهای مستقل معمولاً پیوسته هستند، اما رگرسیون لجستیک از عهده متغیرهای مستقل طبقه‌ای برمی‌آید. به طور کلی رگرسیون لجستیک در مقایسه با تحلیل تشخیصی می‌تواند در گستره وسیع‌تری از موقعیت‌ها به کار رود. تفسیر نتایج رگرسیون لجستیک نیز آسان‌تر از تفسیر نتایج تحلیل تشخیصی است (بریس و همکاران، ۳۹۴:۱۳۹۱).

زمانی که بخواهیم نمره افراد (میزان درآمد) افراد را برحسب سابقه شغلی و میزان تحصیلاتی که دارند پیش‌بینی کنیم از روش رگرسیون چند متغیره استفاده می‌کنیم، اما زمانی که بخواهیم پیش‌بینی کنیم که هر فرد برحسب سابقه شغلی و میزان تحصیلاتی که دارد در کدام گروه درآمدی (درآمد بالا یا درآمد پایین) قرار می‌گیرد، از روش رگرسیون لجستیک بهره می‌گیریم. از رگرسیون لجستیک برای پیش‌بینی عضویت طبقه برای مواردی که از قبل عضویت در آن مشخص نیست استفاده می‌کنیم.

در تعریفی دیگر، زمانی که متغیر وابسته در سطح اسمی است و متغیرهای مستقل، هم اسمی، ترتیبی و فاصله‌ای/نسبی هستند، روش‌های رگرسیون خطی معمولی و تحلیل تشخیصی، مقدار برآوردها را کمتر از مقدار واقعی نشان می‌دهند، در این وضعیت از رگرسیون لجستیک استفاده می‌شود.

رگرسیون لجستیک بر تحلیل تشخیصی برتری دارد و مهم‌ترین دلیلش این است که در تحلیل تشخیصی گاهی اوقات احتمال وقوع یک پدیده خارج از طیف ۰ تا ۱ قرار می‌گیرد و متغیرهای پیش‌بین نیز باید دارای توزیع نرمال چندمتغیره باشند. در حالی که در رگرسیون لجستیک، احتمال وقوع یک پدیده در داخل ۰ تا ۱ قرار دارد و رعایت

⁵⁰ Logistic Regression

⁵¹ Discriminant Analysis

پیش فرض نرمال بودن متغیرهای پیش بین لازم نیست (سرمد، ۱۳۸۴: ۳۳). در مجموع زمانی از رگرسیون لجستیک استفاده می کنیم که شرایط زیر برقرار باشد:

- (۱) متغیر وابسته اسمی (دو یا چندوجهی) باشد.
- (۲) متغیرهای مستقل هم ترتیبی (یا اسمی) و هم فاصله ای/نسبی باشند.

زمانی که پیش فرض های نرمال بودن چندمتغیره و برابری ماتریس های واریانس و کوواریانس تامین شدند، و متغیرهای مستقل همه در سطح سنجش فاصله ای/نسبی باشند؛ در آن صورت پیشنهاد می شود که از روش تحلیل تشخیصی به جای روش رگرسیون لجستیک استفاده کنیم. در جدول ۵-۷ به مقایسه انواع رگرسیون پرداخته ایم.

جدول ۵-۷- مقایسه روش های رگرسیون لجستیک، رگرسیون خطی و تحلیل تشخیصی

نوع روش	سطح سنجش متغیر وابسته	سطح سنجش متغیرهای مستقل	پیش فرض توزیع نرمال
رگرسیون لجستیک	اسمی	همه سطوح سنجش	(ندارد)
تحلیل تشخیصی	اسمی	فاصله ای/نسبی	نرمال بودن چندمتغیره
رگرسیون خطی	فاصله ای/نسبی	فاصله ای/نسبی و ترتیبی	نرمال بودن توزیع باقیمانده ها

دو نوع رگرسیون لجستیک وجود دارد: رگرسیون لجستیک اسمی دو وجهی و رگرسیون لجستیک اسمی چندوجهی. یعنی بسته به این که متغیر وابسته دارای چندوجه (طبقه) باشد، از رگرسیون لجستیک مناسب با آن بهره می گیریم. در این کتاب به آموزش رگرسیون لجستیک دو وجهی پرداخته می شود. در رگرسیون لجستیک دو وجهی، متغیر وابسته دو طبقه دارد، مانند گروه درآمدی بالا و گروه درآمدی پایین یا گروه شاغلان و گروه غیرشاغلان.

پیش فرض ها و ملاحظات

۱- سطح سنجش: متغیرهای مستقل می توانند هم در سطح کمی (فاصله ای/نسبی) و هم در سطح کیفی (اسمی/ترتیبی) طبقه بندی شده باشند. اما چنان چه یک یا چندمتغیر

مستقل در سطح اسمی/ترتیبی بودند، حتماً باید ابتدا این متغیرها را به متغیرهای تصنعی تبدیل کنیم (یعنی کدهای ۰ و ۱). البته در روش رگرسیون لجستیک، کادری به نام ... Categorical وجود دارد که با انتخاب و اجرای آن، متغیرهای اسمی و ترتیبی به‌طور خودکار به متغیرهای تصنعی تبدیل می‌شوند. بنابراین، نیازی به کدگذاری مجدد آن‌ها نیست.

۲- توزیع نرمال: لزوم تبعیت متغیرهای مستقل از توزیع نرمال ضروری نیست. اما چنان چه این متغیرها دارای توزیع نرمال چندمتغیره باشند، در آن صورت برازش مدل بهتر خواهد بود.

۳- هم خطی چندگانه: از دیگر مفروضات رگرسیون لجستیک نبود هم‌خطی چندگانه بین متغیرهای مستقل است. چرا که در صورت هم‌خطی چندگانه بودن این متغیرها، برآوردها دارای اریب بوده و خطاهای استاندارد نوسان زیادی خواهند داشت.

۴- حجم نمونه: اگرچه در ادبیات مربوط به رگرسیون لجستیک، قواعد خاصی برای حجم نمونه پیشنهاد نشده است، اما برخی نویسندگان در حوزه آمار چندمتغیره، حداقل حجم نمونه برای یک تحلیل رگرسیون لجستیک خوب را ۱۰۰ نفر و برخی نیز ۵۰ نفر عنوان کرده‌اند (حبیب‌پور و صفری، ۱۳۸۸: ۷۱۶-۷۱۵).

مثال

در پژوهشی اطلاعات مربوط به سابقه شغلی و میزان تحصیلات کارمندان اداره پست شهر اصفهان را جمع‌آوری کرده‌ایم (داده‌ها فرضی است) و قصد داریم این مسأله را بررسی کنیم که آیا می‌توان با استفاده از سابقه شغلی و تحصیلات کارمندان، پیش‌بینی کنیم که آنان در چه گروه درآمدی (درآمد بالا یا درآمد پایین) قرار می‌گیرند. به بیان دیگر قصد داریم عواملی را که بر نوع گروه درآمدی کارمندان مؤثر هستند را شناسایی کنیم و با استفاده از یک سری متغیرها (سابقه شغلی و تحصیلات)، پیش‌بینی کنیم که آنان در کدام گروه درآمدی قرار می‌گیرند. نحوه سنجش و کدگذاری متغیرهای مستقل و وابسته در جدول ۵-۸ آمده است.

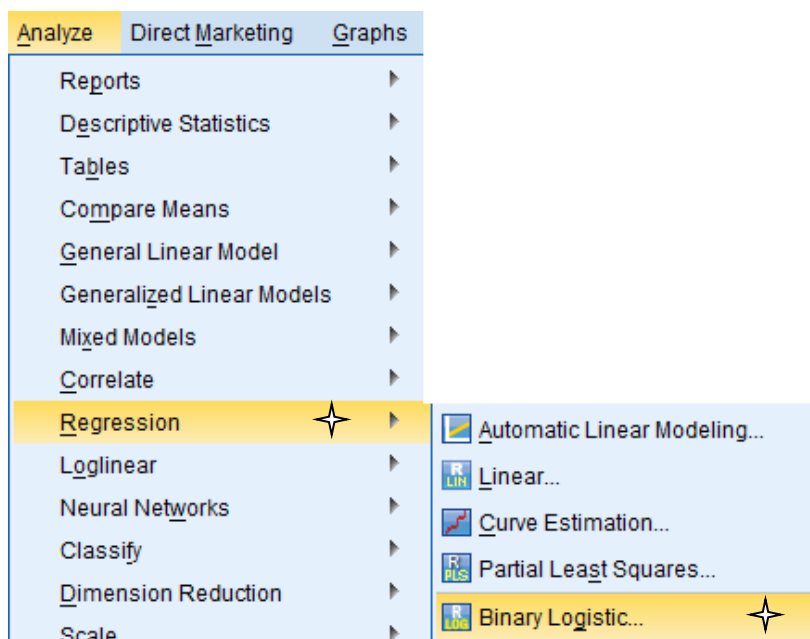
جدول ۵-۸- متغیرهای مستقل و وابسته (مثال کتاب)

متغیرهای مستقل:	سابقه شغلی: (به سال) میزان تحصیلات : شامل چهار طبقه: زیر دیپلم: کد ۱ دیپلم: کد ۲ فوق دیپلم و لیسانس: کد ۳ فوق لیسانس و دکترا: کد ۴
متغیر وابسته:	میزان درآمد: شامل دو طبقه درآمد پایین: کد ۱ درآمد بالا: کد ۲

اجرا:

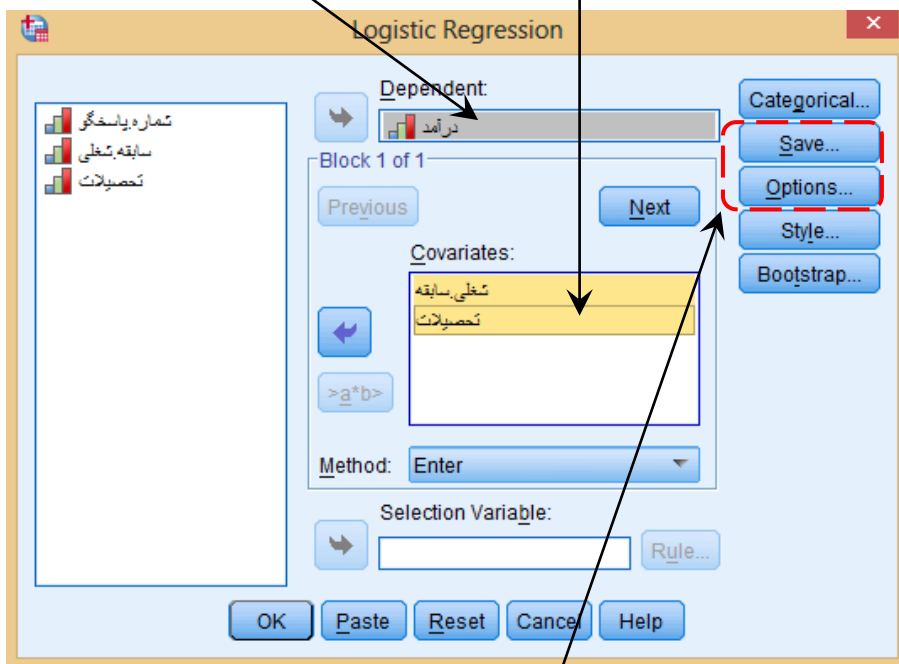
مسیر زیر را دنبال می‌کنیم:

Analyze ---> Regression ---> Binary Logistic



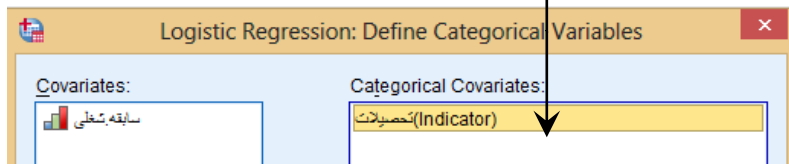
۱- متغیر درآمد را وارد کادر متغیر وابسته می‌کنیم.

۲- متغیرهای تحصیلات و سابقه شغلی را وارد کادر متغیرهای مستقل کمّی (Covariates) می‌کنیم.

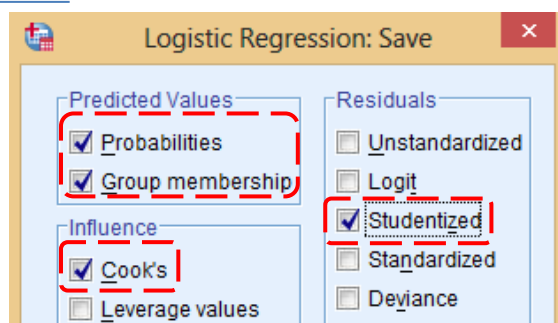


۳- گزینه Categorical را انتخاب می‌کنیم و تنظیمات آن را مطابق شکل بعد انجام می‌دهیم.
در مراحل بعد گزینه‌های Save و Options را می‌زنیم.

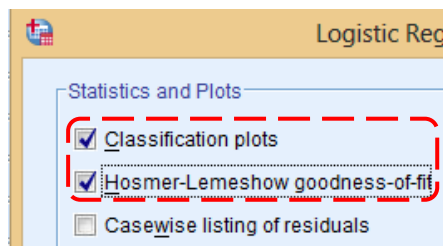
در پنجره متغیر طبقه‌ای (Categorical)، متغیرهای طبقه‌ای (اسمی یا ترتیبی) که در اینجا میزان تحصیلات است را وارد می‌کنیم.
گزینه Continue را انتخاب می‌کنیم.



در پنجره ذخیره کردن
(Save) گزینه‌های زیر را
مطابق شکل فعال می‌کنیم:
Probabilities
Group membership
Cook's
Studentized



در پنجره گزینه‌ها (Options) گزینه‌های
زیر را مطابق شکل فعال می‌کنیم:
Classification plots
Hosmer-Lemeshow goodness...
در انتها بر روی گزینه Continue و در
کادر اصلی بر روی Ok کلیک می‌کنیم.



نتایج:

خروجی‌های آزمون رگرسیون لجستیک در ادامه آورده شده است. در ادامه به توضیح جداول مهم می‌پردازیم و در انتها به تفسیر و نحوه گزارش نتایج خواهیم پرداخت. جدول زیر نشان می‌دهد که حجم نمونه معتبر ۱۱۰ نفر است و داده ناقص یا گمشده‌ای وجود ندارد. به بیان دیگر ۱۰۰ درصد پاسخگویان که برابر با ۱۱۰ نفر می‌شوند در پردازش و تحلیل شرکت داده شده‌اند. لازم به یادآوری است که حداقل حجم نمونه برای تحلیل رگرسیون لجستیک را ۱۰۰ مورد و برخی نویسندگان ۵۰ مورد می‌دانند (حبیب پور و صفری، ۱۳۸۸: ۷۱۶-۷۱۵).

Case Processing Summary

Unweighted Cases ^a		N	Percent
Included in		110	100
Selected Cases	Missing Cases	0	.0
	Total	110	100
Unselected Cases	Missing Cases	0	.0
	Total	110	100

کدگذاری مجدد متغیر وابسته: هنگامی که دستور رگرسیون لجستیک را اجرا می‌کنیم، برنامه SPSS به طور خودکار به کسانی که درآمد پایین دارند و کد ۱ گرفته بودند کد ۰ می‌دهد و به کسانی که درآمد بالا دارند و کد ۲ گرفته بودند در آزمون رگرسیون لجستیک کد ۱ می‌دهد و تحلیل‌ها را بر اساس کدهای صفر و یک انجام می‌دهد. این طبقه‌بندی به صورت خودکار انجام می‌شود.

Dependent Variable Encoding

Original Value	Internal Value
درآمد پایین	0
درآمد بالا	1

کدگذاری مجدد متغیر مستقل طبقه ای (تحصیلات): در روش رگرسیون لجستیک باید متغیر تحصیلات که یک متغیر ترتیبی است به شیوه‌ای کدگذاری مجدد کرد. آزمون رگرسیون لجستیک هر طبقه را با طبقه بالاتر از آن مقایسه می‌کند. به بیان دیگر چون میزان تحصیلات طبقه‌ای بوده و شامل چهار طبقه می‌شود در نتیجه به آن ۳ کد اختصاص داده است. شیوه‌کار بدین صورت است که در مقایسه طبقه اول و دوم با هم، برای طبقه اول کد ۱ و برای طبقه دوم کد ۰ را در نظر گرفته است. سپس در مقایسه طبقه دوم و سوم با هم، برای طبقه دوم کد ۱ و برای طبقه سوم کد ۰ را در نظر می‌گیرد و الی آخر. برای طبقه آخر کدی در نظر گرفته نمی‌شود. مقایسه متغیرهای ترتیبی به صورت زوجی یا دوه‌دو انجام می‌شود.

لازم به ذکر است که ما در فایل داده‌های SPSS متغیرهای تحصیلات و درآمد را کدگذاری کرده‌ایم. این کدگذاری موجب می‌شود که در فایل خروجی رگرسیون لجستیک به جای عدد، نام‌های طبقات درآمد و تحصیلات گزارش شود و فهم جداول آسان‌تر شود.

فراوانی هر کدام از طبقات میزان تحصیلات در جدول ذکر شده است. برای مثال افرادی که میزان تحصیلاتشان زیر دیپلم است ۱۹ نفر هستند.

Categorical Variables Codings

	Frequency	Parameter coding		
		(1)	(2)	(3)
زیر دیپلم	19	1	.000	.000
دیپلم	31	.000	1	.000
تحصیلات فوق دیپلم و لیسانس	46	.000	.000	1
فوق لیسانس و دکترا	14	.000	.000	.000

نتایج جدول زیر نشان می‌دهد که با اطمینان ۵۷.۳ درصد، با استفاده از مجموع ۲ متغیر مستقل یعنی تحصیلات و سابقه شغلی، قادریم تغییرات متغیر وابسته گروه درآمدی (بالا و پایین) را تبیین کنیم.

Classification Table^a

Observed		Predicted (پیش بینی شده)		
		درآمد		Percentage Correct درصد صحیح
		درآمد پایین	درآمد بالا	
Step 0	درآمد پایین	0	47	.0
	درآمد بالا	0	63	100
Overall Percentage (درصد کلی)				57.3

a. The cut value is .50

ارزیابی کل مدل

نتایج آزمون آم نیبوس، ارزیابی کل مدل رگرسیونی لجستیک را نشان می‌دهد. این آزمون به بررسی این موضوع می‌پردازد که مدل (نقش تحصیلات و سابقه شغلی در طبقه‌بندی گروه درآمدی) تا چه اندازه قدرت تبیین و کارایی دارد. مجذور کای برابر با ۲۹.۲۶ است که در سطح خطای کمتر از ۰.۰۵، معنی‌دار است و نشان از برازش مدل دارد و همچنین نشان می‌دهد که متغیرهای مستقل از توانایی لازم در پیش‌بینی عضویت افراد در گروه‌های درآمدی برخوردارند.

Omnibus Tests of Model Coefficients

		Chi-square	df	Sig.
Step 1	Step	29.26	4	.000
	Block	29.26	4	.000
	Model	29.26	4	.000

بررسی برازش مدل

ضرایب جدول بعد یعنی (Cox & Snell R Square و Nagelkerke R Square)، تقریب‌های (معادل‌های) ضریب تعیین (R^2) در رگرسیون خطی هستند که در این‌جا در رگرسیون لجستیک استفاده می‌شوند. در رگرسیون لجستیک، چون محاسبه دقیق مقدار ضریب تعیین دشوار است، بنابراین از مقادیر آماره‌های فوق برای این کار استفاده می‌شود تا مشخص گردد که متغیرهای مستقل چه میزان از واریانس متغیر وابسته را می‌توانند تبیین کنند.

مقادیر بین ۰ تا ۱ نوسان دارد. مقادیر دو آماره برابر با ۰.۲۳۴ و ۰.۳۱۴، به دست آمده است و بدین معناست که دو متغیر مستقل توانسته‌اند بین ۲۳ تا ۳۱ درصد از تغییرات متغیر میزان درآمد پاسخگویان (درآمد بالا یا پایین داشتن) را تبیین کنند. لازم به ذکر است که در آزمون رگرسیون لجستیک، ضریب تعیین نه به صورت یک عدد معین، بلکه به صورت یک طیف از حداقل تا حداکثر میزان ضریب تعیین گزارش می‌شود.

Model Summary

Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	120.89 ^a	.234	.314

جهت تبیین قدرت مدل در تفکیک افراد در طبقات متغیر وابسته، از نتایج جدول بعد استفاده می‌کنیم. از نتایج این جدول می‌توانیم به میزان صحت و درستی در طبقه‌بندی افراد پی ببریم. دقت کل در طبقه‌بندی افراد برابر با ۷۳.۶ درصد بوده است. این دقت در افراد با درآمد پایین برابر با ۷۰.۲ درصد است که نشان می‌دهد ۳۳ نفر از کسانی که درآمد پایین داشته‌اند درست تفکیک شده‌اند (واقعا هم درآمد پایین داشته‌اند) و ۱۴ نفر

به اشتباه تفکیک شده‌اند (درآمد پایین داشته‌اند اما در پیش‌بینی عضویت گروهی به اشتباه در گروه با درآمد بالا قرار گرفته‌اند). همچنین دقت طبقه‌بندی در کارمندان با درآمد بالا ۷۶.۲ درصد است که در این گروه، ۴۸ نفر از کسانی که درآمد بالا داشته‌اند به درستی تفکیک شده‌اند و ۱۵ نفر به اشتباه تفکیک شده‌اند. هر چه افراد بیشتری به درستی تفکیک شوند نشان از دقت بالاتر فرآیند پیش‌بینی عضویت گروهی افراد دارد.

Classification Table^a

Observed		Predicted		
		درآمد		Percentage Correct
		پایین	بالا	
Step 1	پایین درآمد	33	14	70.2
	بالا	15	48	76.2
Overall Percentage				73.6

a. The cut value is .500

جدول بعد مهم‌ترین جدول در تفسیر نتایج مربوط به معنی‌داری و میزان تأثیر هر متغیر مستقل بر متغیر وابسته است. در این جدول:

B همان ضریب رگرسیونی استاندارد نشده است.

S.E همان خطای استاندارد است.

Wald: آماره والد، مهم‌ترین آماره برای آزمون معنی‌داری حضور هر متغیر مستقل در مدل می‌باشد. آماره والد معادل آماره t در رگرسیون خطی است. نتایج نشان می‌دهد که تأثیر متغیر میزان تحصیلات در مجموع معنی‌دار است اما تأثیر متغیر سابقه شغلی بر میزان درآمد معنی‌دار نیست ($P=0.158$).

Exp: معادل ضریب رگرسیونی استاندارد شده در رگرسیون خطی است که برای تفسیر نتایج از آن استفاده می‌شود.

جدول ضرایب

Variables in the Equation

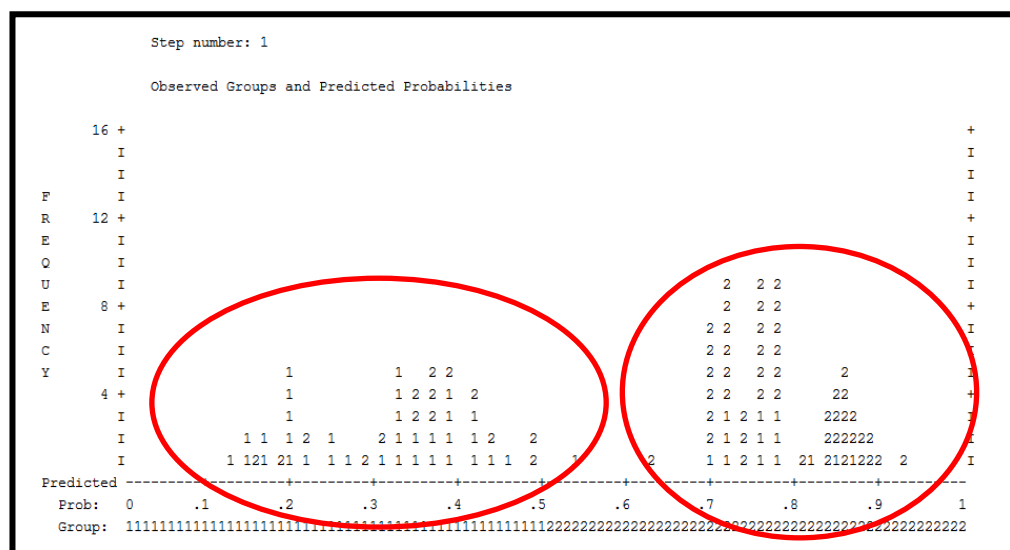
	B	S.E.	Wald	df	Sig.	Exp(B)
تحصیلات			22.82	3	.000	
(1)تحصیلات	-3.29	.968	11.54	1	.001	.037
(2)تحصیلات	-2.44	.866	7.96	1	.005	.087
Step 1 ^a (3)تحصیلات	-.77	.846	.83	1	.362	.463
سابقه شغلی	.10	.069	1.99	1	.158	1.102
Constant (مقدار ثابت)	1.21	.864	1.96	1	.161	3.352

a. Variable(s) entered on step 1: تحصیلات, سابقه شغلی.

شکل ۵-۹ به نمودار طبقه‌بندی (Classification Plot) معروف است و تصویر بصری از صحت طبقه‌بندی را در یک نمودار هیستوگرام پشته‌ای نشان می‌دهد. این نمودار در کنار جدول طبقه‌بندی (ClassificationTable) روش دیگری برای نمایش صحت طبقه‌بندی پاسخگویان است.

این نمودار از دو محور X و Y تشکیل شده است و میزان انطباق احتمالات پیش‌بینی شده با پیامدها را نشان می‌دهد. هر علامت ۱ نشان‌دهنده پاسخگویی است که درآمد پایین دارد و هر علامت ۲ نشان‌دهنده پاسخگویی است که درآمد بالا دارد.

هرچه پاسخگویان یک گروه در یک سمت نمودار و پاسخگویان گروه دیگر در سمت دیگر نمودار قرار بگیرند، صحت طبقه‌بندی و پیش‌بینی مدل بالاتر است. دقیق‌ترین شیوه طبقه‌بندی به‌صورتی است که همه اعداد ۱ (درآمد پایین) در یک سمت نمودار و همه اعداد ۲ (درآمد بالا) در سمت دیگر نمودار قرار بگیرند. بررسی نمودار نشان می‌دهد که در سمت چپ نمودار، عدد ۱ یا افرادی که درآمد پایین دارند، فراوانی بیشتری از عدد ۲ دارند و در سمت راست نمودار عدد ۲ بسیار بیشتر از عدد ۱ به چشم می‌خورد. اگر در یک سمت نمودار فقط عدد ۱ و در سمت دیگر فقط عدد ۲ وجود داشته باشد، دقت گروه‌بندی و پیش‌بینی افراد کامل و صددرصد است.



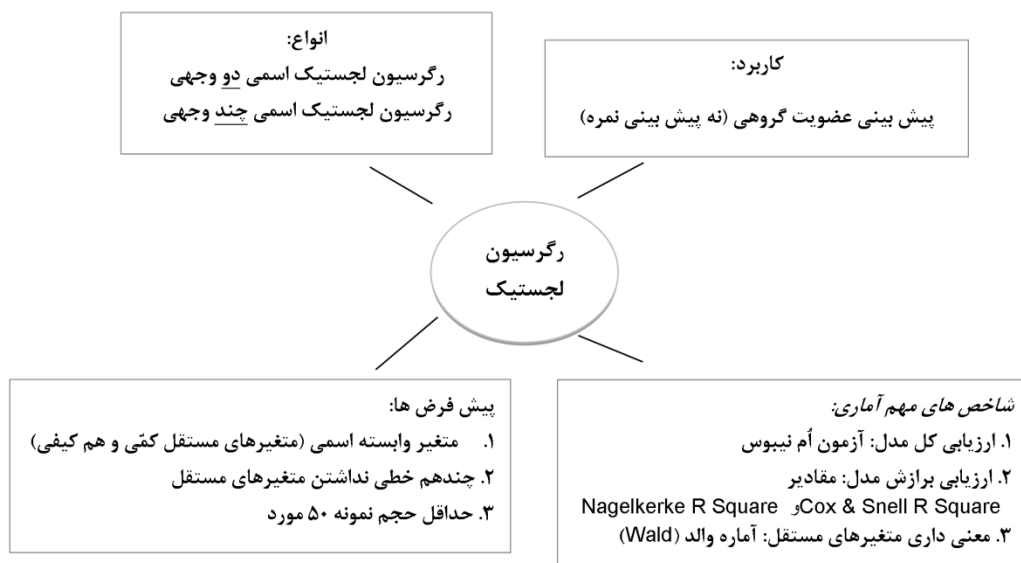
شکل ۵-۹- نمودار طبقه‌بندی گروه درآمدی افراد بر حسب متغیرهای مستقل (مثال)

گزارش:

آزمون تحلیل رگرسیون لجستیک با هدف پیش‌بینی عضویت گروهی انجام شد که در آن، متغیر گروه درآمدی (درآمد بالا و درآمد پایین) به عنوان متغیر وابسته و میزان تحصیلات و سابقه شغلی به عنوان متغیرهای پیش‌بین (مستقل) انتخاب شدند. در کل ۱۱۰ نفر در تحلیل وارد شدند و مدل معنی‌دار بود ($P < .001$ ، $df = 4$ ، $\chi^2 = 29.265$ مجذور کای).

این مدل بین ۲۳ تا ۳۲ درصد از واریانس میزان درآمد (بالا یا پایین) را تبیین می‌کند. ۷۰.۲ درصد از افراد با درآمد پایین درست طبقه‌بندی شده‌اند و ۷۶.۲ درصد از پیش‌بینی‌ها در مورد افراد با درآمد بالا صحیح بود. در کل ۷۳.۶ درصد از پیش‌بینی‌ها درست بود.

{ جدول شماره... (جدول ضرایب)}، ضرایب و آماره Wald و درجات آزادی و مقادیر احتمال برای هر کدام از متغیرهای پیش‌بین را ارائه می‌دهد. نتایج نشان می‌دهد که فقط میزان تحصیلات به طور معنی‌داری توانسته است عضویت افراد را در گروه‌های درآمدی پیش‌بینی کند. جهت این تأثیر منفی است که نتایج نشان می‌دهد که با افزایش میزان تحصیلات کارمندان، احتمال قرارگرفتن آنان در گروه درآمدی پایین هم افزایش می‌یابد (داده‌ها فرضی است)!! نتایج نشان داد که سابقه شغلی افراد پیش‌بینی‌کننده معنی‌داری برای عضویت افراد در گروه‌های درآمدی نیست ($P > 0.5$).



تعریف

به جدول صفحه ۹۶ رجوع کنید:

۱- افراد را برحسب میزان تحصیلات در دو دسته تقسیم بندی و کدگذاری مجدد کنید: افراد با تحصیلات دیپلم و پایین تر و افراد با تحصیلات دانشگاهی. سپس بیازمایید که آیا جنس و سن افراد پیش بینی کننده های خوبی برای عضویت افراد در دو گروه مربوط به تحصیلات هستند؟

بخش چهارم:

آزمون‌های ناپارامتریک

آزمون‌های آماری به دو دسته پارامتریک و ناپارامتریک تقسیم می‌شوند. در آزمون‌های پارامتریک پیش‌فرض‌هایی در ارتباط با جمعیت آماری و داده‌هایی که از جمعیت به دست می‌آید وجود دارد و در این آزمون‌ها پژوهش‌گر می‌تواند پارامترهای جمعیت را برآورد کند. آمار ناپارامتریک یا آزمون‌های توزیع-آزاد^{۵۲} آزمون‌هایی هستند که بر تخمین پارامترها تاکید ندارند و در مورد توزیع متغیرها پیش‌فرضی را مطرح نمی‌کنند. در آزمون‌های ناپارامتریک در مورد جمعیتی که نمونه از آن به دست می‌آید پیش‌فرض-هایی (پیش‌شرط‌هایی) مطرح نمی‌شود. بدین معنا که:

در آزمون‌های پارامتریک حداقل سطح سنجش متغیر باید فاصله‌ای باشد. اما اکثر آزمون‌های ناپارامتریک برای داده‌هایی در سطح سنجش ترتیبی به کار می‌روند و برخی از آزمون‌ها هم برای داده‌هایی در سطح سنجش اسمی.

برای استفاده از آزمون‌های پارامتریک این مفروضات باید برقرار باشد: داده‌ها از جمعیتی با توزیع نرمال به دست آمده باشد (و در نتیجه توزیع داده‌ها در نمونه هم نرمال باشد)، نمونه‌هایی که از جمعیت به دست می‌آیند باید واریانس برابر داشته باشند و پاسخگویان (آزمودنی‌ها) باید مستقل از یکدیگر انتخاب شده باشند. اما در آزمون‌های ناپارامتریک نیازی به برقرار بودن این پیش‌فرض‌ها نیست.

☑ نکته: سطح سنجش (اندازه‌گیری)، از عناصر مهم در تصمیم‌گیری برای انتخاب آزمون‌های پارامتریک یا ناپارامتریک است. اعتقاد بر این است که آزمون‌های پارامتریک برای داده‌های دارای سطوح فاصله‌ای/نسبی استفاده شوند، اما نشان داده شده است که استفاده از تکنیک‌های پارامتریک برای داده‌های ترتیبی به ندرت موجب انحراف در نتایج می‌شود (مونرو، ۱۳۸۹: ۱۳۴). یعنی در مورد متغیرهایی که از ترکیب چند گویه یا سوال ترتیبی حاصل شده‌اند می‌توان (با کمی تسامح و تساهل و در صورت برقرار بودن سایر فرض‌های آزمون‌های پارامتریک) از

⁵² Distribution Free

آزمون‌های پارامتریک استفاده کرد. تقسیم‌بندی آزمون‌های پارامتریک و ناپارامتریک و پیش‌فرض‌های این آزمون‌ها در جدول ۵-۹ گزارش شده است.

جدول ۵-۹- پیش‌فرض‌ها و آزمون‌های مرتبط با آزمون‌های پارامتریک و ناپارامتریک		
نوع آزمون	پیش فرض ها	آمارها و آزمون‌های مرتبط
پارامتریک	سطح سنجش فاصله ای	میانگین و انحراف استاندارد (فراوانی، مد و میانه)
	یا نسبی	آزمون همبستگی پیرسون، آزمون
	توزیع نرمال متغیرها	رگرسیون
	برابری واریانس گروه ها	آزمون‌های مقایسه ای: آزمون های t
ناپارامتریک	استقلال نمونه ها از یکدیگر	، آزمون تحلیل واریانس، آزمون اندازه های مکرر، آزمون آنکووا
		فراوانی، مد و میانه
		ضرایب پیوند: مجذور کای، فی و وی کرامر
		سطح سنجش ترتیبی یا اسمی

در این بخش برخی از آزمون‌های مهم ناپارامتریک آموزش داده می‌شود که شامل آزمون مجذور کای تک‌متغیره، آزمون‌های مقایسه دو گروه مستقل (آزمون مان-ویتنی)، دو گروه همبسته (ویلکاکسون)، چندگروه مستقل (کروسکال-والیس) و چندگروه همبسته (فریدمن) می‌شود.

☑ نکته: در آزمون‌های مقایسه‌ای از نوع پارامتریک، بر تفاوت میانگین تاکید می‌شود و در آزمون‌های مقایسه‌ای از نوع ناپارامتریک، بر تفاوت میانه تاکید می‌شود.

سه دسته آزمون مقایسه‌ای وجود دارد:

آزمون‌های تک نمونه‌ای که یک گروه از نمونه را آزمون می‌کند: مانند آزمون دو جمله‌ای، کولموگروف اسمیرنوف تک نمونه‌ای و مجذور کای تک‌متغیره.

آزمون‌های نمونه‌های وابسته که دو یا چند وضعیت را در نمونه‌های یکسان مقایسه می‌کند: مانند آزمون تی همبسته، اندازه‌های مکرر، ویلکاکسون و فریدمن.

آزمون‌های نمونه‌های مستقل که یک وضعیت را در دو یا چند گروه متفاوت مقایسه می‌کند: مانند آزمون تی مستقل، تحلیل واریانس (آنووا)، مان-ویتنی و کروسکال والیس.

☑ نکته: تمامی آزمون‌های ناپارامتریک (به‌غیر از ضرایب همبستگی) در منوی زیر قابل دسترس هستند: Analyze ---> Nonparametric Tests

آزمون تک‌متغیری مجذور کای

به گونه اختصاصی، آزمون مجذور کای (آزمون χ^2 دو یا کای اسکوئر) برای ارزیابی چهار نوع فرضیه به کار می‌رود: (۱) برازندگی یک توزیع نمونه با توزیع نظری مورد انتظار، (۲) رابطه بین دو یا چند متغیر طبقه‌ای، (۳) مقایسه گروه‌ها بر پایه سطوح مختلف طبقه‌ها یا مقوله‌های اندازه‌گیری شده و (۴) برازندگی مدل با داده‌ها (عسگری و هومن، ۱۳۹۲: ۳). در این بخش به آموزش آزمون مجذور کای در ارتباط با برازندگی یک توزیع نمونه با توزیع نظری مورد انتظار می‌پردازیم.

آزمون تک‌متغیری مجذور کای، که گاه آزمون برازندگی مجذور کای نامیده می‌شود، یکی از مهم‌ترین آزمون‌های ناپارامتری است که در پژوهش‌های علوم رفتاری، به‌ویژه در تعیین برازندگی و الگوی داده‌های تجربی و مقایسه آن‌ها با مدل‌های نظری، موارد استفاده بسیار دارد. آزمون برازندگی مجذور کای برای آزمایش فرضیه‌هایی که پژوهش‌گر قصد دارد بر پایه آن‌ها درباره انطباق یا عدم انطباق توزیع فراوانی مشاهده شده با توزیع نظری تصمیم بگیرد از اهمیت ویژه‌ای برخوردار است. (عسگری و هومن، ۱۳۹۲: ۲۷). آزمون مجذور کای تک‌متغیره (یک بعدی) برای فرضیه‌هایی به کار می‌رود که در آن پژوهش‌گر از یک متغیر ترتیبی برای تنظیم فرضیه استفاده کرده است. هدف اصلی این آزمون، مقایسه فراوانی‌های مشاهده شده با فراوانی‌های مورد انتظار، به‌ویژه از طریق مقادیر و فراوانی‌های باقیمانده است. به عبارتی، آزمون مجذور کای از طریق مقایسه این دو فراوانی (مشاهده‌شده و مورد انتظار) با یکدیگر، به آزمون فرضیه پژوهش می‌پردازد (حبیب‌پور و صفری‌شالی، ۱۳۸۸: ۶۳۳). این آزمون را می‌توان در زمانی که توزیع متغیر در آزمون تی تک‌نمونه‌ای نرمال نیست، جایگزین این آزمون کرد که البته در این حالت، باید متغیرهای فاصله‌ای را به متغیر ترتیبی تبدیل کرد.

با استفاده از این آزمون می‌توانیم وضعیت توزیع یک متغیر کیفی (عموماً ترتیبی) را با یک توزیع نظری مقایسه کنیم. به عنوان مثال، می‌توانیم قومیت کارمندان شرکت ایران خودرو را بررسی کرده و وضعیت توزیع فراوانی کارمندان با قومیت‌های گوناگون (ترک، فارس، کرد و...) را ارزیابی کرده و آزمون کنیم که آیا از نظر آماری کارمندان قومیت‌های گوناگون در این شرکت از توزیع فراوانی یکسانی برخوردارند یا خیر.

همچنین اگر از وضعیت درآمدی ساکنان یک شهر اطلاع داریم و متغیر وضعیت درآمدی در سه طبقه کم درآمد، میان درآمد و پردرآمد سنجیده شده است، می‌توانیم فراوانی هرکدام از طبقات درآمدی را با یکدیگر و با یک توزیع نظری مقایسه کنیم.

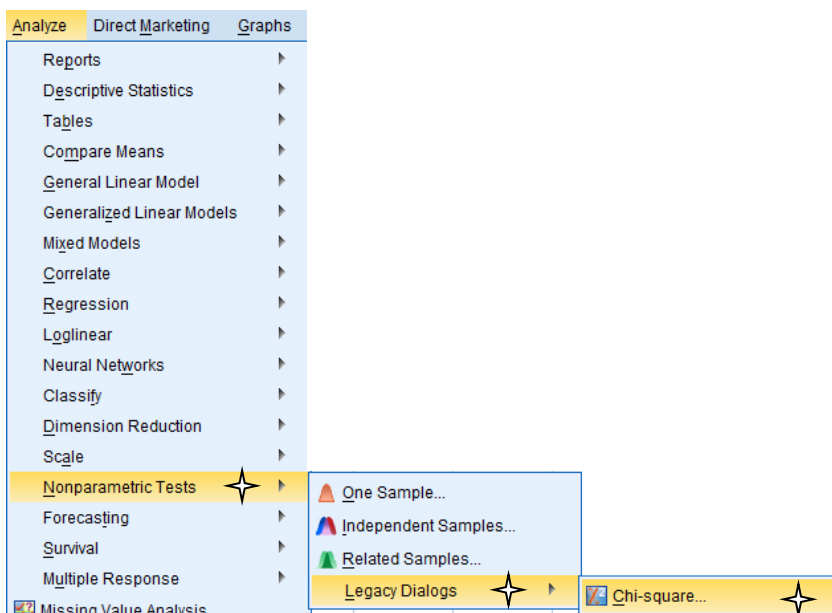
مثال

پژوهشی با هدف بررسی عوامل مؤثر بر ترک اعتیاد انجام شد که در آن از تعدادی از صاحب‌نظران سوالاتی پرسیده شد. در این پژوهش ۱۰۰ نفر از روان‌شناسان و جامعه‌شناسان شرکت داشتند (داده‌ها فرضی است). یکی از سوالاتی که از صاحب‌نظران پرسیده شد این بود که «تا چه حد حمایت عاطفی خانواده معتادین بر ترک موفقیت‌آمیز اعتیاد از سوی معتادین تأثیر دارد». این سوال در قالب طیف لیکرت پنج گزینه‌ای (خیلی کم، کم، متوسط، زیاد و خیلی زیاد) مطرح شد. بر این اساس فرضیه زیر طرح و آزمون شد: فرضیه: حمایت عاطفی خانواده بر ترک موفق اعتیاد از سوی معتادان تأثیرگذار است.

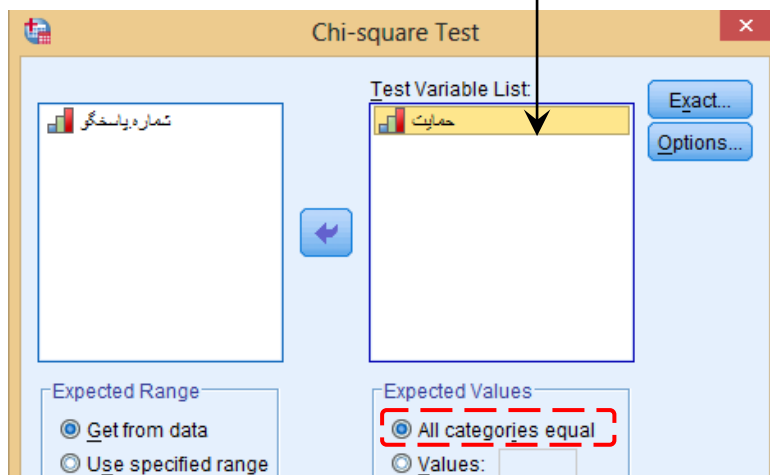
اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Nonparametric Tests ---> Legacy Dialogs ---> Chi Square



متغیر مدنظر را از کادر سمت چپ وارد کادر متغیرهای آزمون (Test Variable List) می‌کنیم. به‌طور پیش‌فرض گزینه همه طبقات برابر باشند (All categories equal) فعال است. این گزینه فراوانی تمامی طبقات متغیر را برابر فرض می‌کند و فراوانی مشاهده شده تمامی طبقات متغیر حمایت عاطفی را با فراوانی نظری (یکسان) طبقات مقایسه می‌کند.



نتایج:

جدول زیر به مقایسه فراوانی‌های مشاهده شده و مورد انتظار می‌پردازد. در ستونی فراوانی‌های مورد انتظار تنها عدد ۲۰ وجود دارد که بدین معناست که ما انتظار داریم که فراوانی تمامی طبقات برابر با ۲۰ باشد. حجم نمونه برابر با ۱۰۰ است و طبقات متغیر حمایت عاطفی شامل ۵ گزینه است که در نتیجه انتظار می‌رود تمامی طبقات فراوانی برابری داشته باشند و فراوانی هر طبقه برابر با ۲۰ باشد. توزیع نظری در اینجا به معنی برابر بودن فراوانی تمامی طبقات است و ما قصد داریم توزیع مشاهده شده را با توزیع نظری مقایسه کنیم.

ستون اول، طبقات متغیر را نشان می‌دهد، ستون دوم، فراوانی مشاهده شده (واقعی) را نشان می‌دهد. ستون سوم فراوانی مورد انتظار (نظری) و ستون چهارم مقدار باقیمانده را که از تفریق فراوانی مشاهده شده از فراوانی مورد انتظار به دست می‌آید را نشان می‌دهد. مقایسه فراوانی مشاهده شده طبقات نشان می‌دهد که فراوانی طبقات با یکدیگر و با توزیع نظری دارای اختلاف است. بین طبقات اختلاف فراوانی زیاد است که نشان می‌دهد توزیع مشاهده شده با توزیع نظری متفاوت است و فراوانی طبقات با یکدیگر برابر نیست.

مقادیر باقیمانده میزان تفاوت مقادیر فراوانی مشاهده شده و فراوانی مورد انتظار را نشان می‌دهند. هرچه مقدار باقیمانده به عدد صفر نزدیک‌تر و از عدد فراوانی موردانتظار (۲۰) کوچکتر باشد نشان از نزدیکی بیشتر فراوانی‌های مشاهده شده و مورد انتظار دارد که به معنای تفاوت کمتر فراوانی طبقات است.

با مقایسه فراوانی‌ها یا بررسی مقادیر باقیمانده دریابیم که طبقات زیاد و متوسط دارای بیشترین فراوانی هستند و بیشترین اختلاف فراوانی‌های مورد انتظار و مشاهده شده در طبقات خیلی کم و زیاد وجود دارد. مثلاً فراوانی طبقه خیلی کم برابر با ۷ است که از مقدار فراوانی مورد انتظار (۲۰) تعداد ۱۳ مورد کمتر است که در ستون باقیمانده با عدد ۱۳- مشخص است که نشان می‌دهد تفاوت مقادیر مشاهده شده و فراوانی مورد انتظار قابل توجه است. اما جهت بررسی معنی‌دار بودن تفاوت فراوانی مشاهده شده و فراوانی مورد انتظار نتایج جدول بعد (Test Statistics) را بررسی می‌کنیم.

حمایت عاطفی

	Observed N فراوانی مشاهده شده	Expected N فراوانی مورد انتظار	Residual باقیمانده
خیلی کم	7	20	-13
کم	8	20	-12
متوسط	30	20	10
زیاد	33	20	13
خیلی زیاد	22	20	2
Total	100		

جدول زیر نتایج آزمون برازندگی مجذورکای را نشان می‌دهد. مقدار مجذورکای به دست آمده ۲۹.۳ است که در سطح خطای کمتر از ۰.۰۵ / معنی‌دار شده است ($P < .05$). به بیان دیگر از جنبه آماری بین فراوانی مشاهده شده و فراوانی مورد انتظار تفاوت وجود دارد و فراوانی طبقات با یکدیگر برابر نیست.

در ادامه باید به مقایسه فراوانی طبقات بپردازیم. بررسی فراوانی طبقات نشان می‌دهد که بیشترین فراوانی به ترتیب متعلق به طبقات زیاد، متوسط و خیلی زیاد است. از مجموع ۱۰۰ پاسخگو، تعداد ۸۵ نفر تأثیر حمایت عاطفی خانواده بر ترک اعتیاد معتادین را متوسط یا بیشتر دانسته‌اند. همچنین ۵۵ نفر از پاسخگویان گفته‌اند که

حمایت عاطفی خانواده به میزان زیاد و خیلی زیاد بر ترک اعتیاد معتادین مؤثر است. در نتیجه فرضیه پژوهش مبنی بر تأثیر حمایت عاطفی خانواده بر ترک موفق اعتیاد به مواد مخدر از سوی معتادان مورد تأیید قرار می‌گیرد.

Test Statistics

	حمایت عاطفی
Chi-Square (کای اسکور)	29.30^a
df (درجه آزادی)	4
Asymp. Sig. (سطح معنی داری)	.000

گزارش:

آزمون برازندگی مجذور کای نشان می‌دهد که از جنبه آماری فراوانی طبقات با یکدیگر متفاوت است و بین فراوانی مشاهده شده و مورد انتظار تفاوت وجود دارد ($P < .001$ ، $df = 4$ ، 29.30 = مجذور کای). فراوانی طبقات متوسط، زیاد و خیلی زیاد بیشتر از طبقات کم و خیلی کم است که گویای این است که از نظر پاسخگویان حمایت خانواده بر ترک اعتیاد معتادین مؤثر است.

آزمون ویلکاکسون

آزمون ویلکاکسون^{۵۳} برای طرح‌های درون‌گروهی (نمونه‌های وابسته) مناسب است. این آزمون از نظر مفهومی مشابه آزمون پارامتریک تی گروه‌های همبسته است. در این آزمون ما یک گروه از افراد داریم که در دو وضعیت یا دو مقطع زمانی مختلف مورد سنجش قرار گرفته‌اند هدف این است که تغییرات نمرات (میانۀ) را در دو وضعیت یا مقطع زمانی مقایسه کنیم. سطح سنجش متغیر در این آزمون باید ترتیبی باشد.

به عنوان مثال می‌توانیم میزان رضایت از زندگی زناشویی تعدادی از زوج‌های جوان را قبل از شرکت در دوره آموزشی مهارت‌های همسراری و بعد از شرکت آن‌ها در این دوره بسنجیم. رضایت زناشویی در این مثال در قالب یک سوال ترتیبی سنجیده شد (یک سوال پنج‌گزینه‌ای از خیلی کم تا خیلی زیاد یا در قالب نمره‌دهی در یک طیف از ۱ تا ۱۰). اگر بین میزان رضایت زوجها از زندگی زناشویی، قبل و بعد از شرکت در دوره مهارت‌های همسراری تفاوت وجود داشته باشد، می‌توانیم نتیجه‌گیری کنیم که دوره آموزشی مهارت‌های همسراری موجب افزایش رضایت‌مندی زوجها از زندگی مشترک شده است. همچنین می‌توانیم شیوه تدریس دو استاد دانشگاه یک کلاس را ارزیابی کنیم و از دانشجویان آن کلاس بپرسیم که شیوه تدریس هر استاد را چقدر می‌پسندید؟ که دانشجویان باید به هر استاد از ۱ تا ۱۰ نمره بدهند. در این مثال (فقط) دانشجویانی که با هر دو استاد کلاس داشته‌اند، به هر استاد نمره می‌دهند و در نهایت نمره‌ای که دو استاد گرفته‌اند را مقایسه می‌کنیم.

مثال

در یک پژوهش به بررسی میزان رضایت از زندگی در بین ۳۰ زوج پرداخته شد. تمامی شرکت‌کنندگان به‌تازگی ازدواج کرده بودند. هدف این پژوهش مقایسه میزان رضایت از زندگی در بین زنان متأهل، قبل و بعد از ازدواج بوده است. در این پژوهش قصد داریم میزان رضایت از زندگی افراد را قبل و بعد از ازدواج مقایسه کنیم.

⁵³ Wilcoxon

در این پژوهش از زنان سوال شد که «در مجموع چقدر از زندگی خودتان رضایت دارید؟» این سوال در قالب طیف لیکرت پنج گزینه‌ای مطرح شد که شامل پاسخ‌های خیلی کم، کم، متوسط، زیاد و خیلی زیاد می‌شود. این سوال یکبار قبل از ازدواج و یکبار بعد از ازدواج از آنان پرسیده شد. لازم به یادآوری است که فقط همان افرادی که در مرحله اول به این سوال پاسخ دادند در مرحله دوم هم شرکت داشته و به این سوال پاسخ دادند.

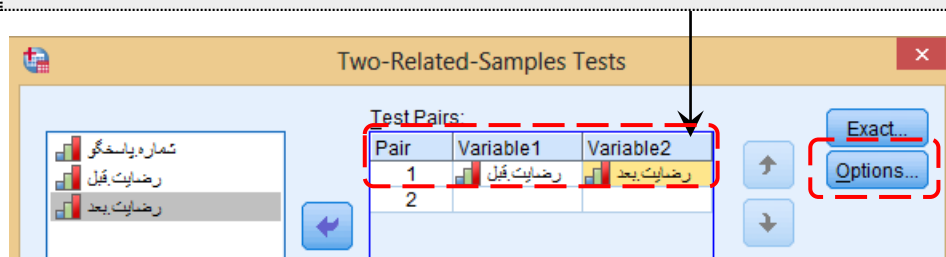
بر این اساس فرضیه زیر طرح و آزمون شد:
فرضیه: میزان رضایت از زندگی در زنان، قبل و بعد از ازدواج متفاوت است.

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Nonparametric Tests ---> Legacy Dialogs ---> **2** Related Samples

دو متغیر را از کادر سمت چپ به ترتیب وارد کادر آزمون جفت‌ها (Test Pairs) کرده و در ردیف اول قرار می‌دهیم.
در کادر متغیر اول (Variable 1)، رضایت قبل از ازدواج و در کادر متغیر دوم (Variable 2)، رضایت بعد از ازدواج را وارد می‌کنیم.
گزینه Options را انتخاب کرده و در پنجره ایجاد شده گزینه Descriptive را انتخاب می‌کنیم تا آمار توصیفی در خروجی ارائه شود. در انتها OK را انتخاب می‌کنیم.



نتایج:

جدول زیر توصیف متغیرها را در دو زمان قبل و بعد از ازدواج نشان می‌دهد. میانگین و انحراف استاندارد متغیرها در این جدول ارائه شده است. میانگین رضایت از زندگی قبل از ازدواج برابر با ۲.۸۰ و بعد از ازدواج برابر با ۳.۷۰ است که نشان می‌دهد رضایت از زندگی در زمان بعد از ازدواج بیشتر از زمان قبل از ازدواج است. جهت بررسی معنی‌دار بودن تفاوت میزان رضایت از زندگی در دو زمان، باید نتایج جدول آخر (Test Statistics) را بررسی کنیم.

Descriptive Statistics

	N	Mean	Std. Deviation	Minimum	Maximum
	تعداد یا فراوانی	میانگین	انحراف استاندارد	مقدار حداقل	مقدار حداکثر
قبل از ازدواج	30	2.80	1.21	1	5
بعد از ازدواج	30	3.70	1.12	2	5

جدول زیر به مقایسه تفاوت رضایت از زندگی در بین شرکت‌کنندگان می‌پردازد. نتایج نشان می‌دهد که میزان رضایت از زندگی ۴ نفر از زنان بعد از ازدواج کاهش پیدا کرده است، رضایت از زندگی ۱۶ نفر از زنان بعد از ازدواج افزایش پیدا کرده است و میزان رضایت از زندگی ۱۰ نفر از زنان قبل و بعد از ازدواج ثابت مانده و تغییری نکرده است.

Ranks

	N	Mean Rank	Sum of Ranks
		میانگین رتبه	مجموع رتبه‌ها
Negative Ranks رتبه‌های منفی	4 ^a	7.13	28.50
Positive Ranks رتبه‌های مثبت	16 ^b	11.34	181.50
Ties موارد برابر	10 ^c		
Total	30		

- a. قبل از ازدواج < بعد از ازدواج
b. قبل از ازدواج > بعد از ازدواج
c. قبل از ازدواج = بعد از ازدواج

جدول زیر مهم‌ترین جدول آزمون ویلکاکسون است و نتایج آزمون فرضیه را نشان می‌دهد. معنی‌دار بودن مقدار Z (Wilcoxon Signed Ranks Test) در سطح خطای کمتر از 0.05 ($P < 0.05$) نشان از تفاوت میزان رضایت از زندگی در بین زنان در دو زمان قبل و بعد از ازدواج دارد. سطح معنی‌داری مقدار Z برابر با 0.004 است که کمتر از مقدار 0.05 است و نشان می‌دهد که از جنبه آماری بین میزان رضایت از زندگی در بین زنان در زمان قبل و بعد از ازدواج تفاوت وجود دارد. نتایج دو جدول بالا نشان می‌دهد میزان رضایت از زندگی زنان پس از ازدواج (با میانگین 3.70) بیشتر از زمان قبل از ازدواج (با میانگین 2.80) است. در صورت وجود رابطه علی بین ازدواج و رضایت از زندگی می‌توانیم نتیجه بگیریم که ازدواج کردن موجب افزایش رضایت از زندگی در زنان می‌شود.

Test Statistics^a

	قبل از ازدواج - بعد از ازدواج
Z	-2.91^b
Asymp. Sig. (2-tailed)	.004

گزارش:

از آزمون ویلکاکسون جهت مقایسه میزان رضایت از زندگی در زنان در دو مقطع قبل و بعد از ازدواج استفاده شد. نتایج نشان داد که میانگین رضایت از زندگی بعد از ازدواج برابر با 3.70 است که بالاتر از میانگین رضایت از زندگی قبل از ازدواج ($M = 2.80$) است ($P < 0.01$) -2.91 $Z =$ نتایج نشان می‌دهد میزان رضایت از زندگی در زنان با ازدواج افزایش می‌یابد.

آزمون فریدمن

آزمون فریدمن^{۵۴} برای طرح‌های درون گروهی (نمونه‌های وابسته) مناسب است. آزمون فریدمن تعمیم یافته آزمون ویلکاکسون است و معادل ناپارامتریک آزمون اندازه‌های مکرر^{۵۵} است. در این آزمون ما یک گروه از افراد یا آزمودنی داریم که در حداقل دو وضعیت یا دو مقطع زمانی مختلف موردسنجش قرار گرفته‌اند. هدف این است که تغییرات نمرات (میانۀ) را در چند (۲ و بیشتر) وضعیت یا مقطع زمانی مقایسه کنیم. سطح سنجش متغیر در این آزمون باید ترتیبی باشد. پژوهش‌گران عموماً از این آزمون جهت رتبه‌بندی یا اولویت‌بندی متغیرها استفاده می‌کنند.

به عنوان نمونه می‌توانیم مثال مطرح شده در آزمون ویلکاکسون را توسعه دهیم و به مقایسه میزان رضایت از زندگی در سه مقطع قبل از ازدواج، یکسال بعد از ازدواج و ده سال بعد از ازدواج بپردازیم. می‌توانیم از مردم بپرسیم که از هرکدام از شرایط فرهنگی، سیاسی، اقتصادی و امنیتی کشور چقدر رضایت دارید و سپس به بررسی این امر بپردازیم که مردم از کدام یک از شرایط رضایت بیشتری دارند. می‌توانیم از تعدادی از صاحب‌نظران اقتصادی بپرسیم که هرکدام از عوامل عدم وابستگی به درآمد نفتی، حمایت از تولید داخلی، افزایش وام‌های خوداشتغالی، بهبود شرایط گردشگری و آموزش نیروی انسانی چقدر در بهبود وضعیت اقتصادی کشور مؤثر هستند و سپس به اولویت‌بندی (رتبه‌بندی) پاسخ‌ها بپردازیم و بررسی کنیم که از نظر صاحب‌نظران کدام یک از عوامل فوق تأثیر بیشتری در بهبود وضعیت اقتصادی دارد.

مثال

پژوهشی با هدف تعیین و رتبه‌بندی "مهم‌ترین ویژگی‌های اساتید دانشگاه" در دانشگاه شیراز انجام شد و در آن از صد دانشجوی دکتری دانشگاه شیراز پرسیده شد که نظر خودشان را در ارتباط با هرکدام از ویژگی‌های مطرح شده اساتید در پرسشنامه ابراز کنند. از دانشجویان پرسیده شد که تا چه حد به هرکدام از ویژگی‌های اساتید اهمیت می‌دهند؟ پنج ویژگی‌های مطرح شده در ارتباط با اساتید عبارتند از: تسلط علمی،

⁵⁴ Friedman

⁵⁵ Repeated Measures

اخلاق خوب در روابط با دانشجویان، ظاهر مرتب و آراسته، فن بیان مناسب و رعایت انصاف در نمره دهی. دانشجویان به هر کدام از ویژگی های مطرح شده در قالب طیف ۱۱ گزینه ای (از ۰ تا ۱۰ که نمره بالاتر به معنای اهمیت بیشتر است) نمره دادند.

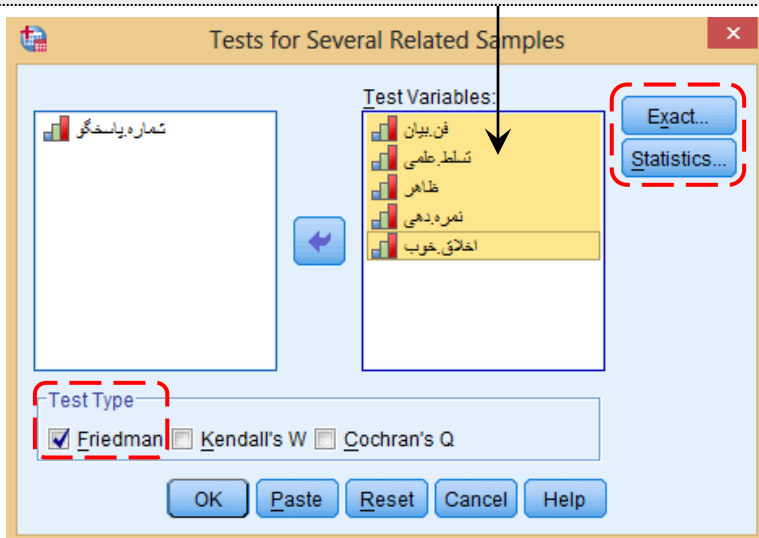
بر این اساس سوال زیر طرح و آزمون شد:
کدام ویژگی های اساتید اهمیت بالاتری برای دانشجویان دارد؟
با استفاده از آزمون فریدمن سوال مطرح شده را ارزیابی می کنیم.

اجرا:

مسیر زیر را دنبال می کنیم:

Analyze ---> Nonparametric Tests ---> Legacy Dialogs ---> **K** Related Samples

متغیرها (ویژگی های اساتید) را از کادر سمت چپ وارد کادر متغیرهای آزمون (Test Variables) می کنیم.
گزینه Statistics را انتخاب کرده و در پنجره جدید گزینه Descriptive را انتخاب می کنیم.



نتایج:

در جدول بعد میانگین و انحراف استاندارد متغیرها ارائه شده است. دامنه میانگین از ۰ تا ۱۰ است. مقایسه میانگین ویژگی های اساتید نشان می دهد که بالاترین میانگین

(۷.۷۹) متعلق به ویژگی انصاف در نمره‌دهی و پایین‌ترین میانگین (۳.۳۱) متعلق به ویژگی ظاهر آراسته است. جهت بررسی معنی‌دار بودن تفاوت بین ویژگی‌های اساتید و رتبه‌بندی مهم‌ترین ویژگی‌های اساتید از نظر دانشجویان، باید نتایج جدول آخر (Test Statistics) را بررسی کنیم.

Descriptive Statistics

	N	Mean	Std. Deviation	Minimum	Maximum
فن بیان	100	4.99	1.83	2	9
تسلط علمی	100	4.48	1.66	2	10
اخلاق خوب	100	5.01	1.79	2	10
انصاف در نمره دهی	100	7.79	1.55	5	10
ظاهر آراسته	100	3.31	2.01	0	8

جدول زیر وضعیت رتبه‌بندی متغیرها (ویژگی‌های اساتید) را نشان می‌دهد. میانگین رتبه (Mean Rank) هرکدام از ویژگی‌ها در جدول گزارش شده است. مقایسه میانگین رتبه‌ها نشان می‌دهد که بالاترین میانگین رتبه (۴.۵۸) به ویژگی انصاف در نمره‌دهی اختصاص دارد و که بدین معناست که مهم‌ترین ویژگی استاد از نظر دانشجویان ویژگی انصاف در نمره‌دهی است. بعد از ویژگی فوق، مهم‌ترین ویژگی اساتید به ترتیب شامل ویژگی‌های فن بیان مناسب، اخلاق خوب در روابط با دانشجویان، تسلط علمی و ظاهر مرتب و آراسته می‌شود. لازم به ذکر است که میانگین رتبه با میانگین حسابی تفاوت دارد و نحوه محاسبه این دو میانگین متفاوت است.

Ranks

	Mean Rank
فن بیان	3
تسلط علمی	2.56
اخلاق خوب	2.93
انصاف در نمره‌دهی	4.58
ظاهر آراسته	1.94

جدول زیر مهم‌ترین جدول آزمون فریدمن است که قبل از تفسیر جداول دیگر نخست باید نتایج این جدول را ارزیابی کرد و در صورت معنی‌دار بودن آزمون فریدمن، به تفسیر نتایج جداول توصیفی و میانگین رتبه بپردازیم. این جدول معنی‌داری آماری را

نشان می‌دهد. مقدار مجذور کای به دست آمده برابر با ۱۶۳.۷ است که در سطح خطای کمتر از ۰.۰۵ قرار دارد ($P < ۰.۰۵$). معنی‌دار بودن آزمون فریدمن بدین معناست که رتبه‌بندی ویژگی‌های اساتید از نظر دانشجویان بامعناست و دانشجویان رتبه‌بندی متفاوتی از ویژگی‌های اساتید دارند.

Test Statistics^a

N	100
Chi-Square	163.75
df	4
Asymp. Sig.	.000

a. Friedman Test

گزارش:

از آزمون فریدمن برای رتبه‌بندی مهم‌ترین ویژگی‌های اساتید استفاده شد. آزمون فریدمن نشان داد که اهمیت و رتبه ویژگی‌های مطرح‌شده در مورد اساتید با یکدیگر متفاوت است ($P < ۰.۰۰۱$). $df = ۴$ ، $۱۶۳.۷۴ =$ مجذور کای). مقایسه میانگین رتبه‌ها نشان می‌دهد که مهم‌ترین ویژگی اساتید از نظر دانشجویان به ترتیب ویژگی انصاف در نمره‌دهی، فن بیان مناسب و اخلاق خوب در روابط با دانشجویان است. میانگین رتبه این ویژگی‌ها به ترتیب ۴.۵۸، ۳ و ۲.۹۳ است (میانگین رتبه همه ویژگی‌ها را در یک جدول گزارش کنید).

آزمون مان-ویتنی

آزمون مان-ویتنی^{۵۶} معادل ناپارامتریک آزمون تی گروه‌های مستقل است با این تفاوت متغیر وابسته در آزمون مان-ویتنی در سطح سنجش ترتیبی است. از این آزمون برای مقایسه دو گروه مستقل استفاده می‌شود و در آن داده‌هایی که از طرح‌های گروه‌های مستقل به دست می‌آیند مورد مقایسه قرار می‌گیرند. در این آزمون مانند ویلکاکسون رتبه‌بندی روی می‌دهد و محاسبات بر روی رتبه‌ها انجام می‌گیرد. این آزمون باید در شرایط زیر به آزمون تی گروه‌های مستقل ترجیح داده شود:

- ۱- هنگامی که داده‌ها فقط به صورت مقیاس اندازه‌گیری ترتیبی هستند.
 - ۲- هنگامی که داده‌ها فاصله‌ای یا نسبی، اما دارای توزیع غیرنرمال هستند (مثلاً کجی شدیدی دارند).
 - ۳- هنگامی که داده‌ها فاصله‌ای یا نسبی هستند، اما واریانس‌های دو نمونه در آزمون واریانس برابر نیستند (بریس و همکاران، ۱۳۹۱: ۱۳۴-۱۳۵).
- ❑ نکته: در آزمون‌های مان-ویتنی و کروسکالیس-والیس گروه‌ها می‌توانند حجم متفاوتی داشته باشند و برابر بودن تعداد اعضای گروه‌ها در این آزمون‌ها ضروری نیست.

مثال

در پژوهشی بر روی ۱۰۰ نفر از زنان و مردان که به طور تصادفی انتخاب شده بودند از آنان سوالاتی پرسیده شد. یکی از سوالات این بود که «تا چه حد به سلامت جسم خودتان اهمیت می‌دهید؟» پاسخ‌ها در قالب طیف لیکرت شش‌گزینه‌ای (اصلاً، خیلی کم، کم، متوسط، زیاد و خیلی زیاد) مطرح شد که در آن به پاسخ‌ها از ۱ (اصلاً) تا ۶ (خیلی زیاد) نمره داده شد.

بر این اساس فرضیه زیر طرح و آزمون شد:

فرضیه: زنان و مردان اهمیت متفاوتی بری سلامت جسمانی خودشان قائل هستند.
با استفاده از آزمون مان-ویتنی فرضیه فوق را می‌آزماییم.

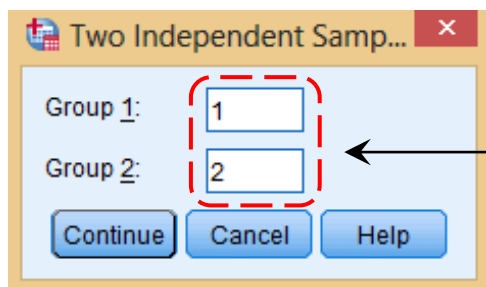
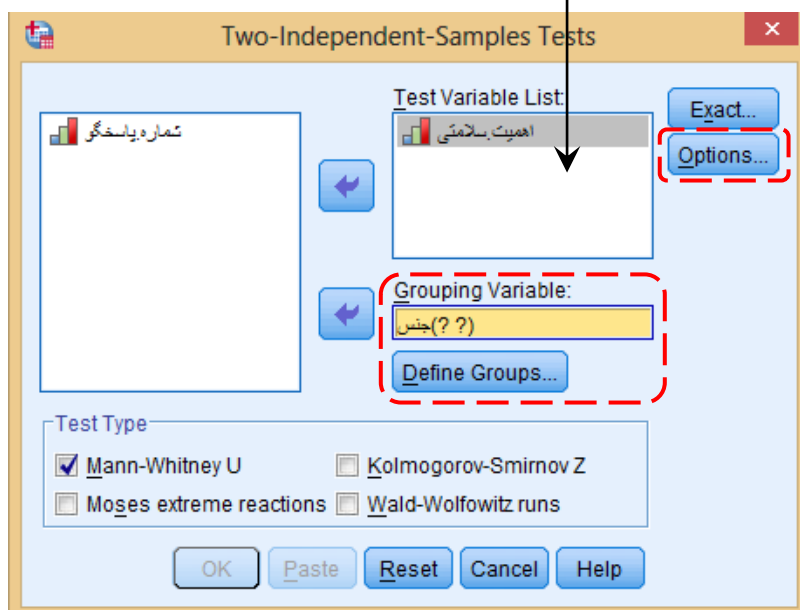
⁵⁶ Mann-Whitney

اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Nonparametric Tests ---> Legacy Dialogs ---> **2**
Independent Samples

متغیر ترتیبی که قصد مقایسه آن در دو گروه را داریم از کادر سمت چپ وارد کادر فهرست متغیرهای آزمون (Test Variable List) می‌کنیم. گزینه Options را انتخاب کرده و در پنجره ایجاد شده گزینه Descriptive را انتخاب می‌کنیم. متغیر جنس که مبنای گروه بندی است را وارد کادر Grouping Variable می‌کنیم. سپس گزینه Define Group را انتخاب می‌کنیم.



با انتخاب گزینه Define Group این پنجره ایجاد می‌شود. باید کدهایی که به دو جنس زن (کد ۱) و مرد (کد ۲) در هنگام تعریف داده اختصاص داده بودیم را وارد کنیم. در نتیجه اعداد ۱ و ۲ را وارد می‌کنیم.

نتایج:

جدول زیر به توصیف متغیرها در آزمون مان-ویتنی می پردازد. مواردی مانند حجم نمونه، میانگین و انحراف استاندارد در جدول گزارش شده است البته برای متغیر جنس افراد که یک متغیر اسمی است آماره‌های میانگین و انحراف استاندارد کاربردی ندارند.

Descriptive Statistics

	N	Mean	Std. Deviation	Minimum	Maximum
اهمیت سلامتی	100	3.84	1.39	1	6
جنس	100	1.58	.50	1	2

جدول زیر رتبه‌های دو جنس زن و مرد را در متغیر اهمیت سلامتی نشان می‌دهد. مردان ۵۸ نفر و زنان ۴۲ نفرند. میانگین رتبه در زنان بیشتر از مردان است. میانگین رتبه اهمیت سلامتی در زنان ۵۳.۲۹ و در مردان ۴۸.۴۸ است که البته این تفاوت در میانگین رتبه‌ها باید از جنبه آماری و با بررسی سطح معنی‌داری ارزیابی شود.

Ranks

جنس	N	Mean Rank	Sum of Ranks
زن	42	53.29	2238
مرد	58	48.48	2812
Total	100		

جدول زیر مهم‌ترین جدول آزمون مان-ویتنی است که مقدار سطح معنی‌داری به دست آمده تأیید یا رد شدن فرضیه را نشان می‌دهد. سطح معنی‌داری به دست آمده برابر با ۰/۳۹۷ است که بیشتر از مقدار مفروض ۰/۰۵ است و در نتیجه می‌توان گفت که زنان و مردان اهمیت تقریباً یکسانی برای سلامت جسمانی خودشان قائل هستند و فرضیه پژوهش رد می‌شود.

Test Statistics^a

	سلامتی
Mann-Whitney U	1101
Wilcoxon W	2812
Z	-.847
Asymp. Sig. (2-tailed)	.397

a. Grouping Variable: جنس

گزارش:

از آزمون مان-ویتنی جهت مقایسه میزان اهمیت سلامتی در بین زنان و مردان استفاده شد. سطح معنی‌داری به دست آمده نشان می‌دهد که زنان و مردان اهمیت تقریباً یکسانی برای سلامتی خودشان قائل هستند ($P = .397$ ، $Z = -.847$). در نتیجه فرضیه پژوهش مبنی بر این که زنان و مردان اهمیت متفاوتی برای سلامت جسمانی خودشان قائل هستند، رد می‌شود.

• آزمون کروسکال والیس

آزمون کروسکال-والیس^{۵۷} تعمیم یافته آزمون مان-ویتنی و معادل ناپارامتریک آزمون تحلیل واریانس یک راهه (آزمون F) است. از این آزمون در طرح‌های گروه‌های مستقل استفاده می‌شود. زمانی از این آزمون استفاده می‌کنیم که تعداد گروه‌های مستقلی که قصد مقایسه آنان را داریم بیشتر از دو گروه باشد. مقیاس اندازه‌گیری در کروسکال-والیس حداقل باید ترتیبی باشد. این آزمون برای مقایسه میانگین‌های بیش از دو نمونه رتبه‌ای (یا فاصله‌ای) به کار می‌رود. فرضیه‌ها در این آزمون بدون جهت است، یعنی فقط تفاوت میانگین را نشان می‌دهد و جهت را (بزرگتر یا کوچکتر بودن گروه‌ها را از نظر میانگین) نشان نمی‌دهد. (حبیب‌پور و صفری، ۱۳۸۸: ۶۸۴).

مثلاً زمانی که بخواهیم میزان رضایت از زندگی (مقیاس ترتیبی) سه گروه جوانان، میانسالان و افراد مسن را با هم مقایسه کنیم از این آزمون استفاده می‌کنیم. یا زمانی که بخواهیم یک متغیر ترتیبی را در بین سه گروه از دانشجویان مقاطع کارشناسی، کارشناسی ارشد و دکترا مقایسه کنیم از این آزمون استفاده می‌کنیم.

مفروضات آزمون مان-ویتنی در مورد آزمون کروسکالیس والیس هم صدق می‌کند (سطح سنجش حداقل ترتیبی، متغیرها از نوع فاصله‌ای/نسبی اما دارای توزیع غیرنرمال و متغیرها از نوع فاصله‌ای/نسبی اما واریانس نابرابر گروه‌ها).

مثال

مثال آزمون مان-ویتنی را گسترش داده و در آزمون کروسکال والیس استفاده می‌کنیم. در پژوهشی بر روی ۱۰۰ نفر که به طور تصادفی انتخاب شده بودند از آنان سوالاتی پرسیده شد. یکی از سوالات این بود که «تا چه حد به سلامت جسم خودتان اهمیت می‌دهید؟» پاسخ‌ها در قالب طیف لیکرت شش گزینه‌ای (اصلاً، خیلی کم، کم، متوسط، زیاد و خیلی زیاد) مطرح شد که در آن به پاسخ‌ها از ۱ (اصلاً) تا ۶ (خیلی زیاد) نمره داده شد. این پژوهش بر روی سه گروه جوانان، میانسالان و افراد مسن انجام شد و هدف

⁵⁷ Kruskal-Wallis

پژوهش بررسی این موضوع بود که کدام یک از گروه‌های سنی اهمیت بیشتری برای سلامت جسمانی خودشان قائل هستند.

متغیر سن در سطح سنجش ترتیبی سنجیده شد و شامل سه گروه جوان، میانسال و مسن می‌شود و متغیر اهمیت سلامت جسمانی در سطح سنجش ترتیبی سنجیده شد.

بر این اساس فرضیه زیر طرح و آزمون شد:

فرضیه: افراد جوان، میانسال و مسن اهمیت متفاوتی برای سلامت جسمانی خودشان قائل هستند.

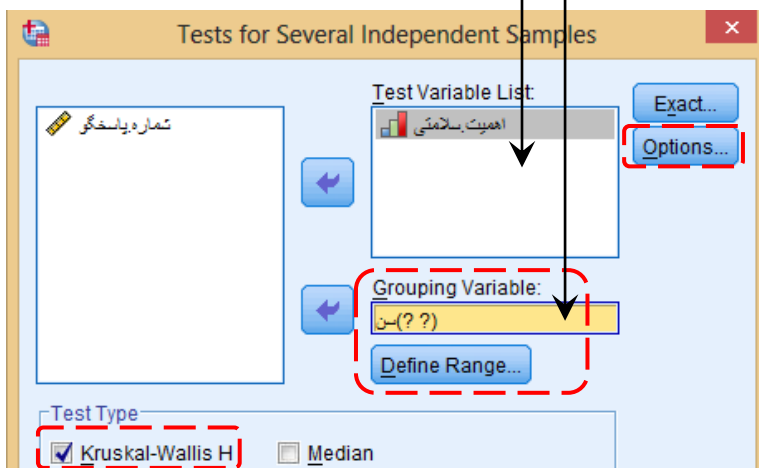
با استفاده از آزمون کروسکال والیس این فرضیه را آزمون می‌کنیم.

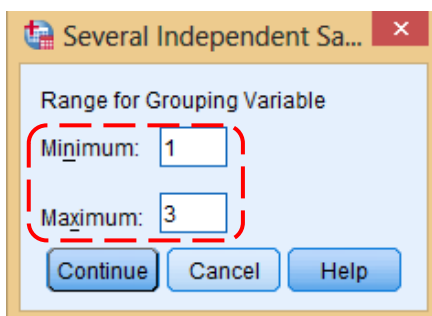
اجرا:

مسیر زیر را دنبال می‌کنیم:

Analyze ---> Nonparametric Tests ---> Legacy Dialogs ---> **K**
Independent Samples

متغیر اهمیت سلامتی که قصد مقایسه آن را در گروه‌ها داریم از کادر سمت چپ وارد کادر فهرست متغیرهای آزمون (Test Variable List) می‌کنیم.
متغیر گروه سنی که مبنای گروه‌بندی است را وارد کادر Grouping Variable می‌کنیم. سپس گزینه معرفی دامنه (Define Range) را انتخاب می‌کنیم.
گزینه Options را انتخاب کرده و در پنجره ایجاد شده گزینه Descriptive را انتخاب می‌کنیم تا آمار توصیفی در خروجی ارائه شود.





با انتخاب گزینه Define Range این پنجره ایجاد می شود. باید دامنه کدهایی که به سه گروه جوانان (کد ۱)، میانسالان (کد ۲) و افراد مسن (کد ۳) داده بودیم در اینجا مشخص کنیم. حداقل کد که ۱ است وارد کادر Minimum و حداکثر کد که ۳ هست وارد کادر Maximum می کنیم. در انتها Continue و OK را انتخاب می کنیم.

نتایج:

جداول زیر مشابه جداول آزمون مان-ویتنی است. جدول زیر به توصیف دو متغیر وارد شده در آزمون مان-ویتنی می پردازد. مواردی مانند حجم نمونه، میانگین و انحراف استاندارد در جدول گزارش شده است البته برای متغیر سن افراد که در این مثال یک متغیر کیفی است، آماره های میانگین و انحراف استاندارد کاربردی ندارند.

Descriptive Statistics

	N	Mean	Std. Deviation	Minimum	Maximum
اهمیت سلامتی	100	3.89	1.49	1	6
جنس	100	1.88	.77	1	3

جدول زیر رتبه های افراد جوان، میانسال و مسن را در متغیر سلامتی نشان می دهد. نتایج نشان می دهد که میانگین رتبه جوانان از سایر گروه ها بیشتر است و بعد از جوانان، میانسالان میانگین رتبه بالاتری دارند. میانگین رتبه جوانان برابر با ۶۰.۷، میانسالان برابر با ۴۶.۹ و افراد مسن برابر با ۴۱.۱ است. جهت بررسی معنی دار بودن تفاوت میانگین رتبه در بین گروه ها، به جدول بعد و سطح معنی داری به دست آمده نگاه می کنیم.

Ranks

سن	N	Mean Rank
جوان	36	60.74
میانسال، اهمیت سلامتی	40	46.91
مسن	24	41.13
Total	100	

جدول زیر مهم‌ترین جدول آزمون کروسکال والیس است. مقدار آزمون مجذور کای برابر با ۸.۰۱ است که در سطح خطای کمتر از ۰.۰۵ / معنی‌دار است ($P < .05$). سطح معنی‌داری به دست آمده برابر با ۰.۰۱۸ است که کمتر از مقدار مفروض ۰.۰۵ است و در نتیجه می‌توان گفت که گروه‌های سنی مختلف (اهمیت متفاوتی برای سلامت جسمانی خودشان قائل هستند و فرضیه پژوهش تأیید می‌شود. نتایج نشان می‌دهد جوانان بالاترین اهمیت را برای سلامت جسمانی خودشان قائل هستند (داده‌ها فرضی است).

☑ نکته: آزمون کروسکال والیس نشان می‌دهد که آیا بین گروه‌های سنی مختلف در میزان اهمیتی که برای سلامتی شان قائل هستند تفاوت وجود دارد یا خیر؛ اما این آزمون مشخص نمی‌کند که بین کدام گروه‌ها تفاوت معنی‌دار وجود دارد. مثلاً آیا اهمیت سلامتی برای افراد میان‌سال و مسن یکسان است؟ برای افراد جوان و مسن چطور؟ در نتیجه بعد از آزمون کروسکال‌والیس و در صورت معنی‌دار بودن، باید به مقایسه جفتی گروه‌ها با هم بپردازیم. در نتیجه برای این کار از آزمون مان-ویتنی استفاده می‌کنیم و به مقایسه تمامی زوج گروه‌ها می‌پردازیم. در مثال کتاب، باید سه آزمون مان-ویتنی انجام بدهیم (یکبار جوانان با میانسالان، یکبار جوانان با افراد مسن و یکبار افراد میانسال با افراد مسن).

Test Statistics^{a,b}

	اهمیت سلامتی
Chi-Square	8.01
df	2
Asymp. Sig.	.018

گزارش:

از آزمون کروسکال-والیس برای مقایسه میزان اهمیت سلامتی در بین سه گروه جوانان، میانسالان و افراد مسن استفاده شد. نتایج نشان داد که در سطح اطمینان ۹۵ درصد بین حداقل دو گروه در میزان اهمیت سلامتی تفاوت وجود دارد ($P < .05$ ، $df = 2$ ، 8.013 = مجذور کای). در نتیجه فرضیه پژوهش تایید می‌شود که بدین معناست که اهمیت سلامتی در بین سه گروه جوانان، میانسالان و افراد مسن تفاوت دارد. افراد جوان بیشتر از افراد میانسال و افراد میانسال بیشتر از افراد مسن به سلامتی اهمیت می‌دهند (جهت مقایسه گروه‌ها به صورت جفتی از آزمون مان-ویتنی استفاده می‌کنیم و نتایج آن را گزارش می‌کنیم).

تعریف

در پژوهشی به بررسی اثر کلاس‌های یوگادرمانی بر میزان اضطراب افراد پرداخته شد. بر این اساس تعداد ۱۴ نفر از افراد به طور تصادفی انتخاب شدند و به بررسی میزان اضطراب آنان در سه مقطع زمانی پرداخته شد: مرحله اول قبل از شروع تمرین‌های یوگا، مرحله دوم یک هفته بعد از تمرین‌های یوگا و مرحله سوم یک ماه بعد از تمرین‌های یوگا. از آزمودنی‌ها خواسته شد تا میزان اضطراب خود را در هر مرحله بر روی یک طیف ۱۰ گزینه‌ای از ۱ (کاملاً مضطربم) تا ۱۰ (کاملاً احساس آرامش می‌کنم) مشخص کنند. پاسخ‌ها بدین صورت است.

شماره آزمودنی	۱۴	۱۳	۱۲	۱۱	۱۰	۹	۸	۷	۶	۵	۴	۳	۲	۱
اضطراب قبل از یوگادرمانی	۹	۱۰	۸	۸	۶	۷	۵	۵	۳	۷	۲	۱	۴	۹
اضطراب یک هفته بعد از	۲	۷	۵	۳	۵	۵	۵	۴	۴	۷	۱	۱	۳	۲
اضطراب یک ماه بعد از شروع	۵	۸	۴	۴	۵	۴	۴	۳	۶	۷	۲	۱	۳	۴

۱- آیا یوگا درمانی موجب کاهش اضطراب افراد شده است؟ یکبار از آزمون ویلکاکسون و یکبار از آزمون فریدمن استفاده کنید. در آزمون ویلکاکسون، مراحل قبل از یوگادرمانی و یک هفته بعد از یوگا درمانی را مقایسه کنید.

به صفحه ۱۸۸ رجوع کنید:

۲- میزان اعتماد به اعضای خانواده در بین شهروندان تهرانی چقدر است؟ آیا شهروندان تهرانی به اعضای خانواده خود اعتماد دارند؟

۳- شهروندان تهرانی به کدام یک از گروه‌ها اعتماد بیشتری دارند؟ میزان اعتماد به هر کدام از گروه‌های اعضای خانواده، بستگان نزدیک، همسایگان و همکاران را اولویت‌بندی کنید.

به جدول صفحه ۱۸۹ رجوع کنید:

۴- میزان اعتماد به همسایگان در بین مردان بیشتر است یا زنان؟ به‌بیان دیگر بررسی کنید که آیا مردان اعتماد بیشتری به همسایگان خود دارند؟

۵- میزان درآمد ۳۰ نفر از کارمندان یک اداره دولتی با میزان تحصیلات مختلف در جدول زیر گزارش شده است. افراد بر حسب میزان درآمدشان در چهار گروه و بر حسب مدرک تحصیلی در سه گروه قرار گرفتند:

مدرک تحصیلی	میزان درآمد
کد ۱: دیپلم	کد ۱: درآمد کمتر از ۵۰۰ هزار تومان
کد ۲: فوق دیپلم و کارشناسی	کد ۲: درآمد بین ۵۰۰ هزار تا ۱ میلیون تومان
کد ۳: کارشناسی ارشد	کد ۳: درآمد بین ۱ تا ۲ میلیون تومان
	کد ۴: درآمد بیشتر از ۴ میلیون تومان

۶- آیا بین میزان درآمد افراد با مدرک تحصیلی مختلف رابطه وجود دارد؟

۱	۱	۱	۱	۱	۱	۱	۱	۱	۲	۲	۲	۲	۲	۲	مدرک تحصیلی
۲	۲	۳	۲	۴	۱	۱	۲	۱	۲	۱	۲	۲	۱	۳	میزان درآمد
۲	۲	۲	۲	۲	۲	۳	۳	۳	۳	۳	۳	۳	۳	۳	مقطع تحصیلی
۲	۲	۳	۲	۲	۳	۴	۴	۴	۳	۳	۴	۳	۳	۴	میزان درآمد

واژه نامه فصل پنجم

Line of Regression	خط رگرسیون	One-Sample T test	آزمون t تک نمونه ای
Relation	رابطه	Independent-Samples T test	آزمون t گروه های مستقل
Statistical Relation	رابطه آماری	Repeated Measures test	آزمون اندازه های مکرر
Linear Relationship	رابطه خطی	Paired-Samples T test	آزمون تی گروه های همبسته
Ranking	رتبه بندی	Independence χ^2 test	آزمون خی دو استقلال
Logistic Regression	رگرسیون لجستیک	Friedman test	آزمون فریدمن
Correlation Coefficient	ضرایب همبستگی	Kruskal-Wallis test	آزمون کروسکال-والیس
Coefficient of Determination	ضریب تعیین	Mann-Whitney test	آزمون مان-ویتنی
Adjusted R Square	ضریب تعیین تعدیل شده	Post Hoc test	آزمون های تعقیبی
Observed Frequency	فراوانی مشاهده شده	Distribution test Free	آزمون های توزیع - آزاد
Chi-Square	کای اسکوئر	Wilcoxon test	آزمون ویلکاکسون
Covariance	کوواریانس (هم تغییری)	One-Way ANOVA	آزمون تحلیل واریانس (یک راهه)
Prediction Variable	متغیر پیش بین	Measures of Association	اندازه های پیوستگی
R Square	مجذور ضریب همبستگی چندگانه	Asymmetric measures	اندازه های نامتقارن
Sum of Squares	مجموع مجذورات	Beta	بتا (ضریب استاندارد شده رگرسیونی)
Zero-Order	مرتبه صفر	Association	پیوستگی
Compare Means	مقایسه میانگین ها	Discriminant Analysis	تحلیل تشخیصی
Spearman's rho Correlation	همبستگی اسپیرمن	Regression Analysis	تحلیل رگرسیون
Pearson Correlation	همبستگی پیرسون	Path Analysis	تحلیل مسیر
Partial Correlation	همبستگی جزئی	Crosstab	جدول تقاطعی

راهنمای انتخاب آزمون آماری مناسب

"آمار فقط ابزاری برای تحلیل است، ولی این ما هستیم که ابزار مناسب برای کارمان را انتخاب می‌کنیم"

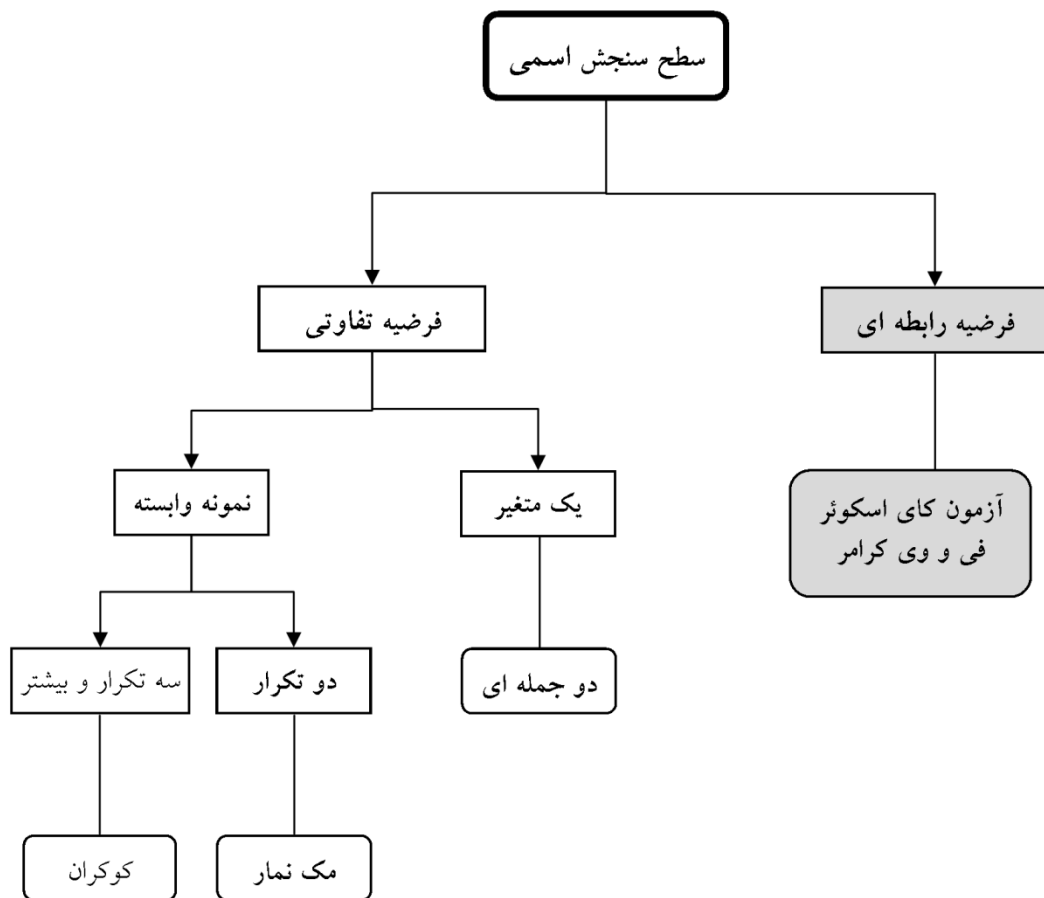
انتخاب آزمون مناسب برای تحلیل همواره یکی از دغدغه‌های پژوهش‌گران بوده است. هر فرضیه یا سوالی، آزمون مناسب با خود را دارد و سوال اصلی این‌جاست که چگونه باید آزمون متناسب را انتخاب کرد؟ انتخاب آزمون مناسب بستگی به چند عامل دارد: سطح سنجش متغیرها، نوع فرضیه، مستقل یا وابسته بودن نمونه‌ها و تعداد گروه‌ها یا متغیرها. در ادامه به طور کوتاه به هر کدام از عوامل بالا اشاره می‌شود.

- ۱- سطح سنجش متغیرها: شامل سه سطح اسمی، ترتیبی و فاصله‌ای/نسبی می‌شود.
- ۲- نوع فرضیه‌ها: شامل دودسته فرضیه‌های رابطه‌ای و تفاوتی می‌شود. در فرضیه‌های رابطه‌ای به بررسی ارتباط بین دو یا چند متغیر با یکدیگر پرداخته می‌شود و در فرضیه‌های تفاوتی یا مقایسه‌ای به مقایسه میانگین (یا میانه یا فراوانی) دو یا چندگروه پرداخته می‌شود.
- ۳- مستقل یا وابسته بودن متغیرها: در آزمون‌های تفاوتی یا مقایسه‌ای، بسته به این که متغیرها وابسته یا مستقل از یکدیگر هستند از آزمون‌های مختلفی استفاده می‌شود.
- ۴- تعداد گروه‌ها یا متغیرها: تعداد گروه‌های مورد تحلیل، بر انتخاب آزمون آماری تأثیرگذار است.

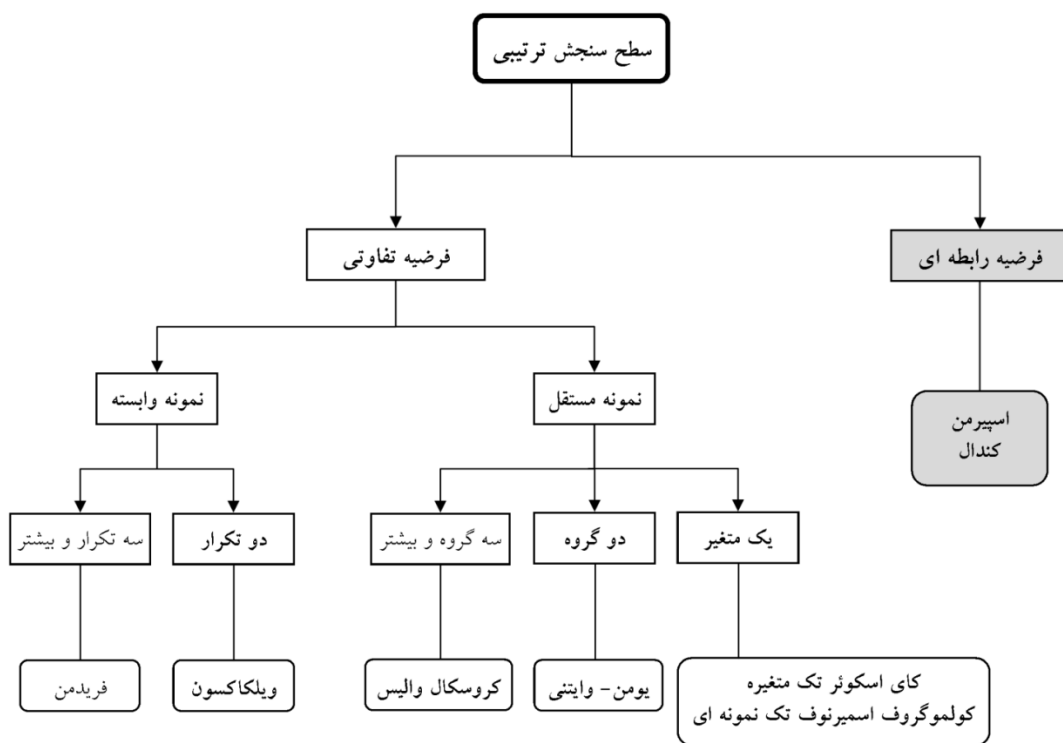
☑ نکته: زمانی که یک متغیر، ترکیبی از چندمتغیر ترتیبی باشد، می‌توان با کمی تسامح و تساهل آن را به عنوان متغیر فاصله‌ای در نظر گرفت. در این مواقع به متغیر ایجاد شده متغیر تراکمی یا شبه‌فاصله‌ای گفته می‌شود و می‌توان از آزمون‌های متناسب با متغیرهای فاصله‌ای/نسبی برای این دسته از متغیرها استفاده کرد.

در ادامه درخت انتخاب آزمون‌ها ارائه شده است.

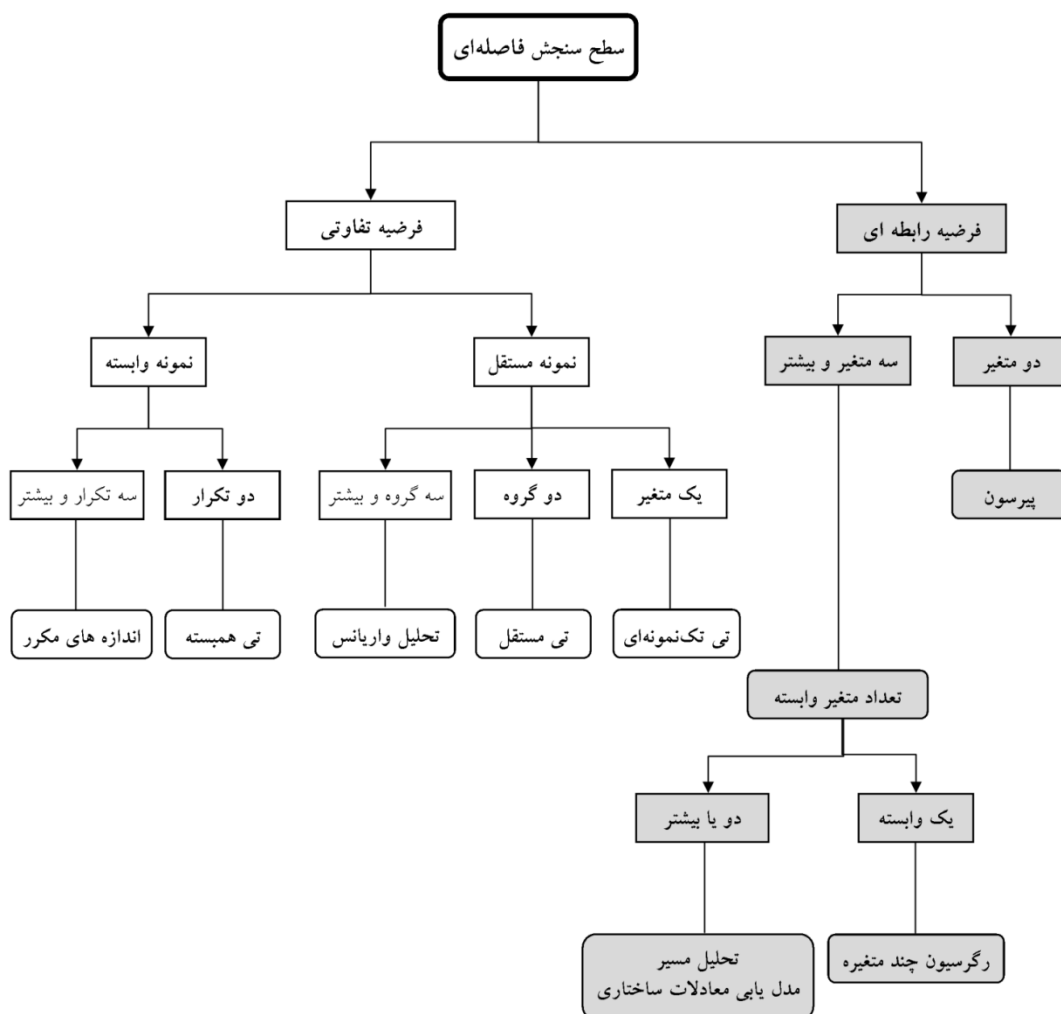
راهنمای انتخاب آزمون‌های آماری در سطح سنجش اسمی



راهنمای انتخاب آزمون‌های آماری در سطح سنجش ترتیبی



راهنمای انتخاب آزمون‌های آماری در سطح سنجش فاصله‌ای/نسبی



منابع:

- ای‌برگ، کریس و لاتین ریچارد دبلیو (۱۳۸۸) روش‌های تحقیق در تندرستی، تربیت بدنی، علوم ورزشی و تفریحات، ترجمه بهروز عبدلی، منصور احمدی و الهام عظیم زاده، تهران: علم و حرکت.
- بی‌بی، ارل (۱۳۸۱)، روش‌های تحقیق در علوم اجتماعی، ترجمه: رضا فاضل، جلد اول، تهران: انتشارات سمت.
- بریس، نیکلا و کمپ، ریچارد و سلنگار، رزمی (۱۳۹۱). تحلیل داده‌های روانشناسی با برنامه SPSS، ترجمه خدیجه علی‌آبادی و سیدعلی صمدی. ویرایش سوم، تهران: دوران.
- حبیب پور، کرم و صفری، رضا (۱۳۸۸) راهنمای جامع کاربرد SPSS در تحقیقات پیمایشی، تهران: نشر لویه.
- دلاور، علی (۱۳۷۸) روش تحقیق در روان‌شناسی و علوم تربیتی، تهران: نشر ویرایش.
- دلاور، علی (۱۳۹۰) روش تحقیق در روان‌شناسی و علوم تربیتی، تهران: نشر ویرایش.
- دواس، دی. ای (۱۳۷۶)، پیمایش در تحقیقات اجتماعی، ترجمه هوشنگ نایبی، تهران: نشر نی.
- ساعی، علی (۱۳۸۸). درآمدی بر روش پژوهش اجتماعی: رویکرد تحلیل کمی در علوم اجتماعی با نرم افزار SPSS، تهران: انتشارات آگاه.
- سرمد، زهره و همکاران (۱۳۸۲)، روش‌های تحقیق در علوم رفتاری، چاپ هفتم، تهران: انتشارات آگاه.
- سیگل، سیدنی (۱۳۷۲) آمار غیرپارامتری برای علوم رفتاری، ترجمه یوسف کریمی، تهران: دانشگاه علامه طباطبائی.
- عسگری، علی و هومن، حیدرعلی (۱۳۹۲) آزمون مجذور کای (مبانی، کاربرد و تفسیر)، تهران: انتشارات سمت.
- کرلینجر و پدهازور (۱۳۶۶) رگرسیون چند متغیری در پژوهش رفتاری، ترجمه حسن سرایی، تهران: انتشارات مرکز نشر دانشگاهی.

- کلانتری، خلیل (۱۳۸۲) پردازش و تحلیل داده‌ها در تحقیقات اجتماعی-اقتصادی، تهران: نشر شریف.
- کلانتری، خلیل (۱۳۸۷) مدل سازی معادلات ساختاری در تحقیقات اجتماعی-اقتصادی، تهران: فرهنگ صبا.
- کلاین، پل (۱۳۸۰) راهنمای آسان تحلیل عاملی، ترجمه سید جلال صدرالسادات و اصغر مینایی تهران: انتشارات سمت.
- گال، مردیت و بورگ، والتر و گال، جویس (۱۳۸۲) روش‌های تحقیق کمی و کیفی در علوم تربیتی و روان‌شناسی، ترجمه احمد رضا نصر و همکاران. تهران: سمت.
- منصورفر، کریم (۱۳۸۴) روش‌های آماری، تهران: دانشگاه تهران.
- منصورفر، کریم (۱۳۸۵) روش‌های پیشرفته آماری همراه با برنامه‌های کامپیوتری، تهران: انتشارات دانشگاه تهران.
- مونرو، باربارا هازارد (۱۳۸۹)، روش‌های آماری در پژوهش مراقبت‌های بهداشتی و کاربرد SPSS در تحلیل داده‌ها، ترجمه و تالیف انوشیروان کاظم نژاد، محمدرضا حیدری و رضا نوروز زاده، تهران: جامعه نگر-سالمی.
- میزر، لاورنس اس و گامست، گلن و گارینو، ا.جی. (۱۳۹۱) پژوهش چندمتغیری کاربردی (طرح و تفسیر)، ترجمه حسن پاشا شریفی و دیگران، تهران: رشد.
- میلر، دلبرت چارلز (۱۳۸۰) راهنمای سنجش در تحقیقات اجتماعی. ترجمه هوشنگ ناییبی. تهران: نشر نی.
- ناییبی، هوشنگ (۱۳۸۷)، برنامه کامپیوتری آمار در علوم اجتماعی SPSS ، تهران: انتشارات سروش.
- ناییبی، هوشنگ (۱۳۸۸)، آمار توصیفی برای علوم اجتماعی، تهران: انتشارات سمت.
- هومن، حیدرعلی (۱۳۸۴)، مدل‌یابی معادلات ساختاری با کاربرد نرم افزار لیزرل، تهران: انتشارات سمت.

- Agresti, A., and Finlay, B.(1997) Statistical methods for the social sciences. Upper Saddle River, NJ: Prentice-Hall.
- Baron, R. M., & Kenny, D. A.(1986). The moderator°mediator variable distinction in social psychological research: Conceptual, strategic, and statistical considerations. *Journal of Personality and Social Psychology*.
- Bryman, Alan. & Cramer, Duncan (2001) Quantitative data analysis with SPSS 10 for Windows: a guide for social scientists.
- Burns,B.Robert. & Burns,P.Robert. & Burns, Richard. (2008) Business Research Methods and Statistics Using SPSS. SAGE.
- Field, Andy P. (2000) Discovering Statistics Using SPSS for Windows: Advanced Techniques for the Beginner. SAGE.
- Hair, JF, Black, WC, Babin, BK & Anderson, RE (2010), *Multivariate Data Analysis*, 7th edn, Pearson, New York.
- Nunnally, JC (1978), *Psychometric Theory*, McGraw Hill, New York.
- Schafer, JL (1997), *Analysis of Incomplete Multivariate Data*. Chapman & Hall, New York.
- Straub, D, Boudreau, M-C & Gefen, D (2004), ‘Validation Guidelines for IS Positivist Research’, *Communications of the Association for Information Systems*, vol. 13, no. 24, pp. 380–427.
- Tabachnick, B & Fidell, LS (2001), *Using Multivariate Statistics*, 4th edn, HarperCollins, New York.
- Tabachnick, B. G., & Fidell, L. S. (2007). *Using multivariate statistics* (5th ed.). Boston: Allyn and Bacon.

پیوست: پرسشنامه

پرسشنامه عوامل مؤثر بر موفقیت دانشجویان کارشناسی ارشد در درس SPSS

- به استاد درس SPSS خود در مقطع ارشد، در هرکدام از ویژگی های زیر از ۰ تا ۱۰ چه نمره ای می دهید؟

شایستگی استاد B	نمره (دور نمره مد نظر خود خط بکشید)
۱. فن بیان مناسب	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۲. تسلط علمی	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۳. ظاهر مرتب و آراسته	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۴. انرژی و نشاط بالا در تدریس	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۵. وقت شناسی در شروع و پایان کلاس	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۶. ارائه جذاب مباحث درسی	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۷. استقبال از سوالات دانشجویان	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۸. نظم در ارائه مطالب درسی	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۹. معرفی منابع آموزشی مناسب (کتاب، مقاله و ...)	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۱۰. ارائه تکالیف درسی مناسب	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۱۱. استفاده از وسایل کمک آموزشی (ویدئو ...)	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۱۲. رعایت انصاف در نمره دهی به دانشجویان	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸
۱۳. اخلاق خوب در برخورد با دانشجویان	۰ ۱ ۲ ۳ ۴ ۵ ۶ ۷ ۸

• سوالات جمعیت شناختی

جنس C: ☐ زن ☐ مرد	تعداد اعضای خانواده A: نفر	سن A: سال
تحصیلات پدر B: ☐ سیکل و پایین تر ☐ دیپلم ☐ فوق دیپلم ☐ لیسانس ☐ فوق لیسانس و بالاتر ☐		
قومیت C: ☐ ترک ☐ فارس ☐ کرد ☐ سایر ☐		
درآمد خانواده B: زیر ۵۰۰ هزار تومان ☐ کمتر از ۱ میلیون ☐ کمتر از ۱.۵ میلیون ☐ ۱.۵ میلیون و بالاتر ☐		

- نظر خود را درباره هر کدام از عبارت های زیر بیان نمایید

کاملاً مخالفم	کاملاً موافقم	بی نظر	مخالقم	موافقم	کاملاً موافقم
					۱. من به آموختن برنامه SPSS علاقه مند هستم
					۲. من برای آموختن برنامه SPSS وقت کافی ندارم
					۳. برای موفقیت در رشته خود باید برنامه SPSS را به خوبی بیاموزم
					۴. برنامه SPSS کاربرد زیادی در رشته ما دارد.
					۵. حاضرم بخشی از وقت خود را صرف آموختن برنامه SPSS کنم

• اطلاعات تحصیلی:

- (۱) معدل مقطع کارشناسی؟ **A**:
- (۲) نمره درس SPSS در مقطع کارشناسی؟ **A**:
- (۳) نمره درس SPSS در مقطع کارشناسی ارشد؟ **A**:
- (۴) به طور متوسط چند ساعت در هفته به مطالعه درس SPSS (در مقطع ارشد) اختصاص می دادید؟ **A**: ساعت
- (۵) نمره بهره هوشی (IQ) شما چند است؟ **A**:
- (۶) به نظر شما شیوه تدریس استاد تا چه حد در آموختن بهتر درس SPSS مؤثر است؟ **B**
خیلی کم ☐ کم ☐ متوسط ☐ زیاد ☐ خیلی زیاد ☐
- (۷) آیا در کل به آموختن درس آمار علاقه مند هستید؟ **C**
بله ☐ خیر ☐

نکته:

A = سطح سنجش فاصله ای/نسبی

B = سطح سنجش ترتیبی (رتبه ای)

C = سطح سنجش اسمی

آموزش های تکمیلی نرم افزار SPSS

و آموزش نرم افزارهای LISREL ، Amos ، PLS

Minitab ، Eviews و... در سایت زیر در دسترس شماست:



www.kharazmi-statistics.ir

www.Kharazmi-statistics.ir

رامین کریمی