

راهنمای کاربردی نرم افزار

SPSS

با تاکید بر روش تحقیق و آمار

براساس سرفصل های مصوب آموزش عالی
همراه با نمونه تفاسیر متعدد از آزمون های آماری



۱۳۸۹

راهنمای کاربردی نرم افزار

SPSS

با تاکید بر روش تحقیق و آمار

براساس سرفصل های مصوب آموزش عالی
همراه با نمونه تفاسیر متعدد از آزمون های آماری

نویسنده:

مجید حیدری چروده

با همکاری

سیمین فروغ زاده



۱۳۸۹

سرشناسه	حیدری چرووده مجید، ۱۳۴۵
عنوان و نام پدیدآور	راهنمای کاربردی نرم افزار SPSS با تاکید بر روش تحقیق و آمار براساس سرفصل‌های مصوب آموزش عالی... / مجید حیدری با همکاری سیمین فروغزاده
مشخصات نشر	تهران: انتشارات جامعه‌شناسان، ۱۳۸۹.
مشخصات ظاهری	: ۲۲۴ ص.
شابک	: ۹۷۸-۶۰۰-۵۹۹۹-۲۵-۹
وضعیت فهرست-نویسی	: فیپا
موضوع	اس. پی. اس. تحت ویندوز
موضوع	: علوم اجتماعی- تحقیق- روش شناسی
موضوع	: علوم اجتماعی- روش‌های آماری- برنامه کامپیوتری
شناسه افزوده	فروغزاده- سیمین
رده‌بندی کنگره	: ۱۳۸۹ ز ۲ / ح ۹ HA۳۲
رده‌بندی دیویی	: ۵۱۹/۵۰۷۸۵۵۳۶۹
شماره کتابشناسی ملی	: ۲۱۵۲۹۸۴



جامعه‌شناسان

فروشگاه و نمایشگاه مرکزی: تهران، خیابان انقلاب، روبروی درب اصلی دانشگاه تهران،

پاساژ فروزنده، همکف، پ ۳۳۳، فروشگاه جامعه‌شناسان، تلفن: ۶۶۰۰۰۵۱

WWW.JAMEESHENASAN.COM / Email: info@jameeshenasan.com

راهنمای کاربردی نرم افزار SPSS با تاکید بر روش تحقیق و آمار

مجید حیدری چرووده

سیمین فروغزاده

چاپ اول ۱۳۸۹ تهران	تعداد: ۱۰۰۰ نسخه
قیمت: تومان	لیتوگرافی و چاپ: طیف‌نگار
شابک: ۹۷۸-۶۰۰-۵۹۹۹-۲۵-۹	ISBN: 978-600-5999-25-9
همه حقوق چاپ و نشر برای ناشر محفوظ است.	Printed in Iran

فهرست مطالب

بخش اول: مقدمات

درس اول: معرفی و راه اندازی برنامه SPSS

۲۲	۱- مقدمه
۲۲	۲- راه اندازی برنامه SPSS
۲۳	۳- کاربرگها
۲۴	۴- تعریف متغیرها
۲۴	۵- نام متغیر (NAME):
۲۵	☐ نوع متغیر (TYPE):
۲۶	☐ تعداد ارقام صحیح (WIDTH):
۲۶	☐ تعداد ارقام اعشار متغیر (DECIMALS):
۲۶	☐ برچسب متغیرها (LABLE)
۲۷	☐ برچسب مقادیر (VALUES)
۲۸	☐ مقادیر نامشخص متغیر (MISSING):
۲۹	☐ اندازه ی ستون نمایش متغیر (COLUMNS):
۲۹	☐ چگونگی نمایش مقادیر در هر خانه (ALIGN):
۳۰	☐ مقیاس اندازه گیری متغیر (MEASURE)
۳۰	۵- آشنایی با ساختار فایل داده
۳۲	۶- نحوه ضبط و فراخوانی فایل داده ها

درس دوم: کدگذاری

۳۸	۱- مقدمه
۳۸	۲- نحوه ی کدگذاری
۳۹	☐ کدگذاری سوالات بسته پاسخ
۴۰	☐ کدگذاری سوالات باز پاسخ
۴۰	☐ کدگذاری سوالات باز پاسخ بر اساس رویکرد تئوریک
۴۱	☐ کدگذاری سوالات باز پاسخ بر اساس رویکرد تجربی

- کدگذاری مقادیر نامشخص (MISSING VALUES) ۴۲
- کدگذاری سوالات بی پاسخ یا پاسخهای مبهم ۴۲
- کدگذاری پاسخ های بی نظر یا نمی دانم ۴۷

درس سوم: ویرایش داده ها

- ۱-مقدمه ۵۲
- ۲-اصول ویرایش داده ها ۵۲
- آیا مقادیر غیر معتبر در جدول خروجی وجود دارد؟ ۵۲
- آیا مقادیر نامشخص (MISSING VALUE) تعریف شده است؟ ۵۴
- آیا برچسب ها (برچسب متغیرها، برچسب مقادیر معتبر و نیز برچسب مقادیر نامشخص) به صورت صحیح تعریف شده است؟ ۵۵
- آیا نحوه کد گذاری داده ها صحیح است؟ ۵۶

درس چهارم: کدگذاری مجدد و دسته بندی داده ها

- ۱- مقدمه ۶۰
- ۲- حالت های مختلف دستور ۶۰
- ۳- کاربردهای دستور ۶۰
- استفاده از دستور RECODE برای اصلاح یا تغییر نحوه ی کد گذاری ۶۲
- استفاده از دستور RECODE برای طبقه بندی متغیرهای کمی یا کیفی ۶۲

درس پنجم: سنجش پایایی

- ۱- مقدمه ۷۰
- ۲- سنجش پایایی از حیث همسازی درونی به روش آلفای کرانباخ ۷۰

درس ششم: شاخص سازی

- ۱- مقدمه ۸۲
- ۲- شاخص و شاخص سازی ۸۲
- ۳- روش شاخص سازی از متغیرهای ناهم دامنه ۸۴
- ۴- شاخص سازی با استفاده از دستور COMPUTE ۸۶

بخش دوم: توصیف نتایج به کمک روش های آمار توصیفی

درس هفتم: توصیف یک متغیری نتایج

- ۱-مقدمه..... ۹۵
- ۲-توصیف یک متغیری..... ۹۸
- توصیف یک متغیر کیفی..... ۹۹
- توصیف یک متغیر کمی..... ۱۰۱

درس هشتم: توصیف دو متغیری نتایج

- ۱-مقدمه..... ۱۱۰
- ۲-توصیف یک متغیر کیفی بر حسب یک متغیر کیفی..... ۱۱۰
- ۳-توصیف یک متغیر کمی بر حسب یک متغیر کیفی..... ۱۱۳

بخش سوم: آمار استنباطی - روابط همبستگی و علی

- ۱- مقدمه..... ۱۱۸

درس نهم: رابطه همبستگی و علی بین متغیرهای اسمی

- ۱- مقدمه..... ۱۲۲
- ۲- آزمون و ضریب مناسب برای سنجش رابطه علی بین متغیرهای اسمی..... ۱۲۳
- آزمون کی دو و ضریب تاثیر لامدا..... ۱۲۳
- ۳- آزمون و ضریب مناسب برای سنجش رابطه همبستگی بین متغیرهای اسمی..... ۱۳۰
- آزمون کی دو و ضرایب همبستگی فی و وی کرامر..... ۱۳۰

درس دهم: رابطه همبستگی و علی بین متغیرهای رتبه ای

- ۱- مقدمه..... ۱۳۶
- ۲- ضرایب مناسب برای سنجش رابطه علی بین متغیرهای رتبه ای..... ۱۳۷
- ضریب تاثیر D سامرز..... ۱۳۷
- ۳- ضرایب مناسب برای سنجش رابطه همبستگی بین متغیرهای رتبه ای..... ۱۴۰
- ضرایب همبستگی تای بی و تای سی کندال..... ۱۴۰

درس یازدهم: رابطه همبستگی بین متغیرهای کمی

- ۱- مقدمه ۱۴۶
- ۲- آزمون فرض غیرنرمال بودن متغیرهای کمی ۱۴۶
- ۳- ضرایب مناسب برای سنجش رابطه همبستگی بین متغیرهای کمی ۱۴۸
 - ضریب همبستگی پیرسون ۱۴۸
 - ضریب همبستگی تفکیکی ۱۵۰
 - ضریب همبستگی اسپیرمن ۱۵۳

بخش چهارم: مقایسه‌ها

درس دوازدهم: مقایسه میانگین‌ها

- ۱- مقدمه ۱۶۰
- ۲- مقایسه میانگین یک متغیر کمی در دو گروه ۱۶۲
 - آزمون برابری میانگینها برای نمونه‌های مستقل ۱۶۲
- ۳- مقایسه میانگین دو متغیر کمی در یک گروه ۱۶۶
 - آزمون برابری میانگینها برای نمونه‌های همبسته ۱۶۶
- ۴- مقایسه میانگین یک متغیر کمی با یک مقدار ثابت ۱۶۹
 - آزمون برابری میانگینها برای یک نمونه ۱۶۹

درس سیزدهم: تحلیل واریانس یک‌طرفه

- ۱- مقدمه ۱۷۴
- ۱- مقایسه میانگین یک متغیر کمی در بیش از دو گروه ۱۷۴
 - آزمون تحلیل واریانس یک‌طرفه ۱۷۴

درس چهاردهم: مقایسه متغیرهای رتبه‌ای

- ۱- مقدمه ۱۸۶
- ۲- مقایسه یک متغیر کیفی رتبه‌ای یا کمی غیرنرمال در دو گروه ۱۸۶
 - آزمون یو مان ویتنی ۱۸۶
- ۳- مقایسه دو متغیر کیفی رتبه‌ای یا کمی غیرنرمال با یکدیگر ۱۸۸

- آزمون ویلکاکسون.....۱۸۸
- ۴-مقایسه یک متغیر کیفی رتبه ای یا کمی غیر نرمال در بیش از دو گروه۱۹۰
- آزمون کراسکال والیس.....۱۹۰
- ۵-مقایسه چند متغیر کیفی رتبه ای یا کمی غیر نرمال در یک گروه۱۹۱
- آزمون فریدمن.....۱۹۱

درس پانزدهم: رگرسیون چندگانه خطی

- ۱- مقدمه۱۹۸
- ۲-رگرسیون چندگانه خطی۱۹۹

فهرست اشکال

شکل شماره (۱) چگونگی راه اندازی برنامه SPSS	۲۲
شکل شماره (۲) معرفی برنامه SPSS (کاربرگ داده ها DATA VIEW)	۲۳
شکل شماره (۳) معرفی برنامه SPSS کاربرگ متغیرها (VARIABLE VIEW)	۲۴
شکل شماره (۴) تصویر پنجره TYPE	۲۵
شکل شماره (۵) تعریف برچسب متغیرها (LABLE)	۲۷
شکل شماره (۶) تعریف برچسب مقادیر (VALUES)	۲۷
شکل شماره (۷) تعریف مقادیر نامشخص در پنجره MISSING VALUE	۲۸
شکل شماره (۸) اندازه های ستون و چگونگی نمایش مقادیر در کاربرگ داده ها	۲۹
شکل شماره (۹) فایل داده های WORLD95 در کاربرگ متغیرها (VARIABLE VIEW)	۳۰
شکل شماره (۹) فایل داده های WORLD 95 در کاربرگ داده ها (DATA VIEW)	۳۱
شکل شماره (۱۰) چگونگی ضبط فایل داده ها	۳۲
شکل شماره (۱۱) چگونگی فراخوانی فایل داده ها	۳۳
شکل شماره (۱۲) زیردستورهای مختلف منوی فایل و کاربردهای آنها	۳۴
شکل شماره (۱۳) انواع واکنش پاسخگو یان به یک سوال پرسشنامه	۴۲
شکل شماره (۱۴) نحوه تعریف مقادیر نامشخص به صورت سه کد متمایز در ستون MISSING VALUE	۴۳
شکل شماره (۱۵) نحوه تعریف مقادیر نامشخص در قالب یک بازه به علاوه یک کدمشخص در ستون MISSING VALUE	۴۴
شکل شماره (۱۶) مقادیر نامشخص تعریف شده متغیر "مدت اشتغال"	۴۴
شکل شماره (۱۷) تعریف مقادیر نامشخص متغیر "مدت اشتغال"	۴۵
شکل شماره (۱۸) اجرای دستور FREQUENCIES از کلیه متغیرهای موجود در فایل داده ها	۵۲
شکل شماره (۱۹) شیوه ی جستجوی مقادیر نامعتبر در فایل داده ها	۵۳
شکل شماره (۲۰) گزینه های موجود در دستور RECODE	۶۰
شکل شماره (۲۱) شیوه ی اجرای دستور RECODE درحالت SAM VARIABLE	
RECODE INTO	۶۱
شکل شماره (۲۲) شیوه ی اجرای دستور RECODE درحالت DIFFERENT	
RECODE INTO VARIABLE	۶۱
شکل شماره (۲۳) پنجره OLD AND NEW VALUES در دستور RECODE	۶۲

شکل شماره (۲۴) چگونگی طبقه بندی متغیر کمی یا کیفی با اجرای دستور RECODE در گزینه	۶۳
..... OLD AND NEW VALUES	
شکل شماره (۲۵) فایل داده های مقیاس "روحیه همکاری دانش آموزان".....	۷۱
شکل شماره (۲۶) شیوه ی اجرای دستور RELIABILITY برای سنجش پایایی مقیاس	
مرحله ۱.....	۷۲
شکل شماره (۲۷) شیوه اجرای دستور RELIABILITY برای سنجش پایایی مقیاس.مرحله	
۲.....	۷۲
شکل شماره (۲۸) تجزیه متغیر طی تعریف عملیاتی و شاخص سازی	۸۲
شکل شماره (۲۹) شیوه استاندارد سازی متغیر (تهیه نمره Z).....	۸۴
شکل شماره (۳۰) نمره استاندارد Z در کاربرگ متغیرها (VARIABLE VIEW).....	۸۵
شکل شماره (۳۱) نمره استاندارد Z در کاربرگ داده ها (DATA VIEW).....	۸۵
شکل شماره (۳۲) شیوه اجرای دستور COMPUTE برای شاخص سازی (مرحله ۱).....	۸۶
شکل شماره (۳۳) شیوه اجرای دستور COMPUTE برای ساختن شاخص (مرحله ۲).....	۸۷
شکل شماره (۳۴) شاخص جدید ساخته شده در کاربرگ متغیرها (VARIABLE VIEW).....	۸۸
شکل شماره (۳۵) شاخص جدید ساخته شده در کاربرگ داده ها (DATA VIEW).....	۸۸
شکل شماره (۳۶) نحوه ساختن شاخص از متغیرهای ناهموزن.....	۹۰
شکل شماره (۳۷) اشکال مختلف تحلیل داده ها در گزارش های پژوهشی.....	۹۷
شکل شماره (۳۸) دیاگرام انتخاب روش مناسب برای توصیف نتایج.....	۹۸
شکل شماره (۳۹) شیوه ی اجرای دستور FREQUENCIES برای توصیف متغیرهای کیفی.....	۹۹
شکل شماره (۴۰) انتخاب نمودار مناسب برای نمایش متغیر.....	۱۰۰
شکل شماره (۴۱) نمودار میله ای توزیع رشته تحصیلی دانش آموزان.....	۱۰۱
شکل شماره (۴۲) شیوه اجرای دستور DESCRIPTIVES برای توصیف	
متغیرهای کمی.....	۱۰۲
شکل شماره (۴۳) شیوه انتخاب شاخص های آماری مناسب برای متغیر.....	۱۰۲
شکل شماره (۴۴) شیوه انتخاب شاخص های آماری مناسب برای متغیر روحیه همکاری دانش	
آموزان در پنجره STATISTICS.....	۱۰۴
شکل شماره (۴۵) شاخص کمی روحیه همکاری دانش آموزان.....	۱۰۶
شکل شماره (۴۶) اجرای دستور CROSSTABS برای توصیف یک متغیر کیفی - کیفی.....	۱۱۰
شکل شماره (۴۷) استفاده از پنجره CELLS برای نمایش اعداد در جدول خروجی.....	۱۱۱
شکل شماره (۴۸) اجرای دستور MEANS برای توصیف یک متغیر کمی بر حسب متغیر کیفی.....	۱۱۳

شکل شماره (۴۹) استفاده از پنجره OPTION برای استفاده از شاخص های آماری مناسب برای توصیف متغیر کمی.....	۱۱۳
شکل شماره (۵۰) نحوه تصمیم گیری در باره روش آماری آزمون فرضیات دارای متغیرهای سطح اسمی.....	۱۲۳
شکل شماره (۵۱) شیوه اجرای آزمون کی دو (مرحله ۱).....	۱۲۴
شکل شماره (۵۲) شیوه اجرای آزمون کی دو (مرحله ۲).....	۱۲۵
شکل شماره (۵۳) شیوه اجرای آزمون کی دو (مرحله ۳).....	۱۲۶
شکل شماره (۵۴) شیوه اجرای آزمون کی دو.....	۱۳۰
شکل شماره (۵۵) راهنمای انتخاب آزمون برای رابطه متغیرهای کیفی.....	۱۳۷
شکل شماره (۵۶) انتخاب آزمونها و ضرایب مناسب.....	۱۳۸
شکل شماره (۵۷) انتخاب آزمونها و ضرایب مناسب.....	۱۴۱
شکل شماره (۵۸) راهنمای انتخاب آزمون مناسب برای رابطه متغیرهای کمی.....	۱۴۶
شکل شماره (۵۹) شیوه اجرای آزمون کولموگروف اسمیرنف.....	۱۴۷
شکل شماره (۶۰) شیوه اجرای آزمون پیرسون.....	۱۴۹
شکل شماره (۶۱) شیوه اجرای ضریب همبستگی تفکیکی (مرحله ۱).....	۱۵۱
شکل شماره (۶۲) شیوه اجرای ضریب همبستگی تفکیکی (مرحله ۲).....	۱۵۱
شکل شماره (۶۳) شیوه اجرای آزمون اسپیرمن.....	۱۵۳
شکل شماره (۶۴) راهنمای انتخاب آزمونها برای مقایسه گروه ها.....	۱۶۱
شکل شماره (۶۵) شیوه اجرای آزمون برابری میانگین ها برای نمونه های مستقل (T-TEST).....	۱۶۲
شکل شماره (۶۶) شیوه اجرای آزمون برابری میانگین ها برای نمونه های همبسته (PAIRED-SAMPLE).....	۱۶۶
شکل شماره (۶۷) شیوه اجرای آزمون برابری میانگین ها با یک مقدار ثابت (ONE-SAMPLE).....	۱۶۹
شکل شماره (۶۸) شیوه اجرای آزمون تحلیل واریانس یک طرفه (مرحله ۱).....	۱۷۵
شکل شماره (۶۹) شیوه اجرای آزمون تحلیل واریانس یک طرفه (مرحله ۲).....	۱۷۵
شکل شماره (۷۰) شیوه اجرای آزمون تحلیل واریانس یک طرفه (مرحله ۳).....	۱۷۶
شکل شماره (۷۱) شیوه اجرای آزمون یومان ویتنی.....	۱۸۷
شکل شماره (۷۲) شیوه اجرای آزمون ویلکاکسون.....	۱۸۸
شکل شماره (۷۳) شیوه اجرای آزمون کراسکال والیس.....	۱۹۰
شکل شماره (۷۴) شیوه اجرای آزمون فریدمن.....	۱۹۲
شکل شماره (۷۵) شیوه اجرای آزمون رگرسیون چند گانه خطی (مرحله ۱).....	۱۹۹

- شکل شماره (۷۶) شیوه اجرای آزمون رگرسیون چند گانه خطی (مرحله ۲) ۲۰۰
- شکل شماره (۷۷) شیوه اجرای آزمون رگرسیون چند گانه خطی (مرحله ۳) ۲۰۱
- شکل شماره (۷۸) نمودار هیستوگرام باقیمانده های استاندارد شده رگرسیون ۲۰۷
- شکل شماره (۷۹) نمودار احتمال نرمال بودن باقیمانده های استاندارد شده رگرسیون ۲۰۷

فهرست جداول

جدول شماره (۱): یک نمونه از تعریف صحیح متغیر.....	۴۵
جدول شماره (۲): یک نمونه از تعریف اشتباه متغیر.....	۴۶
جدول شماره (۳): یک نمونه اشتباه در تعریف متغیر.....	۴۷
جدول شماره (۴): توزیع فراوانی متغیر "نوع مدرسه" در حالت وجود مقادیر غیر معتبر در ستون داده‌های این متغیر.....	۵۳
جدول شماره (۵): نمونه دارای مقدار نامعتبر در تعریف مقادیر نامشخص.....	۵۴
جدول شماره (۶): نمونه بدون مقدار نامعتبر در تعریف مقادیر نامشخص.....	۵۵
جدول شماره (۷): توزیع فراوانی متغیر "مقطع تحصیلی" در حالت نادرست.....	۵۶
جدول شماره (۸): نمونه اشتباه در نحوه کدگذاری.....	۵۷
جدول شماره (۹): توزیع فراوانی کیفی "درصد باسوادان جهان در سال ۱۹۹۵".....	۶۵
جدول شماره (۱۰): توزیع فراوانی کمی "درصد باسوادان جهان در سال ۱۹۹۵".....	۶۶
جدول شماره (۱۱): خلاصه پردازش اطلاعات.....	۷۳
جدول شماره (۱۲): میزان آلفای کرانباخ درمقیاس "روحیه همکاری دانش آموزان".....	۷۳
جدول شماره (۱۳): توصیف میانگین و انحراف معیارگویه های مقیاس "روحیه همکاری دانش آموزان".....	۷۴
جدول شماره (۱۴): مقادیر آماره های مقیاس "روحیه همکاری دانش آموزان".....	۷۵
جدول شماره (۱۵): مقادیر آماره های مقیاس "روحیه همکاری دانش آموزان".....	۷۶
جدول شماره (۱۶): توصیف شاخص "میزان توسعه یافتگی".....	۸۵
جدول شماره (۱۷): خلاصه پردازش اطلاعات.....	۸۸
جدول شماره (۱۸): توزیع فراوانی شاخص کمی "روحیه همکاری دانش آموزان".....	۸۹
جدول شماره (۱۹): توزیع فراوانی رشته تحصیلی دانش آموزان.....	۱۰۰
جدول شماره (۲۰): توصیف شاخص های آماری "روحیه همکاری دانش آموزان".....	۱۰۳
جدول شماره (۲۱): توصیف شاخص های آماری "روحیه همکاری دانش آموزان".....	۱۰۵
جدول شماره (۲۲): ارزیابی دانش آموزان از کیفیت غذای مدرسه برحسب جنسیت.....	۱۱۲
جدول شماره (۲۳): توصیف کیفیت غذای مدرسه بر حسب جنسیت دانش آموزان.....	۱۱۲
جدول شماره (۲۴): وضعیت روحیه همکاری دانش آموزان بر حسب جنسیت.....	۱۱۴
جدول شماره (۲۵): روش های آماری مناسب برای آزمون فرضیات یا سوالات تحقیق.....	۱۱۹
جدول شماره (۲۶): روشهای آماری مناسب برای آزمون فرضیات یا سوالات تحقیق دارای متغیرهای اسمی.....	۱۲۲

جدول شماره (۲۷): خلاصه پردازش اطلاعات.....	۱۲۶
جدول شماره (۲۸): جدول توافقی بین دو متغیر ترجیح شغلی و جنسیت.....	۱۲۷
جدول شماره (۲۹): تعیین سطح معناداری متغیرهای مورد سنجش.....	۱۲۸
جدول شماره (۳۰): تعیین شدت تاثیر جنسیت بر ترجیحات شغلی.....	۱۲۹
جدول شماره (۳۱): جدول توافقی بین دو متغیر ترجیح شغلی و رشته تحصیلی.....	۱۳۱
جدول شماره (۳۲): تعیین سطح معناداری متغیرهای مورد سنجش.....	۱۳۱
جدول شماره (۳۳): تعیین جهت و شدت تاثیر جنسیت بر ترجیحات شغلی.....	۱۳۲
جدول شماره (۳۴): روش‌های آماری مناسب برای آزمون فرضیات یا سوالات تحقیق دارای	
متغیرهای رتبه‌ای.....	۱۳۶
جدول شماره (۳۵): خلاصه پردازش اطلاعات.....	۱۳۸
جدول شماره (۳۶): جدول توافقی بین دو متغیر وضعیت اقتصادی خانواده و تقید به روزه داری	
دانش آموزان.....	۱۳۹
جدول شماره (۳۷): تعیین سطح معناداری، شدت و جهت تاثیر وضعیت اقتصادی بر تقید به روزه	
داری دانش آموزان.....	۱۴۰
جدول شماره (۳۸): خلاصه پردازش اطلاعات.....	۱۴۲
جدول شماره (۳۹): جدول توافقی بین دو متغیر تصویر مذهبی از خود و تقید به روزه داری دانش	
آموزان.....	۱۴۲
جدول شماره (۴۰): تعیین سطح معناداری و شدت و جهتین متغیرهای مورد سنجش.....	۱۴۳
جدول شماره (۴۱): بررسی وضعیت توزیع متغیرهای "درصد شهرنشینی و نرخ باروری".....	۱۴۷
جدول شماره (۴۲): نتایج سنجش همبستگی متغیرهای نرخ رشد جمعیت و درصد شهرنشینی.....	۱۵۲
جدول شماره (۴۳): همبستگی تفکیکی "نرخ مرگ و میرکودکان" و "درصد شهرنشینی" با	
کنترل "درصد باسوادان".....	۱۵۲
جدول شماره (۴۴): نتایج سنجش همبستگی متغیرهای "درصد شهرنشینی" و "نرخ باروری".....	۱۵۴
جدول شماره (۴۵): حالت‌های مختلف مقایسه در فرضیه‌ها یا سوالات تحقیق از حیث سطوح	
سنجش متغیرها.....	۱۶۰
جدول شماره (۴۶): مقایسه میانگین شاخص کمی گرایش به موسیقی در بین زنان و مردان.....	۱۶۳
جدول شماره (۴۷): آزمون مقایسه میانگین شاخص کمی گرایش به موسیقی در بین زنان و مردان.	
جدول شماره (۴۸): مقایسه میانگین حقوق فعلی در بین زنان و مردان.....	۱۶۴
جدول شماره (۴۹): آزمون مقایسه میانگین حقوق فعلی در بین زنان و مردان.....	۱۶۵
جدول شماره (۵۰): مقایسه وضعیت آماره‌های دو متغیر "حقوق فعلی" و "حقوق بدو استخدام".....	۱۶۷
جدول شماره (۵۱): وضعیت همبستگی بین حقوق فعلی و حقوق بدو استخدام.....	۱۶۸

جدول شماره (۵۲): نتیجه آزمون معنی داری برابری میانگین ها برای دو متغیر (حقوق فعلی و حقوق بدو استخدام).....	۱۶۸
جدول شماره (۵۳): توصیف متغیر امید به زندگی مردان	۱۷۰
جدول شماره (۵۴): نتیجه آزمون معنی داری برابری میانگین ها با یک مقدار ثابت	۱۷۰
جدول شماره (۵۵): ویژگی انواع آزمونهای تعقیبی	۱۷۷
جدول شماره (۵۶): توصیف متغیر "حقوق فعلی کارکنان در طبقات شغلی مختلف"	۱۷۸
جدول شماره (۵۷): آزمون مقایسه واریانس حقوق کارمندان در رده های شغلی سه گانه (دفتری، نگهبان، مدیر)	۱۷۸
جدول شماره (۵۸): آزمون مقایسه میانگین حقوق کارمندان در رده های شغلی سه گانه (دفتری، نگهبان، مدیر)	۱۷۹
جدول شماره (۵۹): آزمون تعقیبی توکی برای بررسی تفاوت های زوجی طبقات سه گانه شغلی	۱۸۰
جدول شماره (۶۰): دسته بندی گروه های مورد مقایسه بر اساس آزمون توکی	۱۸۰
جدول شماره (۶۱): آزمونهای مربوط به مقایسه متغیرهای رتبه ای از حیث هدف تحقیق	۱۸۶
جدول شماره (۶۲): نتایج آزمون یو مان ویتنی در باره مقایسه رضایت مشتریان به تفکیک جنس ...	۱۸۷
جدول شماره (۶۳): نتایج آزمون ویلکاکسون در زمینه مقایسه گرایش جوانان به موسیقی کلاسیک و رپ	۱۸۹
جدول شماره (۶۴): نتایج آزمون کراسکال-والیس در زمینه مقایسه ارزیابی گروه های سنی مختلف از موسیقی کلاسیک	۱۹۰
جدول شماره (۶۵): نتایج آزمون فریدمن در زمینه مقایسه رضایت مشتریان از قیمت و تنوع کالا و خدمات یک فروشگاه	۱۹۲
جدول شماره (۶۶): معرفی متغیرهای مستقل و وابسته و روش بکارگرفته شده در انجام رگرسیون	۲۰۱
جدول شماره (۶۷): ضرایب همبستگی و تعیین در گامهای مختلف	۲۰۲
جدول شماره (۶۸): نتایج آزمون معناداری رگرسیون در گام های مختلف	۲۰۳
جدول شماره (۶۹): متغیرهای وارد شده در معادله رگرسیون و ضرایب تاثیر آنها در گامهای مختلف	۲۰۴
جدول شماره (۷۰): متغیرهای خارج شده از معادله رگرسیون	۲۰۵
جدول شماره (۷۱): آماره های مربوط به باقیمانده ها	۲۰۶

مقدمه

امروزه از پژوهش به عنوان کلید توسعه در کشور یاد می‌شود و گرایش به آن در عرصه‌های مختلف، به ویژه در قلمرو مسائل فرهنگی و اجتماعی روندی صعودی دارد. پژوهشگران این حوزه به خوبی به این نکته واقف اند که انجام تحقیقات کمی در اکثر موارد بدون استفاده از نرم‌افزارهای تخصصی امکان پذیر نیست. یکی از مهمترین نرم‌افزارهای تخصصی که عمدتاً در تحقیقات کمی از آن استفاده می‌شود، نرم‌افزار^۱ SPSS است.^۲

تا کنون کتاب‌های متعددی در ارتباط با نحوه استفاده از نرم‌افزار SPSS نوشته شده است. وجه تمایز اصلی و بسیار مهم کتاب حاضر با اکثریت قریب به اتفاق کتاب‌های یاد شده، در محوریت آموزشی و پژوهشی کتاب حاضر و تلفیق روش تحقیق، آمار و SPSS با یکدیگر است. نویسندگان این کتاب با تجربه حاصل از سال‌ها تدریس SPSS و تجزیه و تحلیل تعداد زیادی تحقیق و پایان‌نامه، سعی کرده اند از طریق ارائه‌ی مثال‌های متعدد در بخش‌های مختلف، نمونه‌های متنوعی از تفسیرهای آماری را به خوانندگان ارائه دهند.

استفاده از این نرم‌افزار نیازمند تسلط نسبی در برخی زمینه‌هاست که از آن جمله می‌توان به مهارت در اپراتوری سیستم عامل ویندوز، آشنایی کافی با روش‌های آماری در علوم انسانی و اجتماعی و نیز توانمندی در روش تحقیق در این علوم است. با این حال کتاب حاضر به سبکی نوشته شده است که تا حدی خلاء دانشی و مهارتی دانشجویان در برخی از زمینه‌های فوق را می‌تواند جبران کند.

این کتاب شامل چهار بخش است. بخش اول علاوه بر معرفی این نرم‌افزار و آشنایی با ساختار فایل داده‌ها و کدگذاری، به سه فرآیند مهم ویرایش، دسته‌بندی و شاخص سازی داده‌ها پرداخته است. رکن اصلی بخش دوم کتاب، بیان روش‌های آمار توصیفی به کمک نرم‌افزار SPSS و تفسیر خروجی‌های آن است و بخش سوم و چهارم کتاب به نحوه‌ی استفاده از این نرم‌افزار در روش‌های آمار استنباطی و تفسیر خروجی‌های آن اختصاص دارد. لازم به ذکر است که تصاویر این کتاب اگر چه براساس نسخه‌ی هفده نرم‌افزار SPSS تنظیم شده است

1- Statistical Package for Social Sciences

۲- لازم به ذکر است که کاربرد کامپیوتر در پژوهش منحصر به SPSS نیست. چه نرم افزارهای مهم دیگری نیز وجود دارد که متناسب با نوع نیازها می‌توان از آن استفاده نمود. نظیر نرم‌افزار NCCS برای نمونه‌گیری، LISREL برای ارزیابی معادلات ساختاری SAS, MINITAB و R برای انجام تجزیه و تحلیل‌های آماری.

ولی برای خوانندگانی نیز که از نسخه‌های قبلی این نرم افزار (نسخه ده و بعد از آن) استفاده می‌کنند، می‌تواند مفید باشد.

در پایان بر خود لازم می‌دانیم که از خانم‌ها مریم فردانیان، ندا رضوی‌زاده و مریم اسکافی که زحمت بازخوانی متن را بر عهده داشتند و با پیشنهادهای ارزشمند خود به غنای این کتاب افزودند، صمیمانه تشکر و قدردانی کنیم.

حیدری چروده^۱ -

فروغ‌زاده^۲

تابستان ۱۳۸۹

۱- مجید حیدری چروده پژوهشگر گروه علوم اجتماعی جهاددانشگاهی مشهد و مدرس دانشگاه‌های پیام نور و آزاد اسلامی واحد مشهد است.

۲- سیمین فروغ‌زاده عضو هیات علمی گروه علوم اجتماعی جهاد دانشگاهی مشهد و مدرس دانشگاه‌های پیام نور و امام حسین مشهد است

بخش اول

مقدمات

درس اول

معرفی برنامه SPSS

در این درس انتظار می‌رود دانشجوی

- ۱- نحوه‌ی راه‌اندازی برنامه spss را بداند.
- ۲- کاربرگ داده‌ها را بشناسد و مشخصات آن را بداند.
- ۳- کاربرگ متغیرها را بشناسد و ویژگی‌های ده‌گانه هر متغیر را بداند.
- ۴- بتواند متغیرهای مربوط به یک پرسشنامه را بطور کامل در برنامه spss تعریف کند.
- ۵- نحوه فراخوانی و ضبط فایل‌های spss را بداند.

۲۲ □ راهنمای کاربردی نرم افزار SPSS

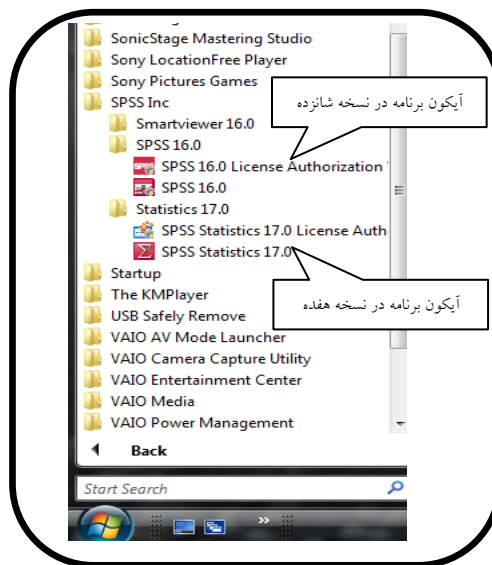
۱- مقدمه

این نرم افزار به گونه ای طراحی شده است که متخصصین رشته های مختلف می توانند از آن استفاده کنند. گستردگی و پیچیدگی آن، به اندازه ای است که کمتر کسی می تواند ادعای تسلط بر تمامی ابعاد و جوانب آن را داشته باشد. در این درس برای پیشگیری از تأثیرات منفی پیچیدگی و گستردگی فوق بر ذهن و اندیشه دانشجوی مبتدی، به معرفی مختصر و مفید این برنامه و کاربرگ های آن پرداخته و با ساختار فایل داده در این نرم افزار آشنا می شوید.

۲- راه اندازی^۱ برنامه SPSS^۲

آسان ترین روش برای راه اندازی SPSS تحت ویندوز، کلیک بر روی دکمه ی Start و جستجو در برنامه های زیرشاخه ی All Programs است. برای شروع کار، همانند شکل زیر روی برنامه SPSS for Windows^۳ در پوشه SPSS کلیک کنید.

شکل شماره (۱) چگونگی راه اندازی برنامه spss

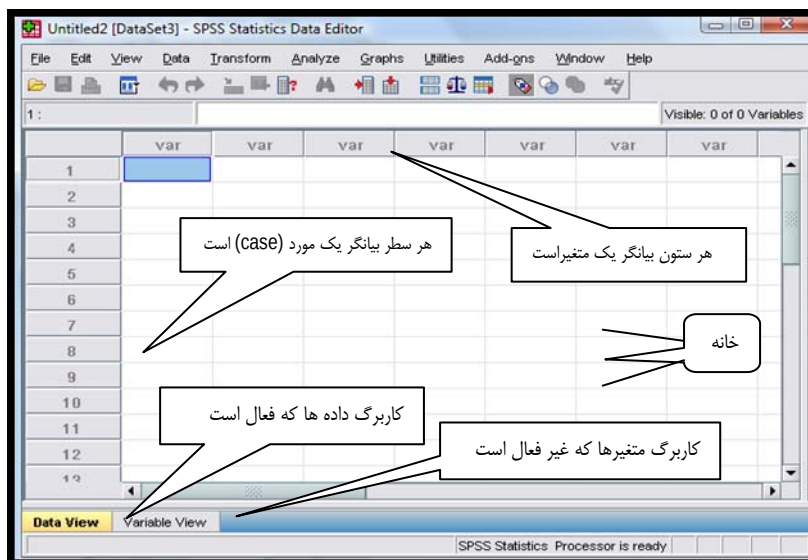


- ۱- برخی از اشکالات رایجی که در هنگام نصب برنامه SPSS یا بعد از نصب آن پیش می آید در ضمیمه شماره یک توضیح داده شده است. توصیه می شود ابتدا آن را بخوانید.
- ۲- دستوراتی که در ادامه برای راه اندازی SPSS ذکر می شود برای سیستم عامل Windows بکار می رود. اگر SPSS را در سیستم دیگری راه اندازی کنید، دستورات ممکن است متفاوت باشند.
- ۳- عنوان برنامه و پوشه ی SPSS در نسخه های مختلف آن اندکی با هم تفاوت دارد که با قدری تأمل به راحتی قابل شناسایی است.

۳- کاربرگ‌ها

با اجرای برنامه SPSS با شکل (۲) مواجه می‌شوید. همانطور که ملاحظه می‌شود در گوشه‌ی سمت چپ و پایین صفحه‌ی اولیه‌ی SPSS، دو کاربرگ وجود دارد که عبارت است از: کاربرگ داده‌ها (Data View) و کاربرگ متغیرها (Variable view). با کلیک بر روی عنوان هر یک از این دو کاربرگ می‌توانید آنها را فعال کرده و بین آنها رفت و برگشت داشته باشید. روشن بودن عنوان هر کاربرگ، نشانه فعال بودن آن است.

شکل شماره (۲) معرفی برنامه spss (کاربرگ داده‌ها Data View)



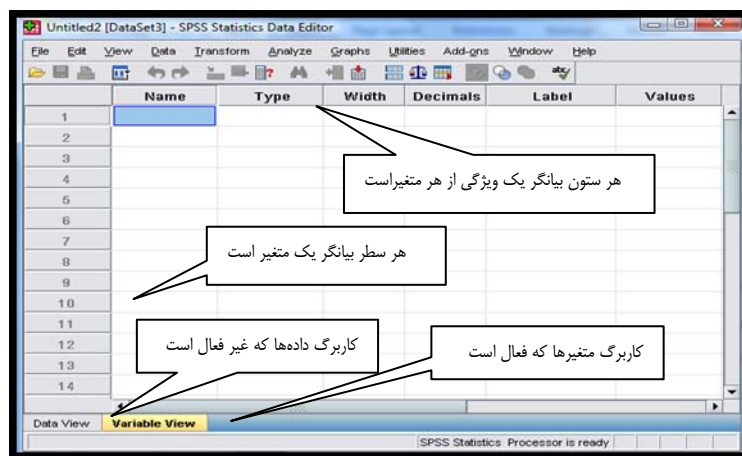
کاربرگ داده‌ها (Data View). مجموعه‌ای از ردیف‌ها و ستون‌ها است که در هر ردیف آن، یک مورد (Case) و در هر ستون آن، یک متغیر (Variable) وجود دارد. Case‌ها واحدهایی هستند که تحقیق در باره‌ی آنها انجام می‌شود و متغیرهای تحقیق به آنها نسبت داده می‌شود. بدیهی است که یک case الزاماً یک شخص نیست. مثلاً اگر روی کتاب‌ها مطالعه می‌کنید، هر مورد یا case که در ردیف‌های پنجره کاربرگ داده‌ها قرار می‌گیرد، یک کتاب است. اگر بر روی قیمت فرش‌ها مطالعه می‌کنید، هر فرش یک مورد به حساب می‌آید. بنابراین مورد، واحدی است که اندازه‌گیری‌ها بر روی آن انجام می‌شود. لازم به ذکر است که محل تلاقی ستون و ردیف، خانه نامیده می‌شود.

۲۴ □ راهنمای کاربردی نرم افزار SPSS

شکل زیر کاربرد متغیرها (Variable view) را نشان می دهد. در این کاربرد، سطرها بیانگر متغیرها و ستونها نشانگر ویژگی های متغیر (مثل نام متغیر، نوع متغیر، برچسب متغیر و...) است.

برای تعریف و شناساندن یک متغیر به نرم افزار SPSS، باید ده ویژگی را در این کاربرد برای هر متغیر تعریف کرد که در ادامه به تشریح آنها می پردازیم.

شکل شماره (۳) معرفی برنامه spss کاربرد متغیرها (Variable view)



۴- تعریف متغیرها

تعریف متغیرها کاری است که باید در کاربرد متغیرها انجام شود. برای تعریف کامل یک متغیر باید ویژگی های آن را به شرح زیر مشخص نمایید.

• نام متغیر (Name):

در این گزینه باید نام متغیرهای تحقیق خود را بنویسید. توصیه می شود:

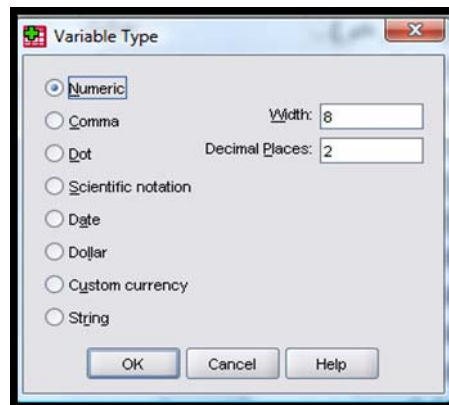
- نام متغیر مختصر و متناسب با شماره سؤال متناظر در پرسشنامه باشد. مثلاً سؤال ۴ پرسشنامه بهتر است نام X4 تعریف شود و چنانچه سؤال ۵ پرسشنامه خود شامل دو سؤال فرعی باشد، بهتر است متغیرهای مربوطه به نام های X5.1 و X5.2 تعریف شود. تناظر نام متغیر با شماره سؤال پرسشنامه موجب تسهیل و تسریع در فرآیند ورود، خواندن و ویرایش داده ها خواهد شد.

معرفی و راه‌اندازی برنامه SPSS □ ۲۵

- از فارسی‌نویسی نام متغیرها خودداری شود. چون فارسی‌نویسی نام متغیر، عملیات استفاده از دستورنویسی به شیوه دستی در فایل^۱ Syntax را با دشواری مواجه می‌سازد.
- یک متغیر به عنوان شماره ردیف پرسشنامه یا سند در ابتدای هر فایل بنویسید و به شماره ردیف خود نرم‌افزار اکتفا نکنید. مزیت متغیر شماره ردیف این است که شما بوسیله آن می‌توانید داده‌ها را به شکل‌های مختلف مرتب کنید. در حالیکه با شماره ردیف پیش‌فرض نرم‌افزار، چنین کاری عملی نیست.
- **نوع متغیر (Type):**

در این قسمت باید نوع متغیر را برای نرم‌افزار مشخص کنید، برای این منظور روی دومین ستون جدول (type) کلیک کنید تا پنجره‌ای به شکل زیر ظاهر گردد:

شکل شماره (۴) تصویر پنجره type



در این پنجره شما می‌توانید با توجه به سطح سنجش متغیر در پرسشنامه یکی از گزینه‌ها را انتخاب کنید. مهمترین انواع متغیرها عبارت‌اند از:

- **Numeric:** این گزینه برای داده‌های کمی بکار می‌رود. در این گزینه می‌توانید تعداد ارقام صحیح و اعشار متغیر خود را انتخاب کنید. پیشنهاد می‌شود از این پیش‌فرض برای همه‌ی متغیرها استفاده شود.^۲

۱- Syntax پنجره‌ای برای نوشتن یا ضبط دستورات آماری در نرم‌افزار spss است.

۲- این توصیه بیشتر به این دلیل است که نحوه‌ی کدگذاری اکثریت قریب به اتفاق متغیرها، در غالب تحقیقات به صورت عددی است. با این حال می‌توان بسته به نوع متغیرها از سایر پیش‌فرض‌ها نیز استفاده نمود.

• **String:** این گزینه برای متغیرهایی انتخاب می‌شود که اطلاعات آنها غیرعددی است (یعنی قرار است به جای عدد، حروف وارد می‌شود). برای مثال اگر برای ورود اطلاعات مربوط به جنسیت افراد نمونه، به جای کدهای ۱ و ۲ از f و m استفاده کنیم، باید نوع متغیر خود را حرفی (String) انتخاب کنید.

• **Dollar:** این گزینه برای آن دسته از داده‌های کمی به کار می‌رود که واحد اندازه‌گیری آنها، دلار است. مثلاً در این حالت برای نرم‌افزار مشخص می‌شود که واحد اندازه‌گیری متغیر موردنظر، دلار است و حداکثر ۶ رقم صحیح و تا ۲ رقم اعشار خواهد داشت.

• **Date:** این گزینه برای آن دسته از داده‌های کمی به کار می‌رود که حاوی تاریخ (سال، ماه و روز) باشند و بسته به نوع اطلاعات می‌تواند انواع مختلفی نیز داشته باشد.

• **تعداد ارقام صحیح (Width):**

در این قسمت می‌توانید تعداد رقم‌های صحیح متغیر خود را با کلیک کردن روی آن افزایش و یا کاهش دهید.

• **تعداد ارقام اعشار متغیر (Decimals):**

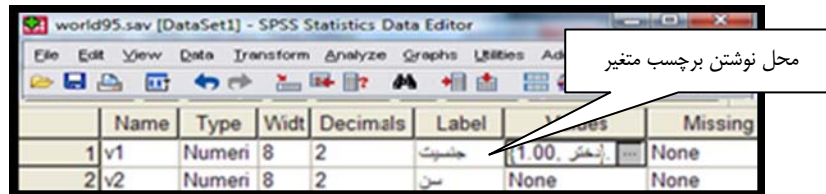
با کلیک کردن در این قسمت می‌توانید تعداد ارقام اعشار را برای متغیر موردنظر افزایش و یا کاهش دهید. توجه داشته باشید که تعداد ارقام صحیح همیشه باید بیشتر از تعداد ارقام اعشاری باشد.

• **برچسب متغیرها (Label):**

برچسب، یک کد شناسایی است که از آن طریق می‌توانید عنوان متغیر را تعریف کنید. مثلاً می‌توانید عنوان جنسیت یا sex را برای متغیر x1 تعریف کنید. برای این منظور، با کلیک کردن بر روی آن، خانه آماده نوشتن می‌شود و می‌توانید برچسب را بنویسید. توصیه می‌شود برای مقاصد آموزشی برچسب متغیرها به صورت فینگلیش^۱ تعریف شود. چون تجربه‌ی آموزشی چند ساله نشان داده است که فارسی نویسی برچسب‌ها در جریان نقل و انتقال داده‌ها بین کامپیوترهای مختلف، عملاً دشواری‌های زیادی ایجاد می‌کند.

۱- مثلاً به جای جنسیت، بنویسید: jensiat

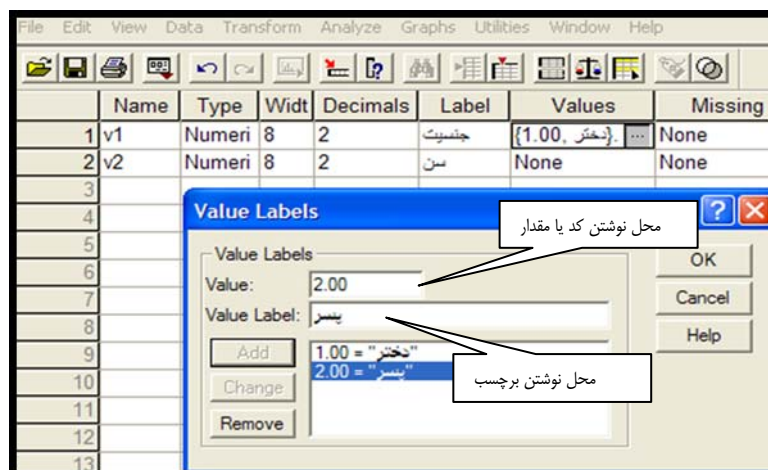
شکل شماره (۵) تعریف برچسب متغیرها (Label)



• برچسب مقادیر (Values):

در این گزینه باید به معرفی مقوله‌ها یا مقادیر هر متغیر (بویژه متغیرهای کیفی) اقدام کنید، چنانچه روی گوشه‌ی سمت راست این گزینه کلیک کنید، پنجره‌ی زیر ظاهر خواهد شد:

شکل شماره (۶) تعریف برچسب مقادیر (Values)



در این پنجره می‌توانید مقوله‌های هر متغیر را معرفی کنید. مثلاً اگر بخواهید برای متغیر جنس کد ۱ را به دختران و کد ۲ را به پسران نسبت بدهید. برای این منظور مراحل ذیل را باید به ترتیب انجام دهید.

– در فضای خالی Value کد ۱ را وارد کنید و پس از آن در فضای خالی گزینه Value Label، کلمه دختر یا فینگلیش آن را (Dokhtar) وارد کنید. سپس گزینه Add را کلیک کنید. در این حالت در مستطیل پایین عبارت «دختر = ۱» اضافه می‌شود.

– سپس در فضای خالی Value کد ۲ را وارد کنید و در قسمت value label عبارت پسر یا فینگلیش آن را (pesar) وارد کنید و روی add کلیک کنید.

گزینه ی Remove برای پاک کردن و Change برای تغییر دادن برچسب های مربوط به مقوله های یک متغیر است.^۱

• **مقادیر نامشخص متغیر (Missing):**

گاهی اوقات پاسخگویان به برخی پرسش ها جواب نمی دهند یا جواب بی ربط می دهند. مقادیر مربوط به این دسته از پرسش ها را اصطلاحاً مقادیر نامشخص می نامند. در این گزینه امکان معرفی این دسته از مقادیر فراهم شده است. با کلیک روی گزینه ی Missing Value شکل زیر ظاهر می شود. همانگونه که ملاحظه می شود چهار گزینه وجود دارد:

شکل شماره (۷) تعریف مقادیر نامشخص در پنجره Missing Value



همانگونه که در شکل شماره (۷) مشاهده می کنید، با انتخاب گزینه ی Discrete Missing values می توانید سه کد را برای مقادیر نامشخص تعریف کنید.^۲ مثلاً:

- کد ۰.۱ برای بی جواب ها
- کد ۰.۲ برای پاسخ های بی ربط و مبهم
- کد ۰.۳ برای آن دسته از سؤالاتی است که شامل حال پاسخگو نمی شود^۳

۱- توضیحات تفصیلی نحوه کدگذاری داده ها، در درس چهارم " کدگذاری داده ها و دسته بندی " آمده است.

۲- نحوه کدگذاری داده ها به تفصیل در درس بعدی توضیح داده شده است.

۳- دقت کنید برای تعریف مقداری که دارای ممیز است، از نقطه (.) به جای (/) استفاده کنید. چون نرم افزار علامت (/) را نمی خواند و آن را به عنوان خطا در نظر می گیرد.

یا در صورت لزوم می‌توانید با انتخاب گزینه‌ی Range plus one ... دامنه‌ای از کدها را برای مقادیر نامشخص تعریف کنید. مثلاً از ۰.۱ تا ۰.۵ را به علاوهِی کد ۰.۹۹ با انتخاب گزینه Discrete value به عنوان مقادیر نامشخص تعریف کنید.^۱

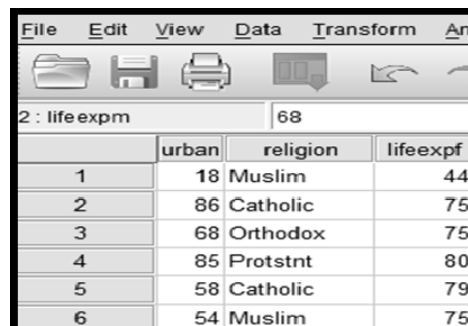
• اندازه‌ی ستون نمایش متغیر (Columns):

این گزینه عرض ستون‌ها را برای نمایش متغیر مشخص می‌کند. با کلیک کردن بر روی آن می‌توانید جای بیشتر یا کمتری به متغیرهای موردنظر اختصاص دهید. مثلاً برای نمایش متغیر جنسیت یک یا دو ستون کفایت می‌کند و نیازی به اختصاص ستون‌های بیشتر نیست.

• چیدمان مقادیر در هر خانه (Align):

این ستون چیدمان مقادیر را در هر ستون مشخص می‌کند. انتخاب گزینه‌های Left, Right و Center باعث می‌شود اعداد به ترتیب در سمت چپ، راست و وسط خانه قرار گیرد.

شکل شماره (۸) اندازه‌ی ستون و چگونگی نمایش مقادیر در کاربرگ داده‌ها



	urban	religion	lifeexpf
1	18	Muslim	44
2	86	Catholic	75
3	68	Orthodox	75
4	85	Protstnt	80
5	58	Catholic	79
6	54	Muslim	75

مثلاً در شکل فوق‌اندازه متغیرهای urban, religion, lifeexpf به ترتیب ۵، ۱۲ و ۸ است و متغیرهای urban و lifeexpf به صورت چپ‌چین و religion در حالت راست‌چین تنظیم شده است.

۱- توضیحات مربوط به انواع داده‌های نامشخص و نحوه برخورد محقق با آنها به تفصیل در فصول بعدی کتاب آمده است

• مقیاس اندازه‌گیری متغیر (Measure):

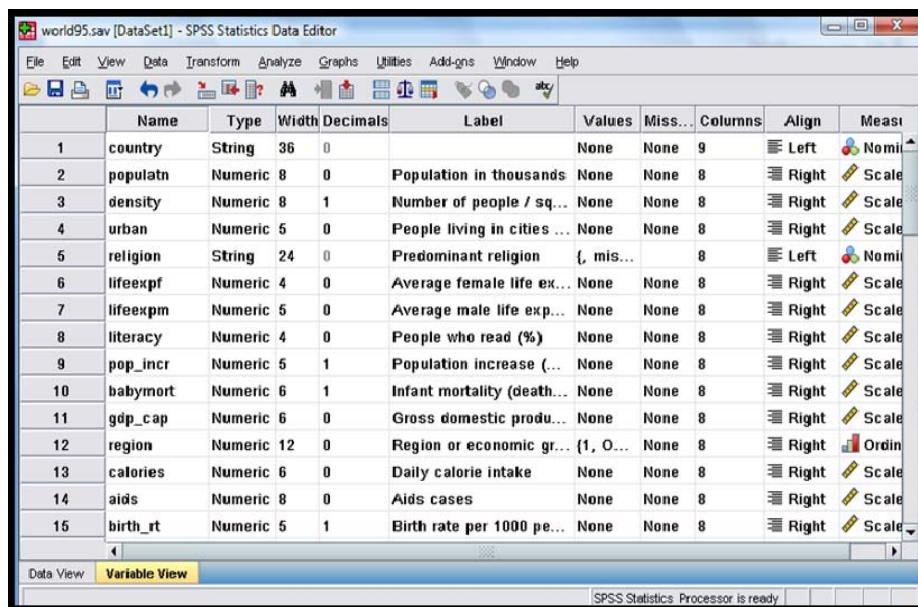
در این گزینه، مقیاس اندازه‌گیری متغیر مشخص می‌شود. به این ترتیب که Nominal برای متغیرهایی با مقیاس اسمی، Ordinal برای متغیرهایی با مقیاس رتبه‌ای و Scale برای متغیرهایی با مقیاس نسبی / فاصله‌ای بکار برده می‌شود.

به این ترتیب، پس از تعریف این ده ویژگی برای متغیر نخست، می‌توانید در ردیف دوم به معرفی متغیر بعدی خود بپردازید. پس از معرفی همه متغیرها در کاربرگ متغیرها (Variable view)، بر روی کاربرگ داده‌ها (Data View) کلیک کنید و چنانکه خواهید دید، متغیرهایی را که در کاربرگ متغیرها تعریف کرده‌اید، در این کاربرگ به صورت یک ستون ساخته شده است. در این حالت نرم‌افزار آماده ورود اطلاعات است و می‌توانید همه‌ی اطلاعات (کدها) مربوط به هر متغیر را در خانه‌ها وارد کنید.

۵- آشنایی با ساختار فایل داده

با تکمیل کاربرگ متغیرها (Variable view) و ورود اطلاعات در کاربرگ داده‌ها (View Data)، فایل داده‌ها آماده می‌شود. در اینجا برای آشنایی بیشتر با ساختار فایل داده‌ها، یک نمونه از این قبیل فایل‌ها معرفی می‌شود.

شکل شماره (۹) فایل داده‌های world95 در کاربرگ متغیرها (Variable view)

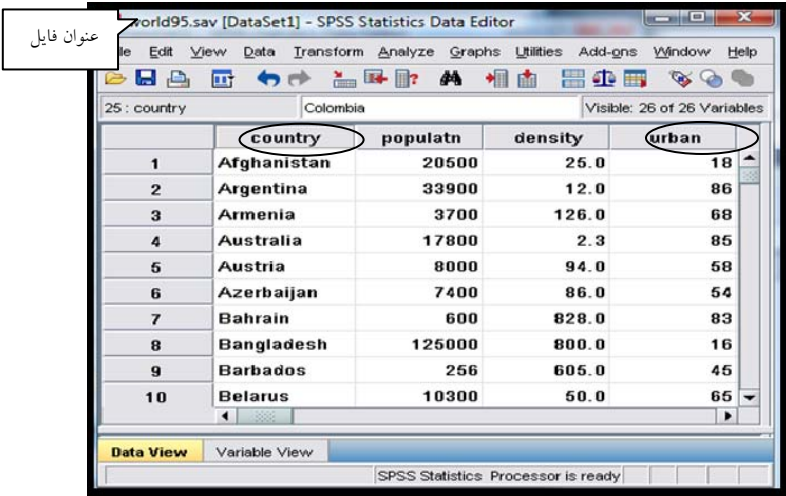


	Name	Type	Width	Decimals	Label	Values	Miss...	Columns	Align	Measur
1	country	String	36	0		None	None	9	Left	Nominal
2	populatn	Numeric	8	0	Population in thousands	None	None	8	Right	Scale
3	density	Numeric	8	1	Number of people / sq...	None	None	8	Right	Scale
4	urban	Numeric	5	0	People living in cities ...	None	None	8	Right	Scale
5	religion	String	24	0	Predominant religion	{, mis...		8	Left	Nominal
6	lifeexpf	Numeric	4	0	Average female life ex...	None	None	8	Right	Scale
7	lifeexpm	Numeric	5	0	Average male life exp...	None	None	8	Right	Scale
8	literacy	Numeric	4	0	People who read (%)	None	None	8	Right	Scale
9	pop_incr	Numeric	5	1	Population increase (...)	None	None	8	Right	Scale
10	babymort	Numeric	6	1	Infant mortality (death...	None	None	8	Right	Scale
11	gdp_cap	Numeric	6	0	Gross domestic produ...	None	None	8	Right	Scale
12	region	Numeric	12	0	Region or economic gr...	{1, 0...	None	8	Right	Ordinal
13	calories	Numeric	6	0	Daily calorie intake	None	None	8	Right	Scale
14	aids	Numeric	8	0	Aids cases	None	None	8	Right	Scale
15	birth_rt	Numeric	5	1	Birth rate per 1000 pe...	None	None	8	Right	Scale

معرفی و راه‌اندازی برنامه SPSS □ ۳۱

عنوان فایل داده‌ی فوق، world 95.sav است که جزء نمونه فایل‌های موجود در پوشه spss است که در هنگام نصب این نرم‌افزار بر روی سیستم ریخته می‌شود.^۱ همانطور که در شکل شماره ۹ ملاحظه می‌شود، در این فایل ۱۵ متغیر نشان داده شده است، اولین آن به نام country و از نوع string است که با حداکثر گنجایش ۳۶ حرف، بدون برچسب متغیر و مقدار و بدون تعریف مقادیر نامشخص و با امکان هشت حرف برای نمایش در کاربرگ داده‌ها، به صورت چپ چین و در سطح اسمی (Nominal) تعریف شده است. یا متغیر ردیف چهارم (Urban) که از نوع کمی است، با حداکثر گنجایش ۵ رقم، که فاقد رقم اعشاری است، برچسب این متغیر درصد شهرنشینی (People Living Cities) و فاقد برچسب مقادیر است، مقادیر نامشخص این متغیر تعریف نشده، اما امکان نمایش در قالب ۸ رقم را داشته است و به صورت راست چین نوشته شده و سطح سنجش آن نسبی (Scale) است. حال با کلیک بر روی کاربرگ داده‌ها، با تصویری همانند شکل شماره ۹ مواجه خواهید شد.

شکل شماره (۹) فایل داده‌های world 95 در کاربرگ داده‌ها (Data View)



	country	populatr	density	urban
1	Afghanistan	20500	25.0	18
2	Argentina	33900	12.0	86
3	Armenia	3700	126.0	68
4	Australia	17800	2.3	85
5	Austria	8000	94.0	58
6	Azerbaijan	7400	86.0	54
7	Bahrain	600	828.0	83
8	Bangladesh	125000	800.0	16
9	Barbados	256	605.0	45
10	Belarus	10300	50.0	65

۱- نمونه‌های متعددی از فایل‌های داده، هنگام نصب نرم‌افزار spss به طور خودکار در پوشه‌ی مربوط به این نرم‌افزار یا زیرپوشه‌ی samples (در نسخه‌های شانزده و هفده) ریخته می‌شود و در هر زمانی قابل استفاده است.

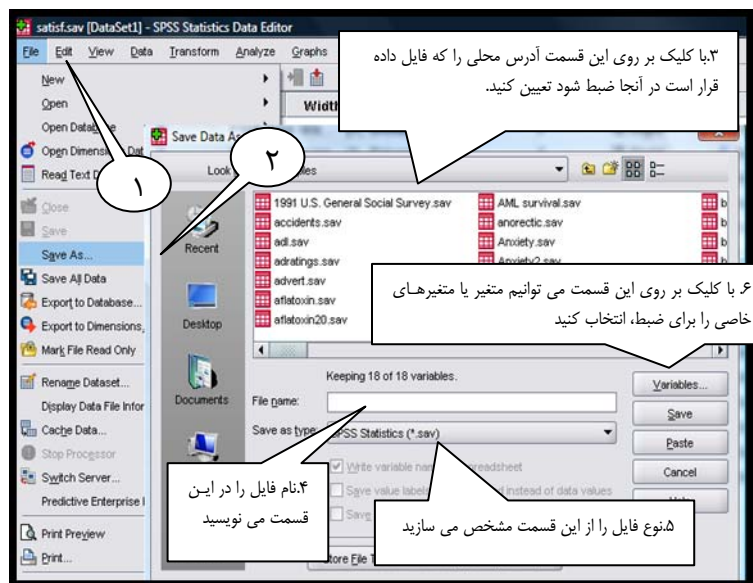
۳۲ □ راهنمای کاربردی نرم افزار SPSS

همانطور که ملاحظه می شود ستون های این کاربرگ، همان متغیرهایی هستند که قبلاً در کاربرگ متغیرها (Variable view) تعریف شده اند. به عنوان مثال ستون اول، داده های مربوط به اولین متغیر (Country) را نشان می دهد که به صورت حرفی (string) وارد شده است و متغیر چهارم (Urban) هم به صورت کمی وارد شده است.

۶- نحوه ضبط و فراخوانی فایل داده ها

پس از اینکه همه ی پاسخ های پاسخگویان وارد نرم افزار شد، تهیه فایل داده ها به اتمام می رسد و فایل داده ها باید به نامی که متناسب با موضوع تحقیق باشد، ضبط شود. برای این کار همان گونه که در تصویر زیر نشان داده شده است، باید گزینه ی save as را از منوی فایل انتخاب کنید. توصیه می شود فایل داده ی هر تحقیقی را در پوشه ی مخصوص به خودش قرار دهید تا پیدا کردن آنها از میان انبوهی از فایل های داده که به مرور زمان با آنها سروکار پیدا می کنید، آسان باشد. ضمناً از فارسی نویسی نام فایل داده و نیز پوشه خودداری کنید، چون در برخی سیستم های عامل، اجرای این دسته از فایل ها با دشواری مواجه می شود.

شکل شماره (۱۰) چگونگی ضبط فایل داده ها

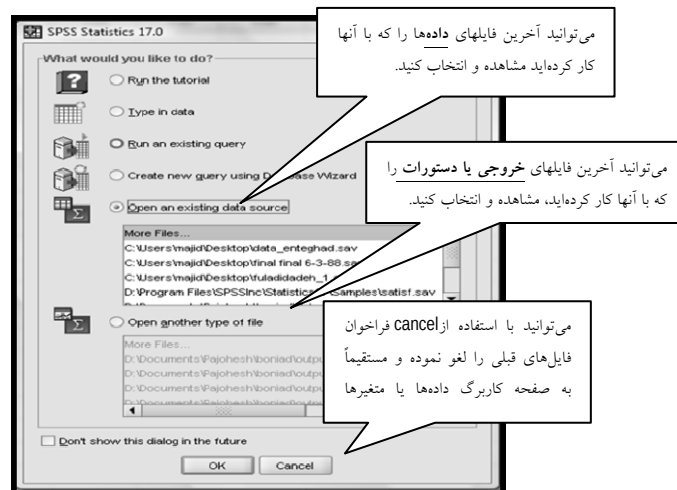


معرفی و راه‌اندازی برنامه SPSS □ ۳۳

چنانچه قصد استفاده از فایل داده‌ها را در محیط‌های دیگری نظیر Excel داشته باشید، می‌توانید نوع فایل خود را با کلیک بر روی قسمت مورد اشاره در تصویر، متناسب با نرم‌افزار اکسل که با پسوند xls ضبط می‌شوند انتخاب کنید^۱.

لازم به ذکر است که پس از ضبط فایل داده و خروج از برنامه spss، چنانچه مجدداً این برنامه را اجرا کنید، به صورت پیش‌فرض چند مورد از آخرین فایل‌های داده یا خروجی یا دستورات را که در این نرم‌افزار بر روی آنها کار شده است، در یک پنجره به کاربر نشان داده می‌شود و در این حالت دیگر نیازی به جستجوی آخرین فایل‌ها در کامپیوتر نیست به شکل زیر توجه کنید.

شکل شماره (۱۱) چگونگی فراخوانی فایل داده‌ها

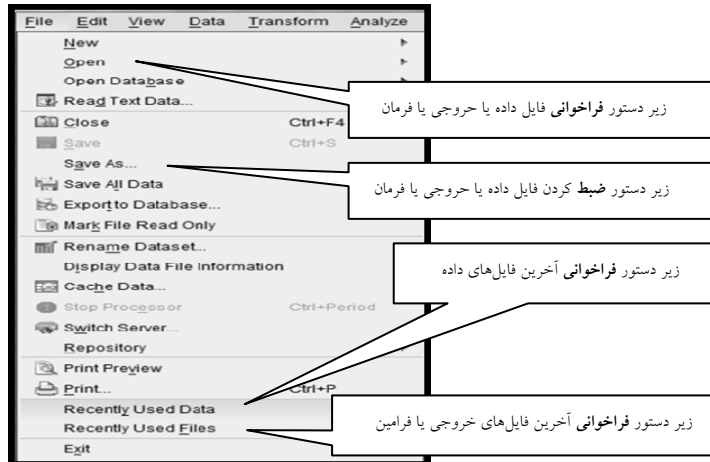


ضمناً عناوین آخرین فایل‌هایی داده‌ای که کاربر از آن استفاده کرده است در منوی فایل در زیر دستور Recently Used Data وجود دارد. علاوه بر آن آخرین فایل‌های خروجی یا فرامین هم در زیر دستور Recently Used Files مشخص است.

۱- گاهی اوقات از شما خواسته می‌شود داده‌هایتان را به گونه‌ای ضبط کنید که توسط نرم‌افزارهای دیگر نظیر اکسل یا اکسس یا سایر ویرایش‌های برنامه spss قابل خواندن باشد. در چنین شرایطی لازم است شما فایل داده‌هایتان را با فرمتی متناسب با نرم‌افزار مورد نظر ضبط کنید. لازم به ذکر است که برنامه spss می‌تواند فایل‌های داده‌ای که توسط نرم‌افزارهای مذکور نوشته شده است را نیز بخواند.

۳۴ □ راهنمای کاربردی نرم افزار SPSS

شکل شماره (۱۲) زیر دستورهای مختلف منوی فایل و کاربردهای آنها



تمرینات درس اول

تمرین ۱: فایل Survey Sample را از پوشه SPSS فراخوانی کنید و با توجه به کاربرگ متغیرهای این فایل، به سؤالات زیر پاسخ دهید.

الف- ویژگی‌های ده‌گانه متغیرهای Marital, Age, Edu را توضیح دهید.

ب- در متغیر Race کدهای ۲ و ۳ چه مفهومی دارند؟

ج- مقادیر نامشخص متغیر income چگونه تعریف شده‌اند؟

د- در متغیر Race، نژادهای سفید و سیاه چه کدی دارند؟

د- نوع، برچسب، مقادیر نامشخص و سطح سنجش متغیر ردیف نهم (born) را مشخص کنید.

تمرین ۲: برنامه spss را اجرا کنید و در کاربرگ متغیرها عملیات زیر را انجام دهید:

الف- متغیر v1 را با برچسب دین به صورت حرفی و در سطح اسمی تعریف کنید.

ب- متغیر v2 را به صورت عددی، با گنجایش ۳ ستون، برچسب گروه سنی و برچسب مقادیر ۱ (زیر ۳۰ سال) و ۲ (بالای ۳۰ سال)، در سطح رتبه‌ای و به صورت وسط چین و با امکان نمایش ۱۰ ستون تعریف کنید.

ج- متغیر v3 را با برچسب وزن، گنجایش ۵ رقم و ۳ رقم اعشار، راست چین و در سطح نسبی تعریف کنید.

د- به کاربرگ داده‌ها بروید و داده‌های مربوط به افراد زیر را وارد کنید:

د-۱- در سطر شماره یک: مسلمان (Muslim) ۲۷ ساله ای که ۷۱.۷۵۶ گرم وزن دارد.

د-۲- در سطر شماره دو: غیرمسلمان (Unmuslim) ۳۳ ساله ای که ۶۴ کیلوگرم وزن دارد.

د-۳- در سطر شماره سه: مسلمان ۵۵ ساله ای که ۵۱.۳۵۱ گرم وزن دارد.

ه- در پنجره‌ی کاربرگ داده‌ها، از منوی View گزینه Value Label را انتخاب کنید. این کار را چند بار انجام دهید. با اینکار چه تغییری در نمایش داده‌ها پیش می‌آید؟ چرا در ستون مربوط به متغیر وزن (v3) هیچ تغییری ایجاد نمی‌شود؟

و- از منوی فایل گزینه‌ی save as را انتخاب کرده و فایل داده‌ی فوق را به نام data_2 ضبط نمایید.

۳۶ □ راهنمای کاربردی نرم افزار SPSS

تمرین ۳: عملکرد هر یک از کلیدهای میانبر شکل زیر را توضیح دهید.



تمرین ۴: چگونه می توانیم اطلاعات مربوط به یک یا چند واحد نمونه، یا مشخصات مربوط به تعدادی از متغیرها را پاک، یا کپی کنیم؟

درس دوم

کدگذاری

در پایان این درس انتظار می‌رود دانشجو

- ۱- مفهوم کدگذاری را بداند و کارکردهای کد را تجزیه و تحلیل کند.
- ۲- بتواند سؤالات بسته پاسخ تک پاسخی و چند پاسخی را کدگذاری کند.
- ۳- بتواند سؤالات بازپاسخ را براساس دو رویکرد تجربی و تئوریک کدگذاری کند.
- ۴- انواع مقادیر نامشخص را بشناسد و نحوه کدگذاری آنها را بداند.
- ۵- تأثیر معرفی یا عدم معرفی مقادیر نامشخص را بر نتایج تحقیق توضیح دهد.
- ۶- نحوه کدگذاری پاسخهای بی‌نظر یا نمی‌دانم را بداند.

۱- مقدمه

کدگذاری که به معنای نسبت دادن کدهای عددی یا حرفی به پاسخهای مندرج در فرم گردآوری اطلاعات (اعم از پرسشنامه، مصاحبه نامه و...) است، یکی از مراحل اجتناب ناپذیر پژوهشهای علمی بویژه در حوزهی تحقیقات اجتماعی و فرهنگی است. انجام عملیات کدگذاری شرط لازم ورود داده ها به نرم افزار SPSS است.

۲- نحوه ی کدگذاری

عملیات کدگذاری برای معرفی مقوله های متغیر کیفی به کار می رود. واضح است که در متغیرهای کمی که پاسخها از نوع عدد و رقم است، دیگر نیازی به کدگذاری نیست.^۱ برای مثال پاسخ افراد به سؤال "سن شما چقدر است؟" یا "چند فرزند دارید؟" که غالباً به صورت عدد است، به کدگذاری نیازی ندارد، اما در مورد متغیرهایی مثل "جنسیت" یا "نوع شغل" که دارای پاسخهای کیفی هستند، برای ورود به نرم افزار SPSS نیاز به کدگذاری دارند مثلاً به صفت مرد کد ۱ و به صفت زن کد ۲، و یا در مورد متغیر شغل، به کارگر کد ۱، کشاورز کد ۲، کارمند کد ۳ و.... نسبت می دهیم.

کدگذاری در فرایند تحقیق کارکردهای مختلفی دارد. کمترین کارکرد کدها این است که بیانگر تمایز هستند. به عبارت دیگر از نحوه ی کدگذاری مشاغل می فهمیم که شغل مربوط به کد ۱ با شغل مربوط به کد ۲ تفاوت دارد. این کارکرد منحصر به کدگذاری متغیرهایی است که سطح سنجش آنها اسمی^۲ است. گفتنی است که کدها علاوه بر تمایز می توانند کارکرد رتبه بندی نیز داشته باشند. مثلاً متغیر تحصیلات را با صفات و کدهای بی سواد (کد ۱)، ابتدایی (کد ۲)، راهنمایی (کد ۳)، متوسطه (کد ۴) و دانشگاهی (کد ۵) در نظر بگیرید. در اینجا کدهای ۴ و ۵ نه تنها به تحصیلات متفاوتی مربوط می شوند (کارکرد تمایزبخشی)، بلکه علاوه بر آن گویای این واقعیت نیز هستند که تحصیلات متناظر با کد ۵ بیشتر از تحصیلات متناظر با کد ۴ است (کارکرد رتبه بندی). این کارکرد هم منحصر به متغیرهایی است که در سطح رتبه ای^۳ هستند.

۱- گفتنی است که متغیرهای دارای سطوح سنجش اسمی و رتبه ای در دسته ی متغیرهای کیفی بوده و سطوح

سنجش فاصله ای و نسبی از جمله متغیرهای کمی محسوب می شوند

۲ - Nominal

3- Ordinal

درس دوم - کدگذاری □ ۳۹

قاعده کلی در کدگذاری متغیرهای کیفی این است که برای هر پاسخ یک کد در نظر بگیرید. با نوشتن کدهای متناظر با هر پاسخ بر روی فرم گردآوری اطلاعات (اعم از پرسشنامه یا مصاحبه نامه یا مشاهده نامه یا...)، کدنامه ساخته می‌شود. کدنامه در حقیقت راهنمای کدگذاری است. بدیهی است که یک کدنامه‌ی شفاف می‌تواند عملیات کدگذاری را تسریع و یکدست کند و در عین حال بر دقت نتایج بیافزاید.

از آنجا که نحوه کدگذاری سؤالات باز و بسته با یکدیگر متفاوت است ابتدا کدگذاری سؤالات بسته پاسخ را تشریح می‌کنیم، سپس به کدگذاری سؤالات باز پاسخ می‌پردازیم.

• کدگذاری سؤالات بسته پاسخ

سؤالات بسته پاسخ به سؤالاتی اطلاق می‌شود که پاسخگو باید از میان گزینه‌های پیشنهادی محقق، گزینه یا گزینه‌های موردنظر خود را به عنوان پاسخ انتخاب کند و معمولاً حق اظهار نظری در خارج از چارچوب پیشنهادی محقق ندارد. این دسته از سؤالات را می‌توان به دو دسته‌ی فرعی تقسیم کرد:

الف) سؤالات بسته‌ای که پاسخگو فقط حق انتخاب یک گزینه را دارد. مثلاً:

وضعیت تأهل شما چگونه است؟ □ مجرد □ متأهل □ مطلقه □ بیوه □ جدا شده

در این نوع سؤالات می‌توان از کدهای ۱ (برای مجرد)، ۲ (برای متأهل)، ۳ (برای مطلقه)، ۴ (برای بیوه)، ۵ (برای جدا شده) استفاده کرد. طبیعی است که در این قبیل کدگذاری، تعداد کدها تابعی از تعداد گزینه‌ها است. مثلاً اگر فردی در پاسخ به این سؤال گزینه‌ی مطلقه را علامت بزند، کد ۳ به او تعلق می‌گیرد.

ب) سؤالات بسته‌ای که پاسخگو حق انتخاب بیش از یک پاسخ را دارد. مثلاً:

اوقات فراغت خود را چگونه می‌گذرانید؟

□ ورزش □ مهمانی □ استراحت □ مطالعه □ خیابانگردی □ تفریح □ سایر

در این نوع سؤالات بهتر است از کدگذاری صفر و یک استفاده شود. بدین معنی که پاسخگویانی که در مقابل هر یک از گزینه‌ها علامت زده اند، کد ۱ و به آنها که علامت نزده‌اند، کد صفر نسبت داده می‌شود.

تفاوت این سبک کدگذاری با کدگذاری سؤالات بسته پاسخ نوع الف، این است که محقق به تعداد گزینه‌ها (مثلاً ورزش، مهمانی، استراحت، مطالعه و...) متغیر در پیش‌رو خواهد داشت. به این معنی که باید هر یک از این گزینه‌ها را همانند یک متغیر تلقی کند.

• کدگذاری سؤالات بازپاسخ

همانطور که قبلاً ذکر شد، سؤالات بازپاسخ، سؤالاتی هستند که پاسخگو در آنها محدود به انتخاب گزینه خاصی نیست، بلکه می‌تواند نظر خود را بر حسب علایق و اطلاعات خود به تفصیل یا به اختصار در اختیار پرسشگر قرار دهد، خواه در قالب مصاحبه و به صورت شفاهی باشد و یا به صورت پرسشنامه کتبی. به سؤالات زیر دقت کنید:

۱. لطفاً عنوان شغل خود را بنویسید.
 ۲. چه مدت است که در این محله سکونت دارید؟
 ۳. نام آخرین کتاب غیردرسی که خوانده‌اید چه بود؟
 ۴. به نظر شما چرا برخی از مردم به مقررات راهنمایی و رانندگی مخصوص عابر پیاده توجه نمی‌کنند؟
- برای کدگذاری پاسخ این دسته از سؤالات دو رویکرد وجود دارد: رویکرد تئوریک و رویکرد تجربی.

■ کدگذاری سؤالات بازپاسخ براساس رویکرد تئوریک

طبق این رویکرد، باید با توجه به مبانی نظری تحقیق، اقدام به تهیه کدنامه برای این دسته از سؤالات کنید. برای مثال اگر قصد تحقیق در باره میزان رعایت بهداشت شهروندان را داشته باشید و پیش‌فرض‌تان این باشد که نوع شغل افراد در میزان رعایت بهداشت موثر است و با اتکا به مبانی نظری و این منطق که افرادی که در حیطه بهداشت و درمان شاغل هستند، در مقایسه با افرادی که مشاغل آموزشی دارند، به میزان بیشتری به مقررات بهداشتی پایبندی دارند و شاغلین در بخش‌های آموزشی در مقایسه با افرادی که دارای مشاغل فنی هستند، تقید بیشتری به مقررات بهداشت دارند، برای کدگذاری پاسخ‌های این سؤال باید به شکل زیر عمل شود:

۱. مشاغل بهداشتی و درمانی (نظیر پزشک، پرستار، ماما، تزریقاتی، کارشناس آزمایشگاه طبی و....) = کد ۱
 ۲. مشاغل آموزشی (نظیر معلم، استاد دانشگاه، مدرس خوشنویسی، مدیر یا معاون مراکز آموزشی و....) = کد ۲
 ۳. مشاغل فنی (نظیر مکانیک، جوشکار، بنا و...) = کد ۳
- از سوی دیگر ممکن است در تحقیقی قصد مطالعه‌ی مشارکت سیاسی را داشته باشید، در اینجا باید با توجه به مبانی نظری تحقیق خود، مشاغل فوق را به دو دسته‌ی دولتی و غیردولتی تقسیم و کدگذاری کنید.

۱. مشاغل دولتی (معلم، استاد دانشگاه، پزشک و...) = کد ۱

۲. مشاغل غیر دولتی (مکانیک، جوشکار، بنا و...) = کد ۲

همانطور که ملاحظه می‌شود این نوع کدگذاری صبغی قیاسی (کل به جزء) دارد و پاسخ‌ها بدون توجه به مصادیق، دسته‌بندی و کدگذاری می‌شود. شاید سؤال شود اگر چنین دسته‌بندی تئوریک وجود دارد چرا از همان ابتدا اقدام به طراحی یک سؤال بسته پاسخ با همان مقولات تئوریک نمی‌شود. در پاسخ باید گفت، طراحی سؤالات بازپاسخ این مزیت بزرگ را دارد که محقق می‌تواند به اشکال مختلفی که برخی از آنها ممکن است پس از گردآوری اطلاعات به ذهن برسد، پاسخ‌ها را دسته‌بندی کند. از سوی دیگر بسته پاسخ کردن چنین سؤالاتی بدین معناست که خود پاسخگو توانایی تشخیص، دسته‌بندی و تخصیص شغل خود را به یکی از مقولات تئوریک دارد که این کار به خودی خود از دقت نتایج تحقیق می‌کاهد. برای مثال فردی که مربی بهداشت یک آموزشگاه است، در تشخیص این موضوع که شغلش در طبقه‌ی آموزشی قرار می‌گیرد یا طبقه‌ی بهداشتی و درمانی دچار سردرگمی خواهد شد. لذا بهتر است که خود محقق عملیات دسته‌بندی و کدگذاری مشاغل را انجام دهد.

■ کدگذاری سؤالات بازپاسخ براساس رویکرد تجربی

در رویکرد تجربی باید ابتدا با مطالعه‌ی تعداد قابل‌توجهی^۱ از فرم‌های گردآوری اطلاعات، کلیه پاسخ‌های پاسخگویان را بر روی برگه‌ای یادداشت کنید، سپس با بررسی وجوه تشابه و تمایز آنها، اقدام به دسته‌بندی کنید. مثلاً ممکن است پاسخ‌های تعدادی از پاسخگویان به این سؤال که "نام آخرین کتاب غیردرسی که خوانده‌اید چه بود؟" شامل موارد زیر باشد:

دنایای سوفی، انسان در جستجوی معنی، جنس ضعیف، اولین و آخرین رهایی، بادبادک باز، در آغوش نور، وضعیت آخر، تحلیلی از مناسک حج، سیاحت غرب، سوپ جوجه برای روح، لطیفه‌های انگلیسی، تجسم خلاق، مجموعه اشعار تارهایی، قورباغه را بخور، روی ماه خداوند را ببوس، مطالعه با تمرکز، راه تحول، اینترنت را آسان یاد بگیرید، دردهای کمر و پشت، خواص میوه‌ها، طالع‌بینی هندی، خودآموز spss، جامعه‌شناسی توسعه و...

۱- پیشنهاد مشخص نگارنده بررسی حدود ده درصد پاسخ‌ها است که به صورت تصادفی انتخاب شده باشد، مشروط بر اینکه از تعداد ۳۰ مورد کمتر نباشد. البته حد بالایی برای تعداد مورد بررسی وجود ندارد. واقعیت این است که هر چه تعداد پرسشنامه‌های مورد بررسی در این مرحله بیشتر باشد، دقت کدنامه بیشتر خواهد شد.

۴۲ □ راهنمای کاربردی نرم افزار SPSS

حال با شناسایی وجوه تشابه و تمایز کتاب‌ها، باید یک دسته‌بندی ارائه دهید.
مثلاً: کتاب‌های رمان = کد ۱، کتابهای مذهبی = کد ۲، کتاب‌های روانشناسی = کد ۳،
کتاب‌های آموزشی = کد ۴ و...

همانطور که ملاحظه می‌شود این رویکرد دارای صبغه استقرایی (جزء به کل) بوده و فاقد طرح و نقشه قبلی برای دسته‌بندی پاسخ‌ها است.

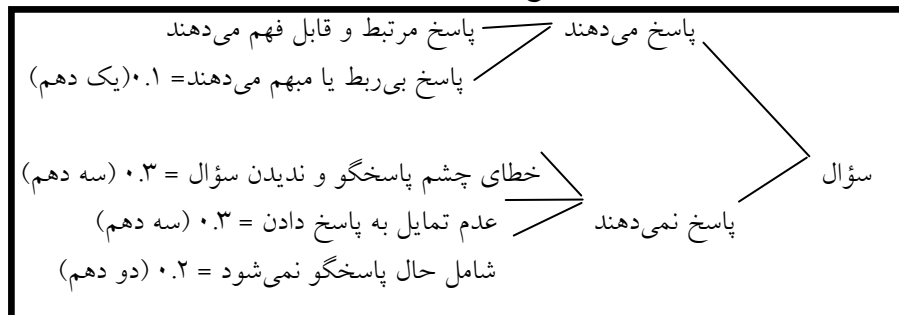
• کدگذاری مقادیر نامشخص (Missing Values)

در غالب موارد، اکثریت قابل توجهی از پاسخگویان به سؤالات محقق پاسخ می‌دهند، اما معمولاً در هر تحقیقی حالت‌های دیگری نیز وجود می‌آید. چنین مواردی اگر چه حکم موارد خاص را دارند، اما در هر صورت، برای ورود به نرم‌افزار SPSS نیاز به کدگذاری دارند که ذیلاً به آن می‌پردازیم.

■ کدگذاری سؤالات بی‌پاسخ یا پاسخ‌های مبهم

تجربه نشان داده است که همیشه فرایند تحقیق به صورت ایده‌آل پیش نمی‌رود و همه پاسخگویان به تمامی سؤالات پاسخ نمی‌دهند. بسیار پیش می‌آید که در تحقیقات مختلف با پاسخ‌های بی‌ربط یا سؤالات بی‌جواب متعددی مواجه می‌شویم. در اینجا می‌خواهیم نحوه کدگذاری این نوع سؤالات را مشخص کنیم. در مجموع می‌توان واکنش پاسخگویان را در مواجهه با یک سؤال به صورت زیر دسته‌بندی کرد:

شکل شماره (۱۳) انواع واکنش پاسخگویان به یک سؤال پرسشنامه



همانگونه که در این دسته‌بندی مشخص است، پاسخگویان یا به سؤالات ما پاسخ می‌دهند یا نمی‌دهند. هر یک از این دو واکنش را نیز می‌توان به دسته‌های جدیدی تقسیم کرد. آنها که پاسخ می‌دهند ممکن است پاسخشان به سؤال مربوط بی‌ربط و مبهم باشد. برای مثال کسی که در پاسخ به سؤال وضعیت تأهل می‌نویسد «عروس رفته گل بچینه!» اگر چه به سؤال مربوطه پاسخ داده است اما پاسخ وی ربطی به جواب‌های مورد انتظار محقق ندارد. لذا همانطور که

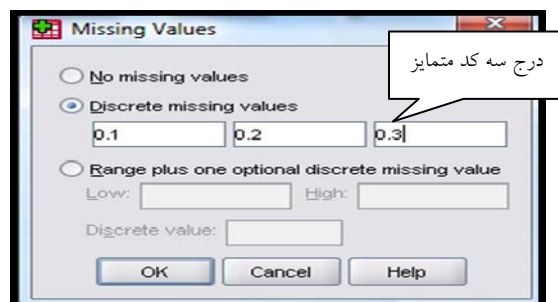
درس دوم – کدگذاری □ ۴۳

برای پاسخ‌های مورد انتظار، کدهای مشخصی در کدنامه منظور می‌کنیم، برای پاسخ‌های بی‌ربط هم باید کدهای مشخصی تعیین کنیم. مشروط بر اینکه کدهای مزبور، از کدهای مربوط به پاسخ‌های مورد انتظار متفاوت باشد. مثلاً نگارنده پیشنهاد می‌کند از آنجا که کدگذاری بسیاری از سؤالات کیفی از صفر یا یک شروع می‌شود و عموماً به صورت گسسته است، کد ۰.۱ (یک دهم) را برای پاسخ‌های بی‌ربط در نظر گرفت.

از سوی دیگر سؤالات بی‌پاسخ را نیز می‌توان به سه دسته تقسیم نمود. برخی از سؤالات بی‌پاسخ باقی می‌مانند چون پاسخگو آنها را ندیده است یا ممکن است نخواسته باشد به آنها جواب بدهد، برخی دیگر نیز به این دلیل بی‌پاسخ است که شامل حال پاسخگو نمی‌شوند. مثلاً پاسخگویان مجرد، سؤال مربوط به تعداد فرزندان را باید بی‌پاسخ بگذارند، چون شامل حال آنها نمی‌شود. نگارنده پیشنهاد می‌کند برای این وضعیت کد ۰.۲ (دو دهم) در نظر گرفته شود و برای وضعیتی که پاسخگو به دلایل مختلف (عمد یا فراموشی یا گنگ بودن سؤال و...) به سؤال جواب نمی‌دهد کد ۰.۳ در نظر گرفته شود. برای معرفی مقادیر نامشخص باید به کاربرگ متغیرها (variable view) مراجعه کرده و روی Missing Value کلیک کنید. در این پنجره به شیوه‌های مختلفی می‌توان کد پاسخ‌های بی‌ربط و مبهم را معرفی کرد. به مثال‌ها و شکل‌های بعد توجه کنید.

مثلاً اگر سؤالی درباره درآمد از افراد پرسیده شود، برای کسانی که به این سؤال پاسخ‌های بی‌ربط یا مبهمی نظیر «خیلی کم» یا «به قدر بخور نمیر» بدهند باید کد ۰.۱ و به کسانی نظیر خانه دارها، بیکاران و سایرین که این سؤال شامل حال آنها نمی‌شود کد ۰.۲ و به کسانی که شاغل هستند به دلایل مختلفی به این سؤال جواب نمی‌دهند، کد ۰.۳ داده می‌شود.

شکل شماره (۱۴) نحوه تعریف مقادیر نامشخص به صورت سه کد متمایز در ستون Missing Value

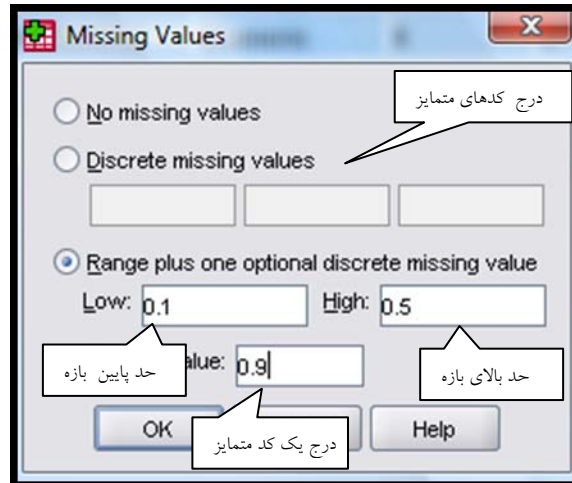


اگر تعداد کدهای مربوط به مقادیر نامشخص بیشتر از سه مورد باشد، می‌توان آنها را به صورت یک بازه تعریف کرد. در چنین مواردی SPSS حتی امکان اضافه نمودن یک کد نامشخص دیگر علاوه بر دامنه مورد نظر را در اختیار محقق قرار می‌دهد. برای مثال در سؤال مربوط میزان درآمد، اگر محقق بخواهد کد ۰.۱ را به پاسخ‌های بی‌ربط یا مبهمی نظیر «خیلی کم» یا «به قدر بخور نمیر»، و کد ۰.۲ را به بیکار، کد ۰.۳ را به خانه‌دار، کد ۰.۴ را به

۴۴ □ راهنمای کاربردی نرم افزار SPSS

بازنشسته و کد ۰.۵ را به از کار افتاده ها و سرانجام کد ۰.۹ را به بی پاسخها اختصاص بدهد در بخش Missing Value باید مقادیر نامشخص را به شکل زیر تعریف کنید.

شکل شماره (۱۵) نحوه تعریف مقادیر نامشخص در قالب یک بازه به علاوه یک کد مشخص در ستون Missing Value



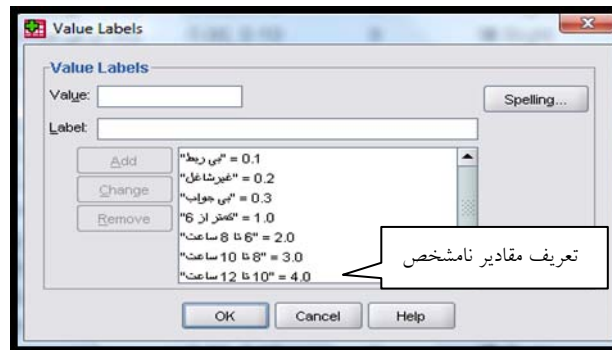
البته همواره باید به یاد داشته باشید که معرفی مقادیر نامشخص در ستون Missing Value اگر چه امری بسیار ضروری است، اما بهتر است برچسب هر یک از مقادیر نامشخص را همزمان با تعریف سایر مقادیر در ستون Value labels هم تعریف کنید. به ستونهای Missing و Values در تصویر پایین دقت کنید. همانطور که ملاحظه می شود این کار در مورد متغیر "مدت اشتغال" انجام شده است.

شکل شماره (۱۶) مقادیر نامشخص تعریف شده متغیر "مدت اشتغال"

	Name	Type	Wi...	De...	Label	Values	Missing
142	v10	Nu...	8	1	مدت اشتغال در ششانه روز	{0.1, بی ربط}...	0.1, 0.2, 0.3

درس دوم – کدگذاری □ ۴۵

شکل شماره (۱۷) تعریف مقادیر نامشخص متغیر "مدت اشتغال"



خروجی این متغیر در قالب جدول توزیع فراوانی به صورت زیر است:
جدول شماره (۱) یک نمونه از تعریف صحیح متغیر

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	کمتر از 6	27	2.2	7.0	7.0
	6 تا 8 ساعت	78	6.5	20.2	27.1
	8 تا 10 ساعت	159	13.2	41.1	68.2
	10 تا 12 ساعت	78	6.5	20.2	88.4
	سایر	45	3.7	11.6	100.0
	Total	387	32.0	100.0	
Missing	بی ربط	1	.1		
	غیرشاغل	785	64.9		
	بی جواب	36	3.0		
	Total	822	68.0		
	Total	1209	100.0		

حال اگر محقق فقط به تعریف برچسب مقادیر اکتفا کند و به تعریف مقادیر نامشخص در ستون Missing Value توجهی نکند، با خروجی زیر مواجه خواهد شد.

جدول شماره (۲) یک نمونه از تعریف اشتباه متغیر

مدت اشتغال در شبانه روز					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	بی ربط	1	.1	.1	.1
	غیرشاغل	785	64.9	64.9	65.0
	بی جواب	36	3.0	3.0	68.0
	کمتر از 6	27	2.2	2.2	70.2
	6 تا 8 ساعت	78	6.5	6.5	76.7
	8 تا 10 ساعت	159	13.2	13.2	89.8
	10 تا 12 ساعت	78	6.5	6.5	96.3
	سایر	45	3.7	3.7	100.0
	Total	1209	100.0	100.0	

تفاوت مهم جدول شماره (۲) با جدول شماره (۱) آن است که عملیات درصدگیری در جدول شماره ۲ بدون توجه به مقادیر نامشخص انجام شده است. حال آنکه در جدول شماره (۱) عملیات درصد گیری یکبار با احتساب مقادیر نامشخص (percent) و یکبار بدون احتساب آنها (valid percent) انجام شده است.^۱

علاوه بر آنچه گفته شد باید به برچسب مقادیر نامشخص نیز توجه داشت. در جدول شماره (۳) علیرغم اینکه مقادیر نامشخص در قسمت Missing Value تعریف شده و در نتیجه در عملیات درصدگیری دخالت داده شده است، اما خواننده این جداول به علت تعریف نشدن برچسب‌های مربوط به مقادیر نامشخص در ستون Value label، مفهوم مقادیر 0.1, 0.2, 0.3 را متوجه نخواهد شد.

۱- ستون‌های percent و valid percent را در جداول شماره ۲ و ۳ با یکدیگر مقایسه کنید.

جدول شماره (۳) یک نمونه اشتباه در تعریف متغیر

مدت اشتغال در شبانه روز					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	کمتر از 6	27	2.2	7.0	7.0
	6 تا 8 ساعت	78	6.5	20.2	27.1
	8 تا 10 ساعت	159	13.2	41.1	68.2
	10 تا 12 ساعت	78	6.5	20.2	88.4
	سایر	45	3.7	11.6	100.0
	Total	387	32.0	100.0	
Missing	0.1	1	.1		
	0.2	785	64.9		
	0.3	36	3.0		
	Total	822	68.0		
	Total	1209	100.0		

■ کدگذاری پاسخ‌های بی‌نظر یا نمی‌دانم

از آنجا که ممکن است برخی پاسخگویان در زمینه‌ی سؤالات ما نظر یا اطلاعی نداشته باشند، لازم است برای واکنش آنان نیز کد خاصی در نظر گرفته و کد مزبور به هنگام تعریف متغیرها در فایل داده‌ها به عنوان مقادیر نامشخص (Missing Value) تعریف شود. در غیر اینصورت کد مربوط به پاسخ بی‌نظر، بسته به مقدار عددی آن تفسیرهای گمراه‌کننده‌ای می‌تواند به دنبال داشته باشد. برای مثال چنانچه مقادیر متغیر موردنظر از نوع رتبه‌ای (خیلی مخالف=۱، مخالف=۲، بینابین=۳، موافق=۴، خیلی موافق=۵) باشد و محقق برای پاسخ بی‌نظر، کد صفر را در نظر بگیرد و این در حالی باشد که در بخش Missing Value آن را معرفی نکند، کد صفر در محاسبات آماری به عنوان یک مقدار واقعی در نظر گرفته می‌شود و معنایی نظیر مخالف تر از خیلی مخالف را در نظر خواهد داشت. بنابر این کدهای مربوط به پاسخ‌های نمی‌دانم و نظری ندارم را باید در کنار کدهای دیگر در بخش Missing Value تعریف کرد. مثلاً همانطور که گفته شد، می‌توان برای این دسته از پاسخ‌ها کد ۰.۳ را در نظر گرفت.

در پایان این مبحث ذکر چند نکته‌ی مهم ضروری است:

۱. برای سهولت کار در ورود اطلاعات، بهتر است کدها را زیر هم و در حاشیه‌ی سمت راست پرسشنامه به صورت ستونی بنویسید (این مورد را حتی به هنگام ساخت پرسشنامه می‌توانیم به یاد داشته باشیم و از قبل جایی را برای آن در نظر بگیریم).

۲. از هرگونه ادغام یا دسته‌بندی داده‌ها (به شیوه‌ی دستی) قبل از ورود به نرم‌افزار spss خودداری کنید. مثلاً اگر از پاسخگو سن او را در قالب سؤال باز پرسیده‌اید و او یک عدد را به عنوان سن، پاسخ داده است، بهتر است از ادغام سن در دسته‌های پنج یا ده سال خودداری نموده و عین پاسخ او را به عنوان کد، وارد کامپیوتر کنید. چون با این کار سطح سنجش متغیرها تقلیل می‌یابد و امکان استفاده محقق از بسیاری از آزمون‌های پیشرفته آماری سلب می‌شود.

۳. از هرگونه ترکیب پاسخ‌ها قبل از ورود داده‌ها به نرم‌افزار spss خودداری کنید. مثلاً اگر ۲۰ گویه برای سنجش رضایت شغلی طراحی شده است، بهتر است به جای آنکه مجموع نمرات پاسخگویان را در کل مقیاس محاسبه و وارد نرم افزار کنید، ابتدا تمامی گویه‌ها را کدگذاری و وارد نرم‌افزار کرده، سپس از طریق دستورهای مناسب، شاخص مجموع نمرات محاسبه کنید. مزیت مهم این کار، امکان محاسبه‌ی شاخص پایایی و شناسایی گویه‌های نامناسب در فرآیند تحقیق است^۱.

با توجه به توضیحاتی که از ابتدای این کتاب آمد، این موضوع روشن شده است که پس از جمع‌آوری اطلاعات، با تعدادی فرم‌های تکمیل شده گردآوری اطلاعات (اعم از پرسشنامه، مصاحبه‌نامه، فرم استخراج از اطلاعات و...) روبرو می‌شوید و اولین گام برای استخراج اطلاعات و تجزیه و تحلیل آن، کدگذاری همه‌ی پاسخ‌ها است و بعد از آن، تعریف کدها برای نرم‌افزار و در نهایت ورود پاسخ‌های پاسخگویان به نرم‌افزار spss است که توضیحات مربوط به آن به تفصیل در صفحات پیشین آمد. برای آشنایی بیشتر به پرسشنامه کدگذاری شده‌ی زیر توجه کنید. این پرسشنامه برای سنجش "میزان استفاده‌ی دانش‌آموزان از مواد خوراکی موجود در فروشگاه مدرسه" تهیه شده است.

۱- در این خصوص در درس "سنجش پایایی" توضیحاتی ارائه می‌شود.

درس دوم – کدگذاری □ ۴۹

نمونه‌ای از پاسخ‌های کدگذاری شده پرسشنامه "میزان استفاده دانش‌آموزان از مواد خوراکی موجود در فروشگاه مدرسه"

سلام دوست عزیز					
پرسشنامه‌ای که پیش روی شماست در مورد بوفه یا همان فروشگاه‌ای است که در مدرسه شما وجود دارد. سؤالاتی در مورد اینکه چه می‌خرید؟ چه چیزی بیشتر می‌خرید و از این قبیل سؤالات است. برای آگاهی از خوراکی‌هایی که دانش‌آموزان از بوفه مدرسه می‌خرند، شما به صورت تصادفی انتخاب شدید. این پرسشنامه علاوه بر شما به دانش‌آموزان دیگر هم داده می‌شود. نیازی نیست که اسم و فامیل خودتان را بنویسید، از اینکه وقت خود را برای پرکردن این پرسشنامه می‌گذارید و با دقت به سؤالات ما جواب می‌دهید.					
متشکریم					
۱. شما معمولاً خوراکی‌های را که در مدرسه می‌خورید، چطور تهیه می‌کنید؟					
□ 1 از خانه می‌آورم	□ 2 از فروشگاه (بوفه) داخل مدرسه می‌خرم	□ 3 هم از خانه می‌آورم و هم از فروشگاه مدرسه می‌خرم	□ 4 از فروشگاه‌های بیرون از مدرسه می‌خرم	□ 5 از راه دیگر (لطفاً بنویسید)	3
۲. روزانه در مدرسه چند بار خوراکی می‌خورید؟					
□ 0 اصلاً نمی‌خورم	□ 1 بار	□ 2 بار	□ 3 بار	□ 4 بار	□ 5 بار
□ 6 بار	□ 7 بار	□ 8 بار	□ 9 بار	□ 10 بار	بیشتر از ۱۰ بار
۳. روزانه حدوداً چند بار از بوفه مدرسه خوراکی می‌خرید؟					
□ 0 اصلاً نمی‌خورم	□ 1 بار	□ 2 بار	□ 3 بار	□ 4 بار	□ 5 بار
□ 6 بار	□ 7 بار	□ 8 بار	□ 9 بار	□ 10 بار	بیشتر از ۱۰ بار
۴. به خاطر دارید آخرین بار بابت خرید از بوفه مدرسه تان چقدر خرج کردید؟ (لطفاً به تومان بنویسید) فکر کنم ۲۰۰ تومان					
۵. نظرتان در مورد فروشگاه مواد خوراکی موجود در مدرسه تان چیست؟					
خیلی خوب 😊😊	خوب 😊	در حد متوسط 😊	بد 😞	خیلی بد 😞😞	
وضعیت تمیزی و بهداشتی بودن بوفه	□ 5	□ 4	□ 3	□ 2	□ 1
از نظر تنوع مواد خوراکی	□ 5	□ 4	□ 3	□ 2	□ 1
ساعت باز بودن بوفه	□ 5	□ 4	□ 3	□ 2	□ 1
قیمت مواد خوراکی	□ 5	□ 4	□ 3	□ 2	□ 1
کیفیت مواد خوراکی	□ 5	□ 4	□ 3	□ 2	□ 1
۶. لطفاً جدول زیر را با علامت زدن، کامل کنید					
در چه مقطع تحصیلی مشغول تحصیل هستید؟			□ 1 ابتدایی □ 2 راهنمایی □ 3 دبیرستان		
شهر محل سکونت شما کدام است؟			□ 1 بجنورد □ 2 اسفراین □ 3 شیروان		
نوع مدرسه			□ 1 دخترانه □ 2 پسرانه		
2 0 0 5 3 3 4 4 1 1 1					

تمرین درس دوم

۱- پاسخ های مندرج در پرسشنامه زیر را که مربوط به انجام تحقیقی در خصوص ورزش های همگانی است، کدگذاری کنید:

۱. آیا در پارک ها و بوستان های نزدیک محل سکونت یا ورزش شما، وسایل ورزشی از طرف شهرداری نصب شده است؟

☐ بلی ☐ خیر

۱.۱. اگر بلی، از چند مورد از وسایل مزبور استفاده می کنید؟

☐ هیچ ☐ یک ☐ دو ☐ سه ☐ چهار ☐ بیشتر

۲. شما معمولاً به تنهایی ورزش می کنید یا همراه با دوستان و آشنایان؟

☐ به تنهایی ☐ همراه با دوستان ☐ همراه با جمعی که لزوماً آشنا نیستند

☐ همراه با اعضای خانواده ☐ سایر

۳. انگیزه اصلی شما از ورزش کردن چیست؟

☐ حفظ تناسب اندام ☐ تأمین سلامت روان (روحیه شادی و نشاط)

☐ درمان بیماری ☐ گذران اوقات فراغت

۴. وضعیت اشتغال شما: ☐ خانه دار ☐ بیکار ☐ شاغل

۴.۱- در صورت اشتغال نوع شغل دقیقاً نوشته شود: سگ دو زدن از کله سحر

۵. جمعاً در طول یک شبانه روز چند ساعت مشغول کار هستید؟

☐ کمتر از ۶ ☐ ۶ تا ۸ ساعت ☐ ۸ تا ۱۰ ساعت ☐ ۱۰ تا ۱۲ ساعت ☐ بیشتر

۶. کارتان از ساعت چند شروع می شود و ساعت چند تمام می شود؟

ساعت تقریبی شروع: ۵ صبح ساعت تقریبی پایان: تا تاریک شدن هوا

۷. در شغل شما تا چه اندازه تحرک بدنی وجود دارد؟

☐ هیچ ☐ بسیار کم ☐ کم ☐ متوسط ☐ زیاد ☐ خیلی زیاد

۸. در منزلتان کدامیک از امکانات و وسایل ورزشی زیر را دارید؟

☐ تردمیل (دوی ثابت) ☐ دوچرخه ثابت ☐ اسکی فضایی

☐ آبکینگ پرو (دراز نشست) ☐ سایر (نام ببرید طناب)

۲- برنامه spss را اجرا کرده و متغیرهای فوق را در کاربرگ متغیرها به صورت کامل تعریف کنید.

۳- سوالات تمرین ۱ را در قالب یک مصاحبه با تعدادی از افراد که در پارک ها و فضاهای سبز به ورزش می پردازند مطرح و پاسخ های آنان را کدگذاری نموده و در فایلی به نام varzesh ضبط کنید.

درس سوم

ویرایش داده‌ها

در پایان این درس انتظار می‌رود دانشجو

- ۱- نحوه‌ی ویرایش داده‌ها و ضرورت آن را بداند.
- ۲- بتواند جدول فراوانی از متغیرها تهیه کند.
- ۳- با ساختار جداول فراوانی آشنا شود.
- ۴- بتواند خطاهای مربوط به ورود داده‌ها را شناسایی و برطرف کند.
- ۵- بتواند خطاهای مربوط به تعریف متغیرها را شناسایی و برطرف کند.

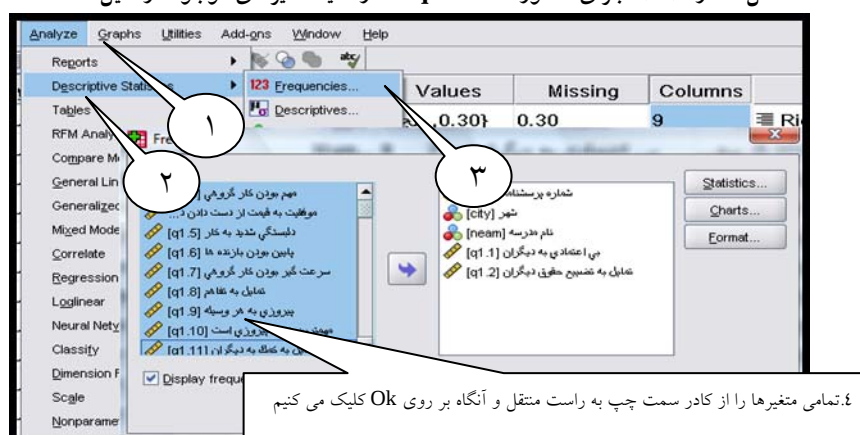
۵۲ □ راهنمای کاربردی نرم افزار SPSS

۱- مقدمه

اطلاعاتی که وارد سیستم شده‌اند، غالباً توأم با خطاهایی هستند که از اعتبار نتایج پژوهش می‌کاهند. لذا باید کوشید تا جایی که امکان دارد میزان خطاها را تقلیل داد. در این مرحله باید مطمئن شوید که کدهای وارد شده به نرم‌افزار SPSS، دقیقاً همان کدهایی باشد که در مرحله کدگذاری برای برنامه تعریف کرده‌اید.

به این منظور باید دستور زیر را برای تمامی متغیرهای موجود در فایل داده‌ها اجرا کرده و به نقد و بررسی جداول فراوانی مربوط به هر یک از متغیرها بپردازید.

شکل شماره (۱۸) اجرای دستور frequencies از کلیه متغیرهای موجود در فایل داده‌ها



۲- اصول ویرایش داده‌ها

جداول فراوانی متغیرها که حاصل اجرای دستور فوق است، باید از حیث موارد زیر بررسی شود.

• آیا مقادیر غیرمعتبر در جدول خروجی وجود دارد؟

مقادیر غیر معتبر به مقادیری اطلاق می‌شود که غیر از کدها یا دامنه‌ی اعدادی باشد که محقق برای متغیرها تعریف کرده است. برای مثال اگر محقق برای متغیر نوع مدرسه کد (۱) را برای پسرانه و کد (۲) را برای دخترانه تعریف کرده باشد، وجود کدهای دیگری غیر از این دو کد (نظیر ۳ یا ۲۲ و...)، یک مقدار غیر معتبر است.

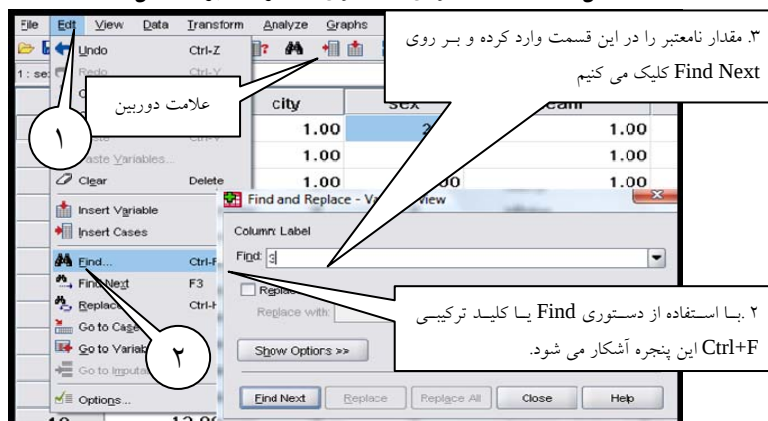
درس سوم - ویرایش داده‌ها □ ۵۳

جدول شماره (۴) توزیع فراوانی متغیر "نوع مدرسه" در حالت وجود مقادیر غیر معتبر در ستون داده‌های این متغیر

نوع مدرسه					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	دختر	120	50.0	50.0	50.0
	پسر	117	48.8	48.8	98.8
	3.00	3	1.3	1.3	100.0
	Total	240	100.0	100.0	

برای تصحیح این خطا در عنوان ستون مربوط به متغیر نوع مدرسه، باید به کاربرگ داده‌ها رفته، در ستون مربوط به این متغیر یک بار کلیک کنید تا ستون مزبور انتخاب شود^۱ و یا می‌توان از طریق دستور ذیل با کادری همانند زیر مواجه شد.
در کادر خالی که با این کار باز می‌شود، مقدار نامعتبر (۳) را وارد و کلید Ok را بزنید تا نرم‌افزار به طور خودکار مقدار خطای مورد درخواست را پیدا کند.

شکل شماره (۱۹) شیوه‌ی جستجوی مقادیر نامعتبر در فایل داده‌ها



پس از یافتن مقدار نامعتبر، شماره‌ی پرسشنامه‌ی متناظر با آن مورد را (که معمولاً به عنوان یکی از متغیرها در فایل معرفی می‌شود) شناسایی و پرسشنامه مورد نظر را پیدا کرده و مقدار واقعی کد را معلوم کنید و در کاربرگ داده‌ها کد صحیح را به جای کد نامعتبر وارد کنید.

۱- به جای کلیک کردن بر روی find از منوی Edit می‌توان روی کلید میانبر علامت دوربین کلیک کرد یا همزمان کلیدهای **ctrl + f** را فشار داد. شما می‌توانید با کمی دقت کلیدهای میانبر به تعدادی از دستورهای این برنامه را مقابل هر دستور در منوهای مختلف پیدا کنید.

بدیهی است که این عملیات را باید برای تمامی موارد خطای موجود در فایل داده ها و در مورد تمام متغیرها ادامه داد.

• آیا مقادیر نامشخص (missing value) تعریف شده است؟

ممکن است در یک یا چند مورد از متغیرها، مقادیر نامشخص را در بخش Missing Value معرفی نکرده باشید. به نتایج مندرج در جدول زیر دقت کنید:
جدول شماره (۵) نمونه دارای مقدار نامعتبر در تعریف مقادیر نامشخص

جدول فراوانی بی اعتمادی				
در این جدول بی جوابها، به عنوان پاسخ معتبر (valid) در نظر گرفته شده اند، که اشتباه است.				
	Frequency	Percent	Valid Percent	Cumulative Percent
Valid	7	.6	.6	.6
کاملاً واقعیت دارد	319	26.3	26.3	26.8
واقعیت دارد	357	29.4	29.4	56.2
کمی واقعیت دارد	411	33.8	33.8	90.0
واقعیت ندارد	75	6.2	6.2	96.2
اصلاً واقعیت ندارد	46	3.8	3.8	100.0
Total	1215	100.0	100.0	

نتایج جدول فوق نشان می دهد که نرم افزار SPSS مقادیری را که با برچسب بی جواب تعریف شده است به عنوان مقادیر نامشخص (Missing Value) نمی شناسد و آن را به عنوان پاسخ معتبر در نظر گرفته است. برای رفع این مشکل محقق باید به کاربرگ متغیرها رفته و در ستون Missing Value مربوط به آن متغیر، مقادیر نامشخص را تعریف کند و اصلاحات لازم را اعمال نماید.^۱

۱- توجه داشته باشید که درصد گیری ستون percent از کل پاسخگویان است که یک درصد گیری ناخالص است اما درصد گیری در ستون valid percent فقط از میان کلیه پاسخگویان است که پاسخ معتبر (missing) به سوال داده اند. همانگونه که ملاحظه می شود در این جدول (شماره ۵) به دلیل عدم تعریف مقادیر نامشخص در ستون Missing Value مقادیر مندرج در این دو ستون یعنی percent , valid (percent) یکسان است که نادرست می باشد.

درس سوم - ویرایش داده‌ها □ ۵۵

با تعریف مقادیر نامشخص در ستون Missing Value، با جدولی همانند جدول شماره ۶ مواجه می‌شوید.

جدول شماره (۶) نمونه بدون مقدار نامعتبر در تعریف مقادیر نامشخص

جدول فراوانی بی‌اعتمادی					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	کاملاً واقعیت دارد	319	26.3	26.4	26.4
	واقعیت دارد	357	29.4	29.6	56.0
	کمی واقعیت دارد	411	33.8	34.0	90.0
	واقعیت ندارد	75	6.2	6.2	96.2
	اصلاً واقعیت ندارد	46	3.8	3.8	100.0
	Total	1208	99.4	100.0	
Missing	بی‌جواب	7	.6		
	Total	1215	100.0		

در این جدول بی‌جوابها، به عنوان Missing Value تعریف شده است که صحیح است

- آیا برچسب‌ها (برچسب متغیرها، برچسب مقادیر معتبر و نیز برچسب مقادیر نامشخص) به صورت صحیح تعریف شده است؟

گاهی اتفاق می‌افتد که محقق تعریف بعضی برچسب‌ها را فراموش می‌کند یا در تعریف آنها مرتکب اشتباهات املایی یا انشایی می‌شود. با دقت در جداول فراوانی می‌توانیم این نوع از خطاها را نیز شناسایی کرده و به کاربرگ داده‌ها رفته، در ستون‌های مربوط، خطاهای مورد نظر را پیدا کرده و تصحیح کنید. به جدول زیر دقت کنید:

جدول شماره (۷) توزیع فراوانی متغیر "مقطع تحصیلی" در حالت نادرست

مقطع تحصیلی					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	راهنمایی	549	45.2	45.4	45.4
	2.00	660	54.3	54.6	100.0
	Total	1209	99.5	100.0	
Missing	.30	6	.5		
	Total	1215	100.0		

در جدول فوق سه مورد از خطاهای برچسب مشاهده می شود که عبارتند از:

۱- در املای برچسب متغیر اشتباه شده است و کلمه تحصیلی به صورت "تحصیلی" نوشته شده است.

۲- برچسب مقادیر به طور کامل تعریف نشده است و به جای آن یکی از برچسبها (۲) آمده است.

۳- اگرچه مقدار نامشخص (Missing Value) برای این متغیر در فایل داده ها تعریف شده است، اما برچسب مربوط به مقدار نامشخص (۰.۳) تعریف نشده است. محقق باید با مراجعه به پرسشنامه و کدنامه، برچسب های صحیح را شناسایی کرده و در کاربرگ متغیرها اصلاحات لازم را اعمال کند.

• آیا نحوه کدگذاری داده ها صحیح است؟

گاهی ممکن است مواردی پیش آید که محقق در نحوه کدگذاری متغیرها مرتکب اشتباه شود. برای مثال به جدول زیر دقت کنید:

جدول شماره (۸) نمونه اشتباه در نحوه کدگذاری

من به دیگران اعتماد دارم					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	۱ کاملاً واقعیت دارد	319	26.3	26.4	26.4
	۲ واقعیت دارد	357	29.4	29.6	56.0
	۳ کمی واقعیت دارد	411	33.8	34.0	90.0
	۴ واقعیت ندارد	75	6.2	6.2	96.2
	۵ اصلاً واقعیت ندارد	46	3.8	3.8	100.0
	Total	1208	99.4	100.0	
Missing	بی جواب	7	.6		
	Total	1215	100.0		

همانطور که ملاحظه می‌شود، پاسخ موافق به گویه "من به دیگران اعتماد دارم" که مضمون مثبتی دارد علی‌القاعده باید نمره بالاتری بگیرد، اما محقق در هنگام کدگذاری و به تبع آن ورود داده‌ها به نرم‌افزار SPSS، به پاسخ "کاملاً واقعیت دارد" کمترین نمره (۱) و به پاسخ "اصلاً واقعیت ندارد"، بالاترین نمره (۵) را داده است.

همانطور که می‌دانیم تصحیح یکایک پرسشنامه‌های وارد شده در کاربرگ داده‌ها به ویژه اگر تعداد آنها زیاد باشد، در عمل کار دشواری است. برای حل این مشکل می‌توان از دستور Recode از منوی Transform استفاده کرد. از این دستور می‌توان برای تغییر نحوه‌ی کدگذاری و نیز دسته‌بندی داده‌های مربوط به متغیرهای کمی استفاده کرد که مباحث آن به تفصیل در درس بعد توضیح داده می‌شود.

تمرین درس سوم

تمرین ۱- در تصویر زیر در متغیر سطر شماره ۲، الف- چند خطا وجود دارد؟ ب- نوع خطا یا خطاها را بنویسید.

	Name	Type	Width	Decimals	Label	Values	Missing
1	ROW	Numeric	8	2		None	None
2	In_rezayat	String	8	2	شخص کئی رضایت شغلی... (خیلی نارضایتی... 1.00)	None	None

تمرین ۲- خطاهای مندرج در جدول زیر را شناسایی کنید.
رشته تحصیلی

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid فاقد رشته	549	45.2	45.2	45.2
انسانی	223	18.4	18.4	63.5
2.00	242	19.9	19.9	83.5
ریاضی	122	10.0	10.0	93.5
فنی و حرفه‌ای	41	3.4	3.4	96.9
کار و دانش	38	3.1	3.1	100.0
Total	1215	100.0	100.0	

تمرین ۳- خطاهای مربوط به فایل داده‌های ورزش (تمرین ۳ از فصل ۲) را شناسایی و تصحیح کنید.

تمرین ۴- سعی کنید داده‌های مربوط به تحقیق یکی از همکلاسان یا اساتیدتان را تهیه و خطاهای آنها را شناسایی کنید.

درس چهارم

کدگذاری مجدد و دسته‌بندی

در پایان این درس انتظار می‌رود دانشجو

- ۱- اهمیت و ضرورت کدگذاری مجدد و دلایل استفاده از آن را بداند.
- ۲- کدگذاری مجدد به دو روش جایگزینی و ساختن متغیر جدید را بداند و مزیت روش دوم بر اول را تشخیص دهد.
- ۳- بتواند از دستور Recode برای اصلاح یا تغییر نحوه کدگذاری متغیرها استفاده کند.
- ۴- بتواند از دستور Recode برای دسته‌بندی متغیرهای کیفی استفاده کند.
- ۵- اصول و مراحل دسته‌بندی متغیرهای کمی به کیفی را بشناسد و با رعایت آنها بتواند متغیرهای کمی را با استفاده از دستور Recode دسته‌بندی کند.

۱- مقدمه

کدهایی که در مرحله کدگذاری به پاسخها نسبت داده می شود، طی مرحله تجزیه و تحلیل داده ها، ممکن است بارها نیاز به ادغام، تغییر یا دسته بندی داشته باشند. برای این کار باید از دستور Recode استفاده کرد.

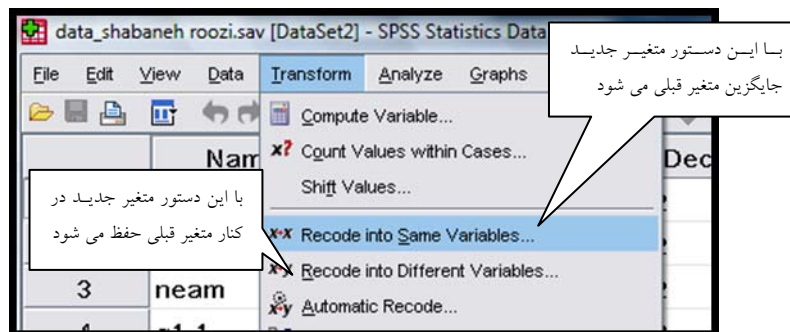
۲- حالت های مختلف دستور Recode

محقق در هنگام انتخاب دستور Recode با دو گزینه مواجه می شود:

- جایگزینی متغیر جدید به جای قبلی (Recode into same Variable)
- ساختن متغیر جدید ضمن حفظ متغیر قبلی (Recode into different Variable)

از آنجا که در حالت Recode into same Variable، متغیر اصلی که مبنای دسته بندی است حذف شده و متغیر جدیدی جایگزین آن می گردد، توصیه می شود در هنگام استفاده از دستور recode از گزینه Recode into different Variable استفاده شود. این کار موجب می شود متغیر اصلی به صورت خام و دسته بندی نشده باقی بماند و در کنار آن متغیر جدید نیز ایجاد شود. به تصویر زیر که دو گزینه مختلف دستور Recode را نشان می دهد، دقت کنید:

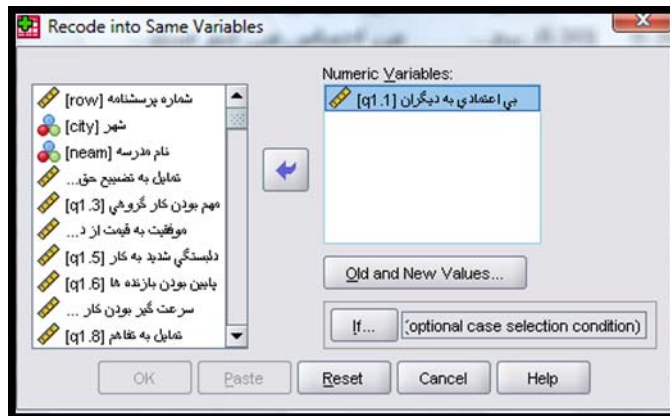
شکل شماره (۲۰) گزینه های موجود در دستور Recode



با اجرای حالت اول (Recode into same Variable) با پنجره ی زیر مواجه می شوید. کافی است متغیری را که باید دسته بندی شود به کادر سمت چپ منتقل کرده و آنگاه با کلیک بر روی Old and New Values، مقادیر قدیم و جدید را به نرم افزار معرفی کنید.

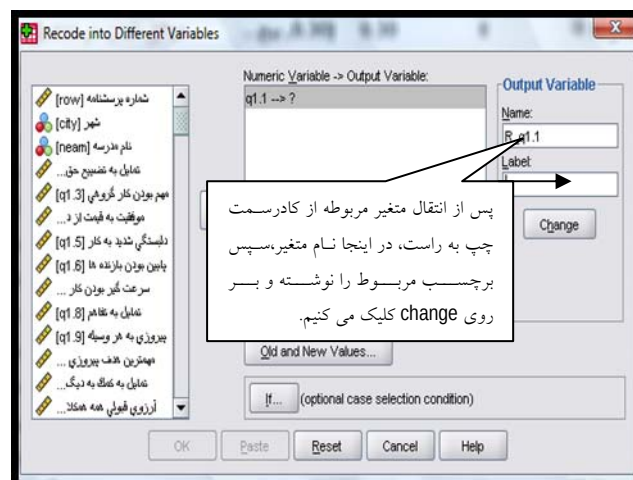
درس چهارم – کدگذاری مجدد و دسته‌بندی □ ۶۱

شکل شماره (۲۱) شیوه‌ی اجرای دستور **Recode** در حالت **Recode into same variable**



با اجرای حالت دوم (**Recode into different Variable**) پنجره‌ی زیر به نمایش در می‌آید. در این وضعیت ابتدا متغیر مربوطه را به کادر سمت راست منتقل و بعد از نوشتن نام و برچسب متغیر جدید، روی **change** کلیک کنید و بعد از آن با کلیک بر روی **Old and New Values**

شکل شماره (۲۲) شیوه‌ی اجرای دستور **Recode** در حالت **Recode into Different variabel**

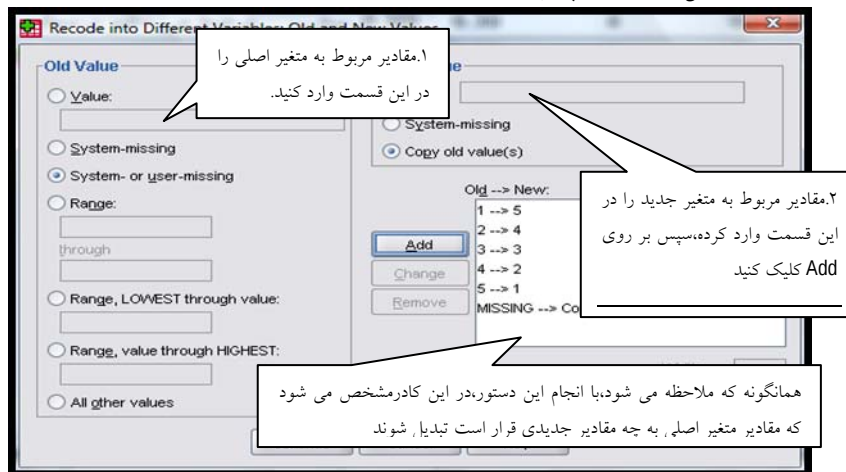


۳- کاربردهای دستور Recode

• استفاده از دستور Recode برای اصلاح یا تغییر نحوه کد گذاری

در پنجره Old and New Values می توان به اشکال مختلفی عمل کرد که بستگی به طرح و نقشه محقق برای دستور Recode دارد. به عبارت دیگر چنانچه قصد تغییر نحوه کدگذاری یک متغیر را داشته باشید، مثلاً بخواهید کدهای مربوط به گویه ی «بود و نبود من در این کلاس هیچ فرقی برایم نمی کند» را که قبلاً به این شیوه کدگذاری شده است: خیلی موافق (۱)، موافق (۲)، بینابین (۳)، مخالف (۴)، خیلی مخالف (۵)، به این شکل تغییر بدهید: خیلی موافق (۵)، موافق (۴)، بینابین (۳)، مخالف (۲)، خیلی مخالف (۱). با کلیک بر روی Old and New Values باید به شکل زیر عمل کنید.

شکل شماره (۲۳) پنجره Old and New Values در دستور Recode



با کلیک بر روی Continue و سپس Ok دستور Recode اجرا می شود. پس از اجرای این دستور می توانید با تهیه ی یک جدول فراوانی از متغیر جدید و مقایسه ی آن با جدول فراوانی متغیر قبلی از چگونگی انجام عملیات Recode اطمینان حاصل کنید.

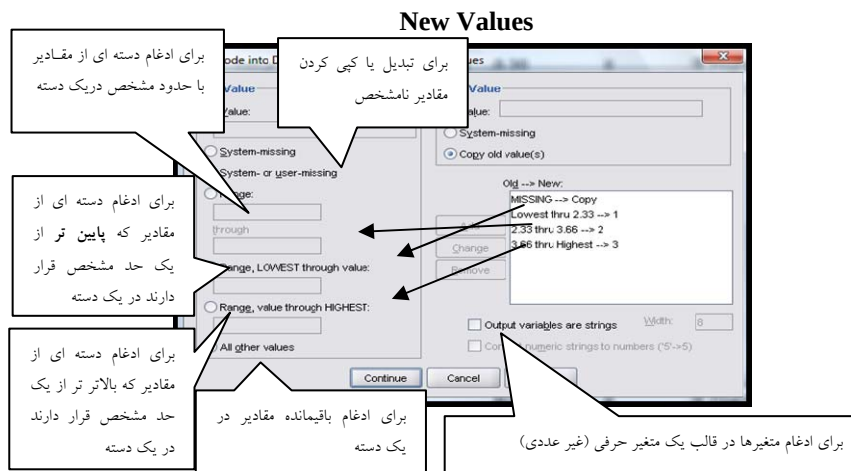
• استفاده از دستور Recode برای طبقه بندی متغیرهای کمی یا کیفی

در مواردی که بخواهید متغیرهای کیفی یا کمی را دسته بندی کنید، دستور Recode قابل استفاده می شود. برای مثال ممکن است بخواهید مشاغل را که کدهای مختلفی دارند، به دو دسته ی مشاغل دولتی و غیردولتی، یا به سه دسته ی مشاغل پایین، متوسط و بالا ادغام کنید.

درس چهارم – کدگذاری مجدد و دسته‌بندی □ ۶۳

علاوه بر این، ممکن است پس از ورود داده‌های مربوط به متغیرهای کمی قصد دسته‌بندی آنها را در دو یا چند طبقه داشته باشید. برای مثال بخواهید متغیر کمی "سن" را در سه دسته (جوان، میانسال، مسن) یا پنج دسته (خیلی بالا، بالا، متوسط، پایین، خیلی پایین)، طبقه‌بندی کنید، در این حالت باید با کلیک بر روی Old and New Values به شکل زیر عمل کنید:

شکل شماره (۲۴) چگونگی طبقه‌بندی متغیر کمی یا کیفی با اجرای دستور Recode در گزینه Old and



به این ترتیب با کلیک بر روی Continue و Ok، دستور Recode اجرا می شود. لازم است پس از اجرای این دستور، با تهیه یک جدول فراوانی از متغیر جدید و مقایسه آن با جدول فراوانی متغیر قبلی از چگونگی انجام عملیات Recode اطمینان پیدا کند. مراحل دسته‌بندی متغیرهای کمی به کیفی را می‌توان به شکل زیر بیان کرد:

۱- ابتدا یک جدول فراوانی از متغیر کمی تهیه و کمترین و بیشترین مقدار آن را شناسایی کنید.

- ۲- با تفاضل کمترین مقدار از بیشترین مقدار، دامنه‌ی تغییرات را محاسبه کنید.
- ۳- دامنه‌ی تغییرات را بر تعداد طبقات (توصیه می‌شود کمتر از ۳ یا بیشتر از ۱۰ نباشد)، تقسیم کنید تا فاصله‌ی طبقات به دست آید.
- ۴- فاصله‌ی طبقات را به کمترین مقدار اضافه کنید تا حدود بالا و پایین طبقه‌ی اول بدست آید.
- ۵- فاصله‌ی طبقات را به حد بالای طبقه‌ی اول اضافه نمایید تا حدود طبقات دوم بدست آید.
- ۶- این کار را تا محاسبه‌ی حدود آخرین طبقه ادامه دهید و حدود هر طبقه را بر روی برگه‌ای یادداشت کنید.

- ۷- دستور Recode را در حالت Recode into different Variable اجرا کنید.
- ۸- نام و برجسب متغیر جدید را وارد و بر روی عبارت change کلیک کنید.
- ۹- حدود طبقات را در کادرهای مربوط در پنجره Old and New Values وارد کنید.
- ۱۰- برای مقادیر نامشخص (System or User Missing Value) از دستور کپی مندرج در پنجره Old and New Values استفاده نمایید.
- ۱۱- برجسب مقادیر متغیر جدید را در کاربرد متغیرها بنویسید.
- ۱۲- از متغیر جدید جدول فراوانی تهیه و با جدول فراوانی قبلی مقایسه کنید تا از صحت دسته‌بندی مطمئن شوید.
- ۱۳- اگر فراوانی نسبی یا درصدی در برخی طبقات متغیر جدید کمتر از ۵ درصد باشد، بهتر است نسبت به ادغام این دسته‌ها در سایر دسته‌ها یا تغییر شیوهی دسته‌بندی از حالت آماری به حالت تئوریک تصمیم‌گیری کنید.

در ارتباط با مرحله‌ی شماره‌ی دوازده، یادآوری می‌شود که دسته‌بندی متغیرهای کمی در قالب متغیرهای کیفی می‌تواند براساس دو رویکرد صورت گیرد. در رویکرد آماری، کافی است دامنه‌ی تغییرات متغیر کمی ($R=MAX-MIN$)، محاسبه و بر تعداد دسته‌های موردنظر تقسیم شود تا فاصله‌ی هر دسته مشخص شود. این رویکرد که تصمیم‌گیری براساس آن امری رایج است گاهی موجب می‌شود همگنی درون دسته‌ها نادیده گرفته شود. اما رویکرد دوم صیغه‌ی تئوریک دارد. بدین معنی که در دسته‌بندی داده‌ها، شباهت یا نزدیکی اعضای درون هر دسته و نیز تمایز و فاصله‌ی بین دسته‌ها در درجه‌ی اول اهمیت قرار دارد و محقق تعهدی نسبت به رعایت فاصله‌ی برابر بین دسته‌ها را ندارد. در اینجا برای شفافیت بیشتر مطالب، همه‌ی مراحل فوق با یک مثال توضیح داده می‌شود

فرض کنید محقق قصد دارد "درصد باسوادان جهان را در سال ۱۹۹۵" در تحقیق خود گزارش دهد. از آنجا که این متغیر از نوع کمی است و فقط با تهیه‌ی جدول فراوانی از آن (که معمولاً دارای جداول فراوانی بلندی نیز هست)، نمی‌توان به خوبی آن را توصیف کرد، در این حالت علاوه بر استفاده از شاخص‌های آماری مختلف^۱، باید این متغیر را (درصد باسوادان جهان را در سال ۱۹۹۵) به یک متغیر کیفی تبدیل کرد تا برای خواننده قابل فهم‌تر باشد. برای این کار به ترتیب مراحل زیر باید انجام شود.

۱. ابتدا باید با اجرای دستور Frequencies، یک جدول فراوانی از این متغیر کمی تهیه کرد.

۱- در درس هفتم "توصیف نتایج" چگونگی توصیف یک متغیر کمی آمده است.

درس چهارم - کدگذاری مجدد و دسته‌بندی □ ۶۵

$$۱۰۰ - ۱۸ = ۸۲ \quad ۲. \text{ دامنه‌ی تغییرات را محاسبه کرد}$$

$$۸۲ \div ۳ = ۲۷.۳ \quad ۳. \text{ فاصله‌ی طبقات را بدست آورد}$$

۴. فاصله طبقات را به کمترین مقدار اضافه کرده تا حدود بالا و پایین طبقه اول بدست آید

$$۱۸ + ۲۷.۳ = ۴۵.۳$$

۵. فاصله طبقات را به حد بالای طبقه اول اضافه کرده تا حدود طبقه دوم بدست آید

$$۴۵.۳ + ۲۷.۳ = ۷۲.۶$$

۶. فاصله طبقات را به حد بالای طبقه دوم اضافه کرده تا حدود طبقه سوم بدست آید

$$۷۲.۶ + ۲۷.۳ = ۱۰۰$$

۷. با اجرای دستور Recode و گزینه Recode into different Variable یک متغیر

کیفی از "درصد باسوادان جهان در سال ۱۹۵۵" برای نرم‌افزار spss تعریف کرد.^۱

۸. در ستون value labels، مقوله‌های متغیر ساخته شده (متغیر کیفی) را برای نرم‌افزار spss تعریف کنید.

۹. با اجرای دستور Frequencies، یک جدول فراوانی از متغیر کیفی تهیه کرده و مقدار کل و همچنین مقادیر نامشخص را با توزیع فراوانی همان متغیر در حالت کمی مقایسه می‌کنیم. انجام این مقایسه باعث می‌شود تا در صورتیکه در دسته‌بندی دچار خطایی شده‌ایم، مراحل قبل را مجدداً بازبینی کنیم.

همانگونه که ملاحظه می‌شود، خروجی زیر توزیع فراوانی متغیر کیفی "درصد باسوادان جهان در سال ۱۹۹۵" است که در مقایسه با خروجی صفحه‌ی بعد (توزیع فراوانی این متغیر در حالت کمی)، بسیار گویاتر است.

جدول شماره (۹) توزیع فراوانی کیفی "درصد باسوادان جهان در سال ۱۹۹۵"

باسوادان					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	کم	14	12.8	13.1	13.1
	متوسط	22	20.2	20.6	33.6
	زیاد	71	65.1	66.4	100.0
	Total	107	98.2	100.0	
Missing	System	2	1.8		
Total		109	100.0		

۱- توضیحات مربوط به اجرای این دستور به تفصیل در درس چهارم "کدگذاری مجدد و دسته‌بندی داده‌ها" ارائه شده است.

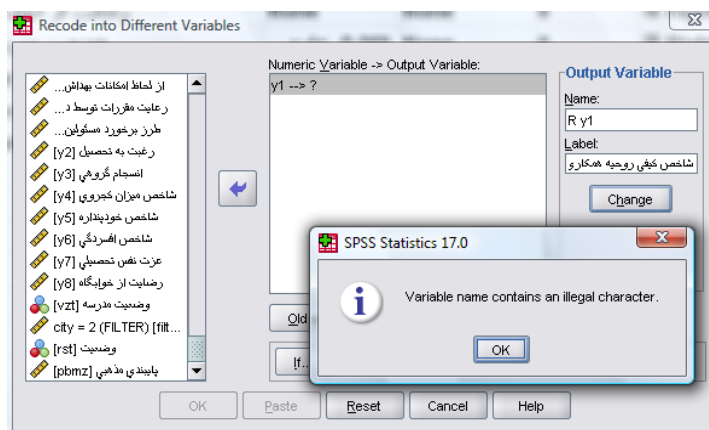
جدول شماره (۱۰) توزیع فراوانی کمی "درصد باسوادان جهان در سال ۱۹۹۵"

درصد باسوادان				
	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 18	1	.9	.9	.9
24	2	1.8	1.9	2.8
27	2	1.8	1.9	4.7
29	1	.9	.9	5.6
35	3	2.8	2.8	8.4
38	1	.9	.9	9.3
40	1	.9	.9	10.3
46	1	.9	.9	11.2
48	2	1.8	1.9	13.1
50	3	2.8	2.8	15.9
51	1	.9	.9	16.8
52	1	.9	.9	17.8
53	1	.9	.9	18.7
54	2	1.8	1.9	20.6
55	1	.9	.9	21.5
57	1	.9	.9	22.4
60	1	.9	.9	23.4
61	1	.9	.9	24.3
62	1	.9	.9	25.2
64	2	1.8	1.9	27.1
68	1	.9	.9	28.0
69	1	.9	.9	29.0
72	1	.9	.9	29.9
73	4	3.7	3.7	33.6
76	1	.9	.9	34.6
77	3	2.8	2.8	37.4
78	3	2.8	2.8	40.2
80	2	1.8	1.9	42.1
81	2	1.8	1.9	43.9
83	1	.9	.9	44.9
85	2	1.8	1.9	46.7
86	1	.9	.9	47.7
87	2	1.8	1.9	49.5
88	5	4.6	4.7	54.2
90	2	1.8	1.9	56.1
91	1	.9	.9	57.0
92	1	.9	.9	57.9
93	5	4.6	4.7	62.6
94	1	.9	.9	63.6
95	2	1.8	1.9	65.4
96	3	2.8	2.8	68.2
97	6	5.5	5.6	73.8
98	3	2.8	2.8	76.6
99	22	20.2	20.6	97.2
100	3	2.8	2.8	100.0
Total	107	98.2	100.0	
Missing System	2	1.8		
Total	109	100.0		

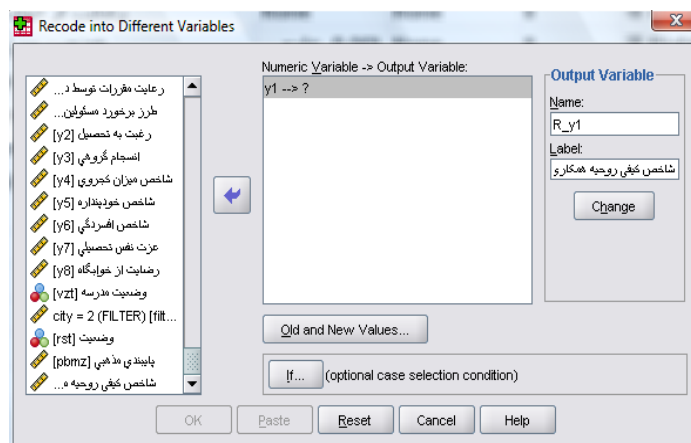
درس چهارم – کدگذاری مجدد و دسته‌بندی □ ۶۷

تمرین درس چهارم

تمرین ۱: پیام زیر در مورد چه خطایی صادر شده است؟



تمرین ۲: چرا در عملیات زیر علیرغم اینکه محقق نام و بر چسب متغیر و نیز دامنه‌ی دسته‌ها را در پنجره‌ی Old and New Values تعریف کرده است، کلید Ok فعال نیست؟



تمرین ۳: جدول زیر بیانگر چه اشکالاتی در نحوه‌ی دسته‌بندی داده‌های مربوط به متغیر روحیه‌ی همکاری است؟

شاخص کیفی روحیه همکاری

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	1.00	5	.4	.5	.5
	2.00	102	8.4	11.1	11.6
	3.00	814	67.0	88.4	100.0
	Total	921	75.8	100.0	
Missing	System	294	24.2		
Total		1215	100.0		

تمرین ۴: فایل world95 را باز کنید و کشورها را از حیث جمعیت (Population) در متغیر جدیدی به نام R_pop به پنج دسته از بسیار کم جمعیت تا بسیار پرجمعیت دسته‌بندی کنید (دسته‌بندی مورد اشاره را هم به صورت آماری و هم به صورت تئوریک انجام دهید و با یکدیگر مقایسه کنید).

تمرین ۵: در همان فایل کشورها را از لحاظ درصد باسوادان (Literacy) به پنج دسته‌ی بسیار زیاد تا بسیار کم تقسیم کنید.

تمرین ۶: در همان فایل کشورها را از حیث دین (Religion) با استفاده از دستور Recode به دو دسته‌ی مسلمان (۱) و غیر مسلمان (۲) تقسیم کنید.

تمرین ۷: در همان فایل درصد باسوادان را در سه دسته کم، متوسط و زیاد فقط در کشورهای مسلمان دسته‌بندی کنید.

درس پنجم

سنجش پایایی

در پایان این درس انتظار می‌رود دانشجو

- ۱- با مفاهیم معرف و مقیاس آشنا شود.
- ۲- مفهوم پایایی را بداند و با سنجش ثبات و همسازی درونی مقیاس آشنا شود.
- ۳- بتواند برای سنجش همسازی درونی از آلفای کرانباخ استفاده کند.
- ۴- تفسیر خروجی آلفای کرانباخ را بداند و معرف‌هایی مناسب را از نامناسب تمیز دهد.
- ۵- عواملی را که موجب کاهش مقدار آلفای کرانباخ می‌شود بشناسد و بتواند آنها را برطرف کند.

۱- مقدمه

پایایی^۱ یا دقت یکی از مباحث کلیدی در اندازه‌گیری مفاهیم است. مفاهیمی را که در پژوهش‌های علمی مورد سنجش قرار می‌گیرند از حیث سادگی یا پیچیدگی ابزار سنجش می‌توان به دو دسته تقسیم کرد.

الف- مفاهیمی که با یک سؤالات یا معرف سنجیده می‌شوند. به عبارت دقیق‌تر مفاهیمی که تنها یک معرف دارند. مثل سن، تحصیلات، جنسیت. یکی از روش‌های ساده‌ی اندازه‌گیری پایایی برای سنجش این دسته از مفاهیم "آزمون مجدد" است. به عبارت دیگر با تکرار اندازه‌گیری می‌توان به دقت ابزار طراحی شده برای سنجش این مفاهیم پی برد.

ب- مفاهیمی که با دو سؤالات (معرف) یا بیشتر سنجیده می‌شوند نظیر پایگاه اجتماعی-اقتصادی، افسردگی، رضایت شغلی. به این مجموعه سؤالات مربوط به یک مفهوم واحد، اصطلاحاً مقیاس گفته می‌شود. مثلاً مقیاس رضایت شغلی به مجموعه‌ای از سؤالات اطلاق می‌شود که ناظر به سنجش مفهوم "رضایت شغلی" است. یکی از روش‌های رایج برای سنجش پایایی این بخش از مفاهیم، آلفای کرانباخ است. آلفای کرانباخ معیاری را در اختیار محقق می‌گذارد که براساس آن می‌توان همسازی درونی مقیاس را شناسایی کرد.

۲- سنجش پایایی از حیث همسازی درونی به روش آلفای کرانباخ

همانطور که قبلاً اشاره شد با استفاده از آزمون آلفای کرانباخ می‌توان با جزئیات بیشتری در مورد دقت یک مقیاس و نیز یکایک گویه‌هایی که قرار است یک مفهوم را بسنجند، قضاوت کرد. به مثال زیر دقت کنید:

محقق قصد دارد "روحیه‌ی همکاری دانش‌آموزان" را در یک مدرسه اندازه‌گیری کند. برای این کار مقیاسی مرکب از ۱۵ گویه به شرح جدول بعد طراحی کرده است. او پس از سنجش‌روایی^۲ مقیاس از طریق توافق داوران، مقیاس فوق را در یک آزمون مقدماتی در اختیار تعدادی از دانش‌آموزان (حداقل ۳۰ الی ۵۰ نفر) می‌گذارد و از آنها می‌خواهد نظر خود را در محل مناسب علامت بزنند.

۱- برای آشنایی با مفهوم پایایی (reliability) روایی و روش‌های سنجش آنها به کتاب‌های روش‌های سنجش روایی و پایایی در تحقیقات فرهنگی و اجتماعی (حیدری چروده، ۱۳۸۹) و کتاب سنجش‌روایی و پایایی اثر کارمینز و زلر، ترجمه (فروغ‌زاده و اسکافی، ۱۳۸۶) مراجعه کنید.

درس پنجم – سنجش پایایی □ ۷۱

مقیاس سنجش "روحیه همکاری دانش‌آموزان"

عبارات جدول زیر نظرات برخی از دانش‌آموزان را در زمینه رقابت و همکاری نشان می‌دهد. ما در اینجا می‌خواهیم بدانیم شما چه نظری دارید. لطفاً آنها را به دقت بخوانید و بعد قضاوت خودتان را با گذاشتن علامت ضربدر در محل

مناسب مشخص نمایید					
ردیف	قضاوتها	کاملاً واقعیت دارد	واقعیت دارد	کمی واقعیت دارد	واقعیت ندارد اصلاً
۱	من به خیلی‌ها اعتماد نمی‌کنم				
۲	برای موفق شدن، می‌توان حق دیگران را زیر پا گذاشت				
۳	کار گروهی مهم‌تر از برنده شدن است				
۴	می‌خواهم موفق شوم، حتی اگر به قیمت از دست دادن دیگری باشد.				
۵	من هر کاری را که بر عهده‌ام باشد طوری انجام می‌دهم که انگار زندگی‌ام به آن وابسته است				
۶	بازنده‌ها همیشه پایین‌تر از دیگران هستند				
۷	کار گروهی از سرعت من در کارها کم می‌کند				
۸	باید با همه به تفاهم برسم				
۹	من برای برنده شدن و پیروز شدن به هر کاری دست می‌زنم				
۱۰	مهم‌ترین کار در بازی برنده شدن است				
۱۱	دوست دارم به دیگران کمک کنم				
۱۲	دوست دارم همه بچه‌های کلاس در امتحان قبول بشوند				
۱۳	دوست ندارم برای رسیدن به هدف دیگران را تحت فشار قرار دهم				
۱۴	ضرر دیگران، نفع من است				
۱۵	هر اندازه بیشتر برنده می‌شوم، احساس قدرت بیشتری می‌کنم.				

حال محقق پس از گردآوری همه ی اطلاعات، کدگذاری پاسخ‌ها و همچنین تعریف فایل متغیرها و ورود داده‌ها به نرم‌افزار SPSS، با چنین فایلی رو برو می‌شود.

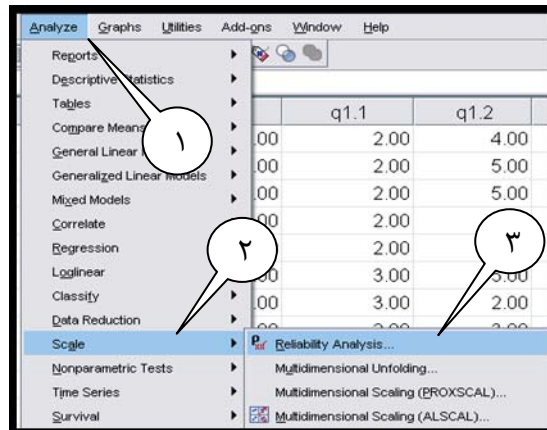
شکل شماره (۲۵) فایل داده های مقیاس "روحیه همکاری دانش آموزان"

Name	Type	Wt...	De...	Label	Values
4	q1.1	Num...	8	۲ بی اعتمادی به دیگران	... , 0.30} بی
5	q1.2	Num...	8	تمایل به تضییع ...	... , 0.30} تمایل
6	q1.3	Num...	8	مهم بودن کار گروهی	... , 0.30} مهم
7	q1.4	Num...	8	موفقیت به قیمت ...	... , 0.30} موفقیت
8	q1.5	Num...	8	دلچسپی شدید به کار	... , 0.30} دلچسپی
9	q1.6	Num...	8	پایین بودن بازنده ها	... , 0.30} پایین
10	q1.7	Num...	8	سرعت گیر بودن کا...	... , 0.30} سرعت
11	q1.8	Num...	8	تمایل به تفاهم	... , 0.30} تمایل
12	q1.9	Num...	8	پیروزی به هر وسیله	... , 0.30} پیروزی
13	q1.10	Num...	8	مهمترین هدف پیرو...	... , 0.30} مهم
14	q1.11	Num...	8	تمایل به کمک به دیگران	... , 0.30} تمایل

حال برای سنجش پایایی مقیاس "روحیه همکاری دانش‌آموزان" به شیوه ی آلفای کرانباخ به ترتیب مراحل زیر باید اجرا شود.

۷۲ □ راهنمای کاربردی نرم افزار SPSS

شکل شماره (۲۶) شیوهی اجرای دستور **Reliability** برای سنجش پایایی مقیاس. مرحله ۱



سپس عملیات را طبق شکل زیر ادامه می دهد:

شکل شماره (۲۷) شیوه اجرای دستور **Reliability** برای سنجش پایایی مقیاس. مرحله ۲



با اجرای این دستور، در فایل خروجی spss شاهد پنج جدول به شرح ذیل خواهید بود^۱.

۱- یاد آوری می شود که دامنه ی شاخص آلفا در شرایط عادی بین صفر تا یک است و مقدار آن هر چه قدر به یک نزدیکتر باشد، بهتر است و نشانه ی پایایی بالای مقیاس است.

درس پنجم - سنجش پایایی □ ۷۳

جدول شماره (۱۱): خلاصه پردازش اطلاعات

Case Processing Summary		
	N	%
Valid	1109	91.3
Excluded(a)	106	8.7
Total	1215	100.0

a Listwise deletion based on all variables in the procedure.

این جدول وضعیت نمونه‌ی مورد بررسی را از حیث پاسخ به مجموعه گویه‌های مقیاس "روحیه همکاری دانش آموزان" نشان می‌دهد. همانگونه که ملاحظه می‌شود از مجموع

۱۲۱۵

عضو نمونه مورد بررسی که انتظار می‌رفت به همه‌ی گویه‌های طرح شده برای این مقیاس پاسخ دهند، ۹۱/۳ درصد (۱۱۰۹ نفر) به تمام گویه‌ها پاسخ داده‌اند و ۸/۷ درصد (۱۰۶ نفر) حداقل به یک گویه پاسخ نداده‌اند و به همین دلیل از فرآیند تحلیل پایایی خارج شده‌اند

جدول شماره (۱۲): میزان آلفای کرانباخ در مقیاس "روحیه همکاری دانش آموزان"

Reliability Statistics	
Cronbach's Alpha	N of Items
0.618	15

این جدول میزان آلفای کرانباخ بدست آمده از مقیاس "روحیه همکاری دانش آموزان" را نشان می‌دهد. همانگونه که مشاهده می‌شود در صورتیکه تمام گویه‌ها در مقیاس حضور داشته باشند، میزان پایایی مقیاس "روحیه همکاری دانش آموزان" براساس آلفای کرانباخ ۰/۶۱۸ است که بیانگر پایایی نسبتاً بالای این مقیاس است.

۷۴ □ راهنمای کاربردی نرم افزار SPSS

جدول شماره (۱۳): توصیف میانگین و انحراف معیار گویه های مقیاس "روحیه همکاری دانش آموزان"

Item Statistics			
گویه ها	Mean	Std. Deviation	N
۱. من به خیلی ها اعتماد نمی کنم.	2.3003	1.04372	1109
۲. برای موفق شدن، می توان حق دیگران را زیر پا گذاشت.	4.3616	.99457	1109
۳. کارگروهی مهمتر از برنده شدن است.	3.7773	1.20692	1109
۴. می خواهم موفق شوم، حتی اگر به قیمت از دست دادن دیگری باشد.	4.0884	1.23624	1109
۵. من هر کاری را که بر عهده ام باشد طوری انجام می دهم که انگار زندگی ام به آن وابسته است.	2.1118	1.01972	1109
۶. بازنده ها همیشه پایین تر از دیگران هستند	3.6853	1.30646	1109
۷. کار گروهی از سرعت من در کارها کم می کند.	3.5600	1.21708	1109
۸. باید با همه به تفاهم برسم.	3.7944	1.13650	1109
۹. من برای برنده شدن و پیروز شدن به هر کاری دست می زنم.	3.4617	1.39028	1109
۱۰. مهمترین کار در بازی برنده شدن است.	3.0388	1.33183	1109
۱۱. دوست دارم به دیگران کمک کنم.	4.6213	.67239	1109
۱۲. دوست دارم همه بچه های کلاس در امتحان قبول بشوند.	4.6312	.75055	1109
۱۳. دوست ندارم برای رسیدن به هدف دیگران را تحت فشار قرار دهم.	3.7764	1.32952	1109
۱۴. ضرر دیگران ، نفع من است.	4.4581	.94109	1109
۱۵. هر اندازه بیشتر برنده می شوم، احساس قدرت بیشتری می کنم.	2.4148	1.34460	1109

این جدول وضعیت میانگین و انحراف معیار هر گویه را به تفکیک نشان می دهد. همانطور که ملاحظه می شود میانگین درمورد بیشتر گویه ها حول عدد ۳ و ۴ است و تنها میانگین سه گویه ۱، ۵ و ۱۵ از سایر میانگین ها متفاوت است. به عبارتی دیگر به غیر از این سه گویه، سایر گویه های مقیاس با یکدیگر هم وزن هستند. شاخص های پراکندگی (انحراف معیار یا واریانس) بیانگر قدرت تمیز گویه ها هستند. هرچه قدر فاصله ی نمرات پاسخگویان در یک گویه از میانگین آن گویه بیشتر باشد، بیانگر توانایی آن گویه در تمایز بخشی بین پاسخگویان است. البته گفتنی است که بالا بودن انحراف معیار گویه با تنهایی کافی نیست. بلکه از یکسو گویه ها باید با همدیگر هماهنگ باشند (همسازی درونی) و از سوی دیگر پراکندگی یا همان قدرت تمیز شاخص کلی که از ترکیب گویه های مقیاس به دست می آید را افزایش بدهند. به نتایج مندرج در جداول زیر دقت کنید.

جدول شماره (۱۴): مقادیر آماره‌های مقیاس "روحیه همکاری دانش‌آموزان"

Item-Total Statistics				
ستون (۵)	ستون (۴)	ستون (۳)	ستون (۲)	ستون (۱)
Cronbach's Alpha if Item Deleted	Corrected Item-Total Correlation	Scale Variance if Item Deleted	Scale Mean if Item Deleted	گویه‌ها
.629	.052	44.904	51.7809	۱. من به خیلی‌ها اعتماد نمی‌کنم.
.576	.422	40.402	49.7196	۲. برای موفق شدن، می‌توان حق دیگران را زیر پا گذاشت.
.597	.272	41.053	50.3039	۳. کارگروهی مهمتر از برنده شدن است.
.577	.378	39.325	49.9928	۴. می‌خواهم موفق شوم، حتی اگر به قیمت از دست دادن دیگری باشد.
.647	-.086	46.885	51.9693	۵. من هر کاری را که بر عهده‌ام باشد طوری انجام می‌دهم که انگار زندگی‌ام به آن وابسته است.
.598	.267	40.577	50.3959	۶. بازنده‌ها همیشه پایین تر از دیگران هستند
.589	.316	40.347	50.5212	۷. کار گروهی از سرعت من در کارها کم می‌کند.
.626	.089	44.084	50.2867	۸. باید با همه به تفاهم برسم.
.589	.311	39.368	50.6195	۹. من برای برنده شدن و پیروز شدن به هر کاری دست می‌زنم.
.588	.318	39.622	51.0424	۱۰. مهمترین کار در بازی برنده شدن است.
.597	.332	43.330	49.4599	۱۱. دوست دارم به دیگران کمک کنم.
.600	.285	43.342	49.4500	۱۲. دوست دارم همه بچه‌های کلاس در امتحان قبول بشوند.
.617	.161	42.171	50.3048	۱۳. دوست ندارم برای رسیدن به هدف دیگران را تحت فشار قرار دهم.
.588	.352	41.564	49.6231	۱۴. ضرر دیگران، نفع من است.
.595	.279	40.165	51.6664	۱۵. هر اندازه بیشتر برنده می‌شوم، احساس قدرت بیشتری می‌کنم.

میانگین مقیاس در صورت حذف گویه

واریانس مقیاس در صورت حذف گویه

همبستگی هر گویه با مقیاس

آلفای کل مقیاس در صورت حذف هر گویه

جدول شماره (۱۵): مقادیر آماره‌های مقیاس "روحیه همکاری دانش‌آموزان"

Scale Statistics			
Mean	Variance	Std. Deviation	N of Items
54.0812	46.721	6.83526	15

میانگین کل مقیاس

واریانس کل مقیاس

انحراف معیار کل مقیاس

تعداد کل گویه‌های مقیاس

جدول شماره ۱۵ میانگین، واریانس و انحراف معیار نمره‌ی کل مقیاس را نشان می‌دهد. نمره‌ی کل مقیاس از طریق مجموع نمرات هر پاسخگو به همهی گویه‌های مقیاس محاسبه می‌شود. به عنوان مثال چون این مقیاس از ۱۵ گویه تشکیل شده است و نمره‌ی هر گویه در دامنه‌ای بین ۱ تا ۵ قرار دارد، می‌توان نتیجه گرفت نمره‌ی هر پاسخگو بین ۱۵ تا ۷۵ قرار خواهد گرفت. از این‌رو طبیعی است که میانگین مقیاس نیز در همین دامنه قرار گیرد.^۱

حال اگر نتایج جدول ۱۴ و ۱۵ را با یکدیگر مقایسه کنیم درمی‌یابیم که با حذف گویه‌های ۱، ۵ و ۱۵ در مقایسه با سایر گویه‌ها، میانگین مقیاس به میزان کمتری تحت تأثیر قرار می‌گیرد (ستون ۲ از جدول ۱۴).

تأمل در ستون واریانس مقیاس در صورت حذف گویه (ستون ۳ از جدول ۱۴)، نشان می‌دهد تمامی گویه‌ها در صورت حذف، واریانس را از حالت فعلی (۴۶/۷۲۱) کاهش می‌دهند، با حذف گویه ۵، واریانس مقیاس افزایش می‌یابد که خود نشانه دقیق نبودن این گویه است.

نگاهی به نتایج ستون همبستگی هر یک از گویه‌ها با کل مقیاس بیانگر این واقعیت است که به استثنای گویه‌های ۱، ۵، ۱۵، نمره‌ی پاسخگویان در هر یک از گویه‌ها با نمره کل مقیاس

۱- به این دلیل که پاسخگو نظر خود را بر روی یک طیف پنج قسمتی از کاملاً واقعیت دارد تا اصلاً واقعیت ندارد علامت می‌زند. اگر پاسخگویی از هر گویه کمترین نمره را بگیرد، حداقل نمره او ۱۵ خواهد شد و اگر پاسخگویی از هر گویه‌ای بیشترین نمره را بگیرد، حداکثر نمره وی ۷۵ خواهد شد. چون ۱۵ گویه داریم و به عبارتی دیگر: $۱۵ * ۱ = ۱۵$ (حداقل نمره) و $۱۵ * ۵ = ۷۵$ (حداکثر نمره).

همبستگی قابل قبولی دارد که نشانگر همسازی درونی اکثر گویه‌های تشکیل‌دهنده‌ی مقیاس است. اما همبستگی گویه‌ی ۵ با سایر نمره کل مقیاس منفی است که نشان‌دهنده‌ی ناهمسازی شدید این گویه با سایر گویه‌های مقیاس است.

شاید سؤالات شود ملاک اصلی برای تشخیص تناسب یا عدم تناسب یک گویه چیست؟ پاسخ این سؤالات را باید در ستون ۵ جدول ۴ جستجو کرد. این ستون میزان آلفای کل مقیاس را در صورت حذف هر گویه نشان می‌دهد. مقدار آلفای کل مقیاس "روحیه همکاری دانش‌آموزان" طبق جدول خروجی شماره (۱۲) معادل با ۰.۶۱۸ بود. همانطور که در ستون ۵ جدول (۱۴) مشاهده می‌شود با حذف گویه‌های ۵، ۱ و ۸، مقدار آلفای کل مقیاس افزایش می‌یابد. در حالیکه حذف سایر گویه‌ها، مقدار آلفای کل مقیاس را کاهش می‌دهد. پیشنهاد می‌شود از میان گویه‌هایی که حذف‌شان آلفای کل مقیاس را افزایش می‌دهد، تنها نسبت به حذف گویه‌ای اقدام کنیم که:

الف) حذف آن، مقدار آلفای کل را بیشتر از حذف گویه‌های دیگر افزایش می‌دهد (در اینجا گویه ۵).

ب) میزان افزایش مزبور لااقل ۳ درصد باشد (در اینجا گویه ۵).

در صورتیکه طبق استنباط محقق، تعداد گویه‌هایی که قرار است حذف شود زیاد باشد، روایی محتوای مقیاس کاهش می‌یابد. طبیعی است که در چنین در مواقعی، محقق موظف است مقیاس جدیدی طراحی کند و مجدداً به سنجش روایی و پایایی آن پردازد.

نکته: گاهی اوقات پیش می‌آید که در سنجش پایایی از طریق آلفای کرانباخ، مقدار آلفای کل منفی می‌شود. به نظر می‌رسد برای حل مشکل بهتر است موارد زیر را واریسی کنید:

۱. ممکن است مقادیر پرت و دورافتاده در متغیرها وجود داشته باشد که گاهی به خاطر اشتباه در ورود داده‌ها است. مثلاً ممکن است کدگذار به جای کد ۲، مقدار ۲۲ را وارد کرده باشد. لذا برای شناسایی چنین مواردی همانگونه که در درس پیش گفته شد، ابتدا باید از تمامی متغیرها جدول فراوانی تهیه و موارد خطا را پیدا و تصحیح کرد.

۲. ممکن است مقادیر نامشخص (Missing value) تعریف نشده باشد. گاهی پیش می‌آید دانشجویان تعریف مقادیر نامشخص در ستون برچسب‌ها (Value) را حمل بر تعریف Missing value می‌کنند که اشتباه است. برای این مهم باید ستون Missing value مربوط به تمامی گویه‌ها را واریسی و موارد خطا را تصحیح کرد.

۳. ممکن است در نحوه‌ی کدگذاری گویه‌های مثبت یا منفی (مخصوصاً در مقیاس لیکرت و مقیاس‌های مشابه) اشتباهی صورت گرفته باشد.^۱ که باید از طریق واریس پرسشنامه‌های کدگذاری شده آن را اصلاح کرد.

۴. اگر با وجود بررسی تمامی موارد فوق، مقدار آلفا همچنان منفی یا ضعیف باشد، نتایج بدست آمده بیانگر همسازی درونی بسیار ضعیف مقیاس و در نتیجه نامناسب بودن میزان پایایی آن است. در چنین شرایطی لازم است به طراحی مجدد مقیاس و سنجش دوباره‌ی روایی و پایایی آن پرداخت.

۱- نحوه کدگذاری گویه‌هایی که دارای مضامین مثبت یا منفی هستند، باید طوری باشد که منطقاً همبستگی مثبت باهم داشته باشند.

درس پنجم - سنجش پایایی □ ۷۹

تمرین درس پنجم

تمرین ۱- نتایج زیر مربوط به سنجش پایایی مقیاس "نگرش نسبت به پوشش مدارس دخترانه" است. با توجه به نتایج مندرج در جداول به سؤالات بعد پاسخ دهید.

Case Processing Summary

		N	%
Cases	Valid	586	94.7
	Excluded ^a	33	5.3
	Total	619	100.0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
.722	6

Item Statistics

	Mean	Std. Deviation	N
یکنواختی پوشش مدرسه	2.3140	1.43356	586
شادی پوشش مدرسه	2.7765	1.42128	586
متانت پوشش مدرسه	3.9283	1.21598	586
دست و پاگیر بودن پوشش مدرسه	3.2235	1.54240	586
سختی پوشش مدرسه	3.2628	1.51884	586
زیبایی پوشش مدرسه	3.0887	1.44494	586

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha if Item Deleted
یکنواختی پوشش مدرسه	16.2799	24.923	.276	.735
شادی پوشش مدرسه	15.8174	21.418	.569	.648
متانت پوشش مدرسه	14.6655	27.119	.184	.750
دست و پاگیر بودن پوشش مدرسه	15.3703	21.639	.481	.675
سختی پوشش مدرسه	15.3311	20.198	.617	.630
زیبایی پوشش مدرسه	15.5051	20.732	.616	.633

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
18.5939	30.925	5.56106	6

الف- حجم کل نمونه‌ی مورد بررسی چند نفر بوده است و در عمل چند در صد از پاسخگویان از فرآیند سنجش پایایی حذف شده‌اند؟

۸۰ □ راهنمای کاربردی نرم افزار SPSS

- ب- مقدار کل پایایی مقیاس در چقدر است؟
- ج- کدام گویه‌ها ناهم وزن هستند؟
- د- کدام گویه‌ها در مقایسه با بقیه قدرت تمیز کمتری دارند؟
- ه- میانگین کل مقیاس چقدر است؟
- و- کمترین و بیشترین نمره‌ی مقیاس به لحاظ نظری چند می‌توانست باشد؟
- ز- با حذف کدام گویه‌ها قدرت تمیز مقیاس افزایش می‌یابد؟
- ح- کدام گویه‌ها همسازی درونی کمتری با نمره‌ی کل مقیاس دارند؟
- ط- با حذف کدام گویه‌ها مقدار آلفای کل مقیاس افزایش می‌یابد؟
- ی- برای ساختن شاخص نگرش نسبت به پوشش مدرسه باید از کدام گویه‌ها استفاده کرد؟

تمرین ۲- با توجه به فایل GSS93 SUBSET VALUE پایایی مقیاس گرایش به موسیقی (متغیرهای ردیف ۵۱ (bigband) تا ۶۱ (hvymental) را به روش آلفای کرانباخ محاسبه و تفسیر کنید.

تمرین ۳- با توجه به فایل satisf.sav پایایی مقیاس رضایت مشتری را با استفاده از متغیرهای ردیف ۱۳ (price) تا ۱۸ (overall) به روش آلفای کرانباخ محاسبه و تفسیر کنید.

تمرین ۴- با توجه به فایل داده‌های world95 پایایی مقیاس سطح بهداشت کشورهای جهان را با استفاده از متغیرهای امید به زندگی مردان و زنان و مرگ و میر کودکان و درصد مبتلایان به ایدز به روش آلفای کرانباخ محاسبه کنید.

درس ششم

شاخص سازی

در پایان این درس انتظار می رود دانشجو

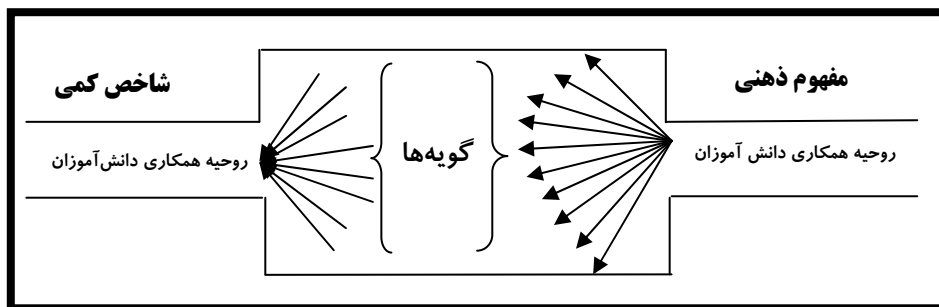
- ۱- مفهوم شاخص و مراحل شاخص سازی را بداند.
- ۲- با استفاده از دستور Compute بتواند شاخص بسازد.
- ۳- در هنگام شاخص سازی بتواند از عملگرها استفاده کند.
- ۴- در هنگام شاخص سازی به ماهیت و سطوح سنجش متغیرها توجه کند.
- ۵- بتواند از متغیرهای ناهموزن و ناهم دامنه شاخص تهیه کند.
- ۶- مفهوم نمره استاندارد و طرز ساختن آنها را بداند.
- ۷- بتواند از نمرات استاندارد در ساختن شاخص استفاده کند.
- ۸- بتواند صحت و سقم فرآیند شاخص سازی را بررسی کند.

۱- مقدمه

همانطور که در درس قبل اشاره شد، بسیاری از مفاهیم علوم انسانی به سادگی قد، وزن، سن، جنسیت نیستند که تنها با یک سؤال سنجیده شوند، بلکه پیچیدگی بسیاری از آنها به اندازه‌ای است که باید آنها را با بیش از یک سؤال سنجید. محققان موظف هستند بعد از سنجش پایایی و کسب اطمینان از دقت مقیاس و گویه‌های تشکیل‌دهنده‌ی آن، گویه‌ها را با یکدیگر ترکیب کرده و شاخص^۱ بسازند این فرایند ترکیب گویه‌ها را اصطلاحاً شاخص‌سازی می‌نامند.

برای مثال اگر محقق‌ی نمرات پاسخگویان را در مورد هر یک از گویه‌های مقیاس "روحیه همکاری دانش‌آموزان" با یکدیگر جمع کند و نمره‌ای با عنوان میزان روحیه همکاری بسازد، عملی را که انجام داده است، شاخص‌سازی می‌نامیم.

شکل شماره (۲۸) تجزیه متغیر طی تعریف عملیاتی و شاخص‌سازی



همچنین اگر قرار باشد مفهوم پایگاه اجتماعی با معرف‌های تعداد سال‌های تحصیل، میزان درآمد ماهانه، نوع شغل و همچنین عضویت در تشکلهای سیاسی- اجتماعی سنجیده شود، گویی مفهوم پایگاه تجزیه شده است. در نهایت باید بر مبنای قاعده‌ای معرف‌های تهیه شده را با یکدیگر ترکیب کرد. حاصل این فرآیند را شاخص می‌نامند.

۲- شاخص^۲ و شاخص‌سازی

باید به خاطر داشته باشیم که شاخص‌سازی یک ترکیب کردن ساده یا فی‌البداهه نیست بلکه طبق منطق خاصی انجام می‌شود و نیاز به پشتوانه‌ی نظری دارد. مثلاً نباید تصور کرد که

۱ - Index

۲- شاخص را نباید مترادف معرف (Indicator) تلقی کرد. شاخص از ترکیب معرف‌ها تهیه می‌شود. برای اطلاع بیشتر در این زمینه به کتاب روش‌های سنجش روایی و پایایی (حیدری چروده، ۱۳۸۹) مراجعه کنید.

درس ششم - شاخص □ ۸۳

شاخص همیشه از میانگین نمرات ساخته می‌شود، چنانچه بخواهید وضعیت اقتصادی خانواده را از طریق درآمد اعضای شاغل در خانواده بسنجید، باید به جای میانگین درآمد اعضای خانواده، مجموع درآمد اعضای شاغل خانواده را محاسبه کند. یا اگر بخواهید تأثیر میزان سواد زن و شوهر را بر ناسازگاری بین آنها بسنجد، باید به جای ساختن شاخصی با عنوان میانگین تحصیلات زوجین، شاخصی تحت عنوان فاصله‌ی تحصیلات زوجین تهیه کنید. از آنجا که شاخص‌سازی یکی از مراحل بسیار حساس در تحقیق است، از این رو ضروری است که قبل از شاخص‌سازی به موارد زیر توجه داشته باشید.

- **ماهیت متغیرها:** یکی از نکات مهم در ابتدای شاخص‌سازی، بررسی ماهیت متغیرهایی است که قرار است طی فرآیند شاخص‌سازی با هم ترکیب شوند. در این مرحله باید اطمینان یافت که متغیرها به لحاظ نظری و محتوایی قابلیت ترکیب دارند تا سرانجام متغیر جدیدی که در فرآیند شاخص‌سازی ساخته می‌شود، موجه باشد.

به عنوان مثال می‌توان از متغیرهای "امید به زندگی زنان و مردان در جهان"، مرگ‌ومیر کودکان در جهان" و "نسبت مبتلایان به ایدز در جهان"، شاخصی به عنوان "سطح بهداشت کشورها" ساخت. همان‌طور که می‌بینید متغیر نخست، ماهیتاً مثبت است و بالا بودن آنها نشانه‌ی وضعیت مناسب سطح بهداشت است، اما در متغیرهای دوم و سوم، پایین بودن میزان‌های مرگ‌ومیر کودکان و ابتلا به ایدز، مطلوب بوده و نشانه وضعیت مناسب سطح بهداشت است. لذا باید برای ترکیب این سه متغیر و ساختن شاخص "سطح بهداشت کشورها" آنها را با یکدیگر هماهنگ و هم جهت کرد.

- **سطح سنجش:** متغیرهایی که قرار است از ترکیب آنها یک شاخص واحد تهیه شود، باید سطوح سنجش یکسانی هم داشته باشند. یعنی نمی‌توان یک متغیر سطح اسمی را با یک متغیر رتبه‌ای و یک متغیر فاصله‌ای ترکیب کرد. مثلاً برای ساختن پایگاه اجتماعی اقتصادی دانشجویان، میزان پول توجیبی روزانه‌ی آنان را به عنوان یک متغیر فاصله‌ای، با نوع وسیله‌ی نقلیه مورد استفاده‌ی آنها در سطح شهر (شخصی، تاکسی، اتوبوس و ...) و دانشگاه محل تحصیل‌شان (دولتی، آزاد، پیام نور و علمی کاربردی)، نمی‌توان به سادگی با هم ترکیب نمود و ادعا کرد که پایگاه اجتماعی اقتصادی دانشجویان ساخته شده است. به عبارت دیگر محقق قبل از ساختن شاخص و حتی به هنگام طراحی سؤالات برای سنجش یک مفهوم، باید بر نحوه شاخص‌سازی مفهوم توجه داشته باشد.

- **دامنه تغییرات:** متغیرهایی که قصد دارید با یکدیگر ترکیب کنید باید دامنه‌ی تغییرات یکسانی داشته باشند. برای مثال برای ساختن شاخص توانایی علمی دانشجویان نمی‌توان نمرات مربوط به درس زبان را که به صورت درصدی اندازه‌گیری شده است با نمره‌ی آمار که از صفر تا بیست است و نمره SPSS که از صفر تا ده است، با یکدیگر ترکیب کرد، یا برای ساختن شاخص وضعیت اقتصادی نمی‌توان درآمد را که دارای دامنه‌ای از صد هزار تا دو

۸۴ □ راهنمای کاربردی نرم افزار SPSS

میلیون تومان است با بهره‌مندی از امکانات زندگی که در دامنه‌ای بین صفر تا ۱۵ قرار دارد، با هم ترکیب کرد. در چنین مواردی باید با تقلیل سطح سنجش^۱ (از طریق دستور (Recode)) متغیرها یا در صورت امکان با ساختن نمرات استاندارد Z^۲ یا T^۳، متغیرها را هم سطح کرده، سپس با هم ترکیب کرد.

۳- روش شاخص‌سازی از متغیرهای ناهم دامنه

همانگونه که اشاره گردید، اگر عناصر تشکیل‌دهنده‌ی شاخص با یکدیگر هم دامنه نباشند، بهتر است برای هر یک از متغیرهای تشکیل‌دهنده‌ی شاخص، یک نمره استاندارد (Z) تهیه کنید، سپس برای شاخص‌سازی با استفاده از متغیرهای استاندارد شده، اقدام لازم را انجام دهید. برای این کار باید مطابق تصویر زیر عمل کنید.

شکل شماره (۲۹) شیوه استاندارد سازی متغیر (تهیه نمره Z)



۱- این روش به دلیل اینکه قدرت تمیز متغیرها را کاهش می‌دهد در شرایطی که امکان ساختن متغیرهای استاندارد (Z) وجود دارد، توصیه نمی‌شود. با استاندارد سازی اکثریت قریب به اتفاق اعضای نمونه‌ی در دامنه خاصی (مثلاً در نمره Z بین -۳ و +۳) قرار می‌گیرند. توجه داشته باشید که استاندارد سازی، وضعیت "توزیع" متغیر را تغییر نمی‌دهد (نرمال نمی‌کند)

$$2- z = \frac{x_i - \bar{x}}{s}$$

$$3- T = 10z + 50$$

درس ششم - شاخص □ ۸۵

اگر محققى بخواهد از ترکیب سه متغیر "درصد باسوادان زن"، "نرخ باروری" و "امید به زندگی زنان"، شاخص "میزان توسعه یافتگی وضعیت زنان" را بسازد. چون این متغیرها با یکدیگر هم دامنه نیستند، باید از این سه متغیر، نمره استاندارد Z تهیه کند که چگونگی تهیهی آنها در تصویر قبل نشان داده شد.

همانگونه که در تصاویر زیر نشان داده شده است، با اجرای این دستور در آخرین ردیف کاربرگ متغیرها (variable view) و همچنین در آخرین ستون کاربرگ داده‌ها (Data view)، متغیرهای استاندارد شده با علامت Z به مجموعه متغیرها اضافه شده است.

شکل شماره (۳۰) نمره استاندارد Z در کاربرگ متغیرها (variable view)

	Name	Type	W...	De...	Label	Values
27	Zurban	Numeric	11	5	Zscore: People living in cities...	None
28	Zfertility	Numeric	11	5	Zscore: Fertility: average number of children...	None
29	Zlit_fema	Numeric	11	5	Zscore: Females who read (literacy)...	None

شکل شماره (۳۱) نمره استاندارد Z در کاربرگ داده‌ها (Data view)

	Zurban	Zfertility	Zlit_fema
1	-1.59184	1.75399	-1.86171
2	1.21769	-0.40110	0.96972
3	0.47399	-0.19610	1.14450

همچنین با اجرای این دستور، این خروجی زیر ارائه خواهد شد که آماره‌های توصیفی هر یک از متغیرهای تشکیل‌دهنده شاخص "میزان توسعه یافتگی" را نشان می‌دهد.
جدول شماره (۱۶) توصیف شاخص "میزان توسعه یافتگی"

	N	Minimum	Maximum	Mean	Std. Deviation
درصد باسوادان	107	18	100	78.34	22.883
امید به زندگی زنان	109	43	82	70.16	10.572
نرخ باروري	107	1.3	8.2	3.563	1.9025
Valid N (listwise)	105				

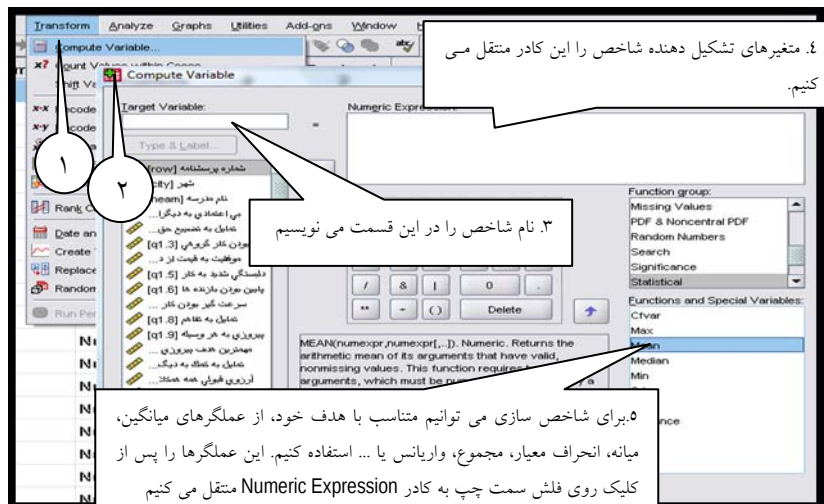
برای ساختن شاخص از متغیرهای ناهم‌دامنه به ادامه درس توجه کنید.

۴- شاخص سازی با استفاده از دستور Compute

رایج ترین ابزار برای شاخص سازی در SPSS استفاده از دستور Compute از طریق منوی Transform است. مثال درس قبل را بخاطر بیاورید که محقق به منظور سنجش روحیه همکاری دانش آموزان، مقیاسی مرکب از ۱۵ گویه را طراحی کرد. از آنجا گویه شماره ۵، در فرآیند سنجش پایایی نامناسب تشخیص داده شد، لذا برای ساختن شاخص کمی روحیه همکاری دانش آموزان، باید تمامی گویه ها به غیر از گویه ۵ را مد نظر داشت باشد و برای تهیه شاخص، آنها را با یکدیگر ترکیب کند.

برای اجرای این دستور باید به ترتیب مندرج در تصویر عمل شود.

شکل شماره (۳۲) شیوه اجرای دستور Compute برای شاخص سازی (مرحله ۱)



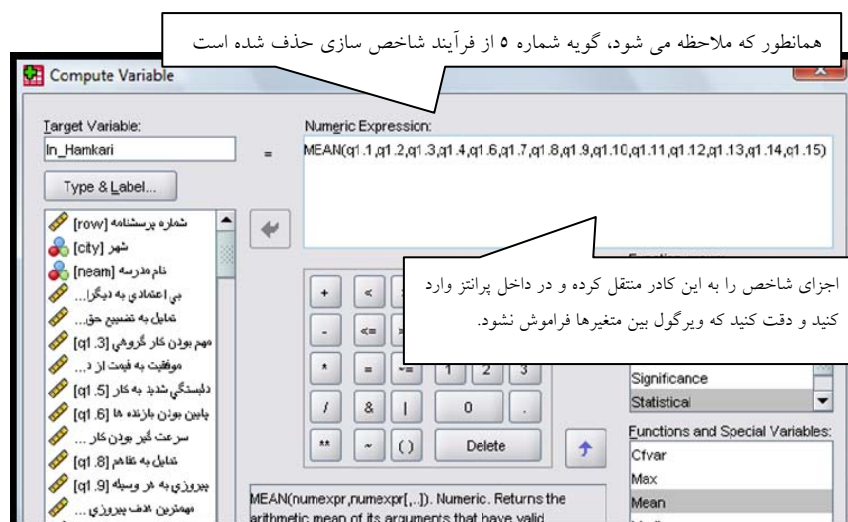
در شاخص سازی به شیوه ابتدایی بسیاری از افراد عنوان متغیرها را در کادر Numeric Expression نوشته و به تناسب قاعده شاخص سازی، عملیات جبری مربوط را در آن جا می نویسند. مثلاً برای میانگین گرفتن از پنج متغیر q1 تا q5 به این شکل عمل می کنند: $(q1+q2+q3+q4+q5/5)$

اما در این شکل از شاخص سازی، نمره شاخص تنها برای آن دسته از اعضای نمونه محاسبه می شود که به تمام سؤالات مربوط به متغیرهای q1 تا q5 پاسخ داده باشند، لذا آن دسته از پاسخگويانی که به یک یا چند مورد از سؤالات مربوطه پاسخ نداده اند از فرآیند شاخص سازی حذف می شوند و در نتیجه برای این دسته از پاسخگويان به جای نمره شاخص، در فایل داده ها مقادیر نامشخص (system missing) منظور می شود.

درس ششم - شاخص □ ۸۷

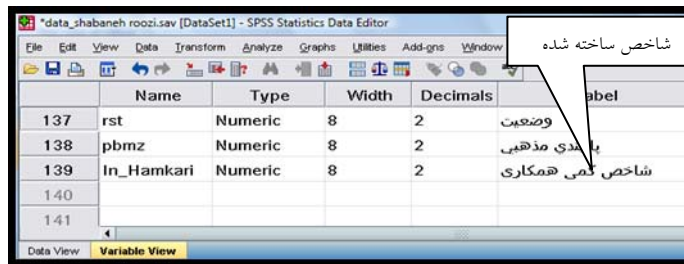
اما با استفاده از کادرهای **Function Group** و **Functions and Special Variables** می‌توان این مشکل را حل کرد. به این ترتیب که اگر پاسخگویی به ۳ سؤال از ۵ سؤال پاسخ داده باشد، برای محاسبه‌ی میانگین، جمع نمرات او تقسیم بر ۳ می‌شود و اگر کسی فقط به دو سؤال جواب داده باشد، مجموع نمرات او بر عدد ۲ تقسیم خواهد شد. به این ترتیب همه‌ی افراد نمونه، نمره‌ای در شاخص خواهند داشت مگر کسانی که به هیچ سؤالی جواب نداده باشند.

در مثال مربوط به مقیاس "روحیه همکاری دانش‌آموزان"، چنانچه محقق بخواهد این شاخص را از طریق میانگین گویه‌های مقیاس محاسبه کند، طبق راهنمای زیر باید عمل کند. شکل شماره (۳۳) شیوه اجرای دستور **Compute** برای ساختن شاخص (مرحله ۲)



سپس با کلیک بر روی **Ok** شاخص مزبور به عنوان یک متغیر جدید در آخرین ستون کاربرگ داده‌ها و نیز در آخرین سطر کاربرگ متغیرها نظیر شکلی که در تصویر زیر می‌بینید، اضافه خواهد شد. به تصاویر بعد دقت کنید.

شکل شماره (۳۴) شاخص جدید ساخته شده در کاربرگ متغیرها (variable view)



شکل شماره (۳۵) شاخص جدید ساخته شده در کاربرگ داده‌ها (Data view)

	rst	pbmz	In_Hamkari
1	0.00	5.56	2.57
2	0.00	12.78	4.43
3	0.00	7.78	4.21

سپس برای اطمینان از صحت و سقم شاخص تهیه شده، یک جدول فراوانی از آن تهیه کرده و مقادیر محاسبه شده را باید بررسی کرد تا خطایی در آن وجود نداشته باشد. همانگونه که ملاحظه می‌شود اولین خروجی نشان می‌دهد شاخص کمی "روحیه همکاری دانش‌آموزان" برای تمامی اعضای نمونه محاسبه شده است. حال آنکه طبق اولین خروجی آلفای کرونباخ در درس قبل تنها برای ۱۱۰۹ نفر از پاسخگویانی که به تمامی سؤالات پاسخ داده بودند، محاسبه شده بود. تفاوت مزبور، حاصل استفاده از عملگرهای محاسباتی کادر **Function Group** است که توضیحات آن در صفحات پیشین آمد. همانگونه که گفته شد استفاده از این عملگرهای محاسباتی باعث می‌شود حتی اگر پاسخگویی فقط به یک سؤال جواب داده باشد، نمره‌ی شاخص برای او محاسبه شود.

جدول شماره (۱۷) خلاصه پردازش اطلاعات

Hamkari Statistics R		
N	Valid	1215
	Missing	0

درس ششم - شاخص □ ۸۹

جدول شماره (۱۸) توزیع فراوانی شاخص کمی "روحیه همکاری دانش‌آموزان"

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 1	1	.1	.1	.1
1 .20	1	.1	.1	.2
1 .67	2	.2	.2	.3
1 .73	.	.	.	.
.	.	.	.	.
.	.	.	.	.
4	5	.4	.4	99.6
4 .60	1	.1	.1	99.7
4 .71	4	.3	.3	100.0
4 .73				
Total	1215	100.0	100.0	

توصیه می‌شود برای اطمینان بیشتر، نمره‌ی شاخص را برای دو یا سه نفر از اعضای نمونه به شیوه‌ی دستی محاسبه و با نمرات محاسبه شده توسط نرم‌افزار در ستون آخر پنجره‌ی کاربرگ داده‌ها برای همان افراد نیز مقایسه کرد.

لازم به ذکر است که چنانچه اجزای شاخص هم‌وزن نباشند^۱ می‌توان به متغیرهای موردنظری که در داخل کادر **Numeric Expression** وارد شده‌اند، ضریب مناسب داد.^۲ به عنوان مثال محقق‌ی را در نظر بگیرید که قصد دارد از ترکیب نمرات آمار (X1)، روش تحقیق نظری (X2) و روش تحقیق عملی (X3)، شاخصی با عنوان مهارت پژوهش بسازد اما به دلایل تئوریک یا بنابه صلاحدید استاد ناظر تحقیق باید برای نمره درس‌های روش تحقیق عملی ضریب ۳، روش تحقیق نظری ضریب ۲ قائل شده و آنگاه شاخص مهارت پژوهش را از طریق محاسبه میانگین نمرات وزن داده شده دروس مورد اشاره تهیه کند. شکل زیر بیانگر همین محاسبات است.

۱- ناهم‌وزن بودن ناظر به شرایطی است که برخی اجزای شاخص اهمیت بیشتری از سایر اجزا داشته باشند. مثلاً ممکن است دلایل متقاعدکننده‌ای وجود داشته باشد که در ساختن شاخص رضایت شغلی، رضایت از ارتقاء یا درآمد اهمیت بیشتری در مقایسه با رضایت از همکاران یا شرح وظایف داشته باشد.

۲- برای ضریب دادن باید متغیر موردنظر را داخل پرانتز گذاشته و بین ضریب و نام متغیر موردنظر علامت ضرب یا همان ستاره را قرار داد.

۹۰ □ راهنمای کاربردی نرم‌افزار SPSS

شکل شماره (۳۶) نحوه ساختن شاخص از متغیرهای ناهموزن



در پایان این مبحث مراحل مختلف عملیات شاخص‌سازی را فهرست وار مرور می‌کنیم:

- ۱- ابتدا از تمامی متغیرهایی که قرار است شاخص از آنها تهیه شود، جدول فراوانی تهیه کنید تا متغیرها از حیث ماهیت، هم‌دامنه بودن و سطح سنجش آنها واری‌شود.
- ۲- سپس پایایی اجزای مزبور را از طریق آلفای کرانباخ محاسبه کرده و متغیرهای مناسب را شناسایی کنید.
- ۳- قاعده‌ی ترکیب متغیرها را معین کنید و از عملگرهای مناسب آماری استفاده کنید.
- ۴- براساس قاعده‌ی ترکیب و با استفاده از دستور **Compute** شاخص مربوطه را تهیه کنید.
- ۵- از شاخص کمی، جدول فراوانی گرفته و مقادیر بدست آمده را واری‌شود.
- ۶- برای دو یا چند نفر از اعضای نمونه نمره‌ی شاخص را به شیوه‌ی دستی محاسبه سپس با نمره‌ی بدست آمده از طریق نرم‌افزار مقایسه کنید.

درس ششم - شاخص □ ۹۱

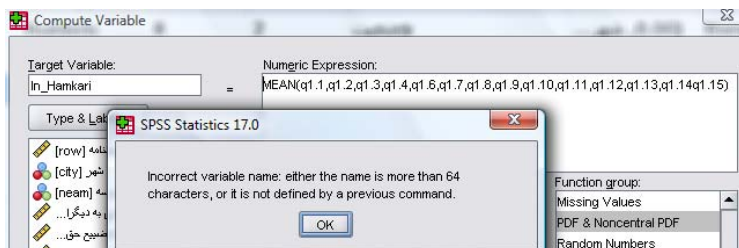
تمرین درس ششم

تمرین ۱- در تحقیقی که در باره‌ی مصرف کالاهای فرهنگی است، محقق می‌خواهد شاخص کمی مطالعه را از طریق ترکیب سه معرف "مدت مطالعه روزنامه"، "مدت مطالعه مجله" و "مدت مطالعه کتاب" تهیه کند. شما چه شیوه‌ی ترکیبی برای این سه متغیر پیشنهاد می‌کنید؟ به عبارت دقیق‌تر از عملگرهای کادر Functions and Special Variables کدامیک را پیشنهاد می‌کنید؟ چرا؟

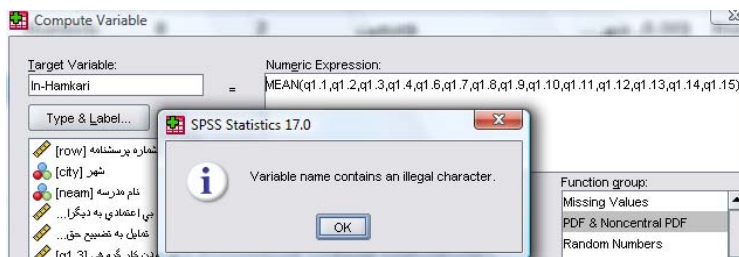
تمرین ۲- در تحقیقی در باره‌ی نگرش سیاسی دانشجویان، محقق می‌خواهد برای انسجام بین مواضع سیاسی سه‌گانه‌ی دانشجویان نسبت به قوه‌های مجریه، مقننه و قضائیه (به عنوان سه متغیر جداگانه)، شاخص تهیه کند. شما برای این کار کدامیک از عملگرهای کادر Functions and Special Variables را پیشنهاد می‌کنید؟ چرا؟

تمرین ۳- پیام‌های زیر در پاسخ به چه خطایی صادر شده اند؟

الف-



ب-



۹۲ □ راهنمای کاربردی نرم افزار SPSS

تمرین ۴- با توجه به فایل GSS93 SUBSET VALUE شاخص کمی "متوسط گرایش به موسیقی" را از متغیرهای bigband, bluegrass, country, blues, musicals, classicl, folk, jazz, opera, rap, hvymetal در بین پاسخگویان بسازید و شاخص مزبور را به صورت یک شاخص کیفی جدید در قالب ۳ طبقه، دسته‌بندی کنید.

تمرین ۵- با توجه به فایل survey_sample.sav شاخص "بالاترین سطح سواد در خانواده" را از طریق متغیرهای Educ, paeduc, maeduc, speduc بسازید.

تمرین ۶- با توجه به فایل satisf.sav شاخص کمی و کیفی (در پنج طبقه)، همگنی (یا یکدست بودن) و رضایت مشتری را از طریق متغیرهای price, numitems, org, services, quality, overall بسازید.

تمرین ۷- با توجه به فایل داده‌های world95 شاخص کمی "سطح بهداشت کشورهای جهان" را از طریق متغیرهای امید به زندگی مردان و زنان و "مرگ‌ومیر کودکان" و "درصد مبتلایان به ایدز" بسازید.

بخش دوم

**توصیف نتایج
به کمک روش‌های آمار توصیفی**

درس هفتم

توصیف یک متغیری نتایج

در پایان این درس انتظار می رود دانشجو

- ۱- بتواند اهمیت و ضرورت تهیه پیش‌نویس فهرست مطالب یافته‌های تحقیق را بداند.
- ۲- نحوه تهیه جدول فراوانی یک متغیری را بداند و بتواند آن را تفسیر کند.
- ۳- نمودارهای یک متغیری دایره‌ای، میله‌ای و مستطیلی را تهیه کند.
- ۴- بتواند برای متغیرها و شاخص‌های کمی تحقیق خود، آماره‌های مرکزی، پراکندگی، وضعی، چولگی و کشیدگی را تهیه و تفسیر کند.

۱- مقدمه

پس از ورود و ویرایش داده‌ها و نیز تهیه شاخص‌های کمی و کیفی، نوبت به تجزیه و تحلیل داده‌ها، می‌رسد. در این مرحله محقق باید براساس سؤالات و فرضیات تحقیق و همچنین با استفاده از قدرت خلاقیت خود به تجزیه و تحلیل اطلاعات بپردازد. یکی از شایع‌ترین اشتباهات دانشجویان و محققان مبتدی آن است که پس از ورود داده‌ها، غالباً فی‌البداهه و بدون هیچ طرح و نقشه‌ای، از متغیرهای مختلفی که در فایل داده‌ها وجود دارد، خروجی گرفته و شروع به گزارش‌نویسی می‌کنند. اما این رویکرد نسبت به تجزیه و تحلیل اطلاعات موجب اتلاف زمان و نیز بی‌نظمی و گاهی دوباره کاری می‌شود.

از این رو توصیه می‌شود که قبل از هر گونه تحلیل آماری، ابتدا پیش‌نویس فهرستی از مطالب مربوط به فصل یافته‌های تحقیق را آماده کنید و روش‌های آماری لازم و نیز نوع نمودارهای مورد نیاز را نیز با ذکر جزئیات مهم آن در مقابل هر یک از بندهای فهرست بنویسد. سپس بهتر است فهرست مزبور را در اختیار استاد ناظر یا راهنما و نیز افراد صاحب‌نظر قرار داده تا نسبت به محتوا و نیز ترتیب و نحوه‌ی تنظیم آن قضاوت کنند. سرانجام پس از نهایی شدن فهرست، اقدام به گرفتن خروجی‌های موردنیاز از طریق نرم‌افزار SPSS کنید. با انجام این مهم، از پراکندگی و سرگردانی ذهنی نجات یافته و به نظامی منطقی در توصیف نتایج خواهید رسید.

همانطور که می‌دانید ساختار یافته‌های اغلب تحقیقات فرهنگی و اجتماعی دارای دو بخش

عمده است :

۱- توصیف نتایج ۲- تحلیل نتایج

در بخش توصیف عموماً به تنظیم جداول فراوانی یک یا دو متغیری مربوط به ویژگی‌های جمعیت‌شناختی نمونه یا متغیرهای اصلی تحقیق و تفسیر شاخص‌های مرکزی و پراکندگی پرداخته می‌شود و بخش تحلیل نتایج، به بررسی روابط بین متغیرها و پاسخگویی به سؤالات و آزمون فرضیات اختصاص دارد.

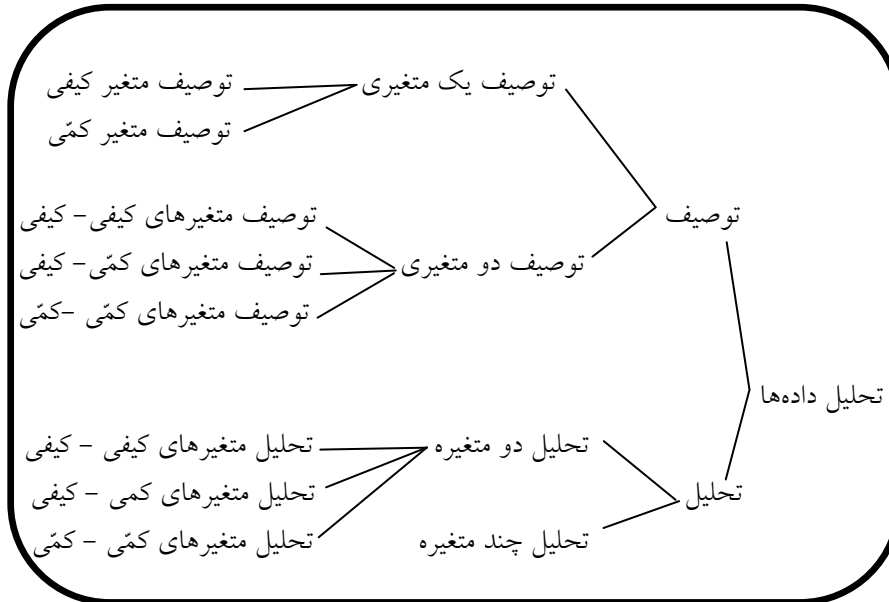
جدول زیر یکی از انواع پیش‌نویس‌های مربوط به یافته‌های تحقیق است. تهیه چنین فهرستی اگر چه خیلی وقت گیر است، اما تجربه نشان داده است که وجود چنین چارچوبی موجب تسریع در گزارش نویسی نهایی و تجزیه و تحلیل داده‌ها می‌شود. بدیهی است شرط لازم برای تهیه چنین فهرستی، تسلط نسبی به آمار توصیفی و استنباطی یا بهره‌گیری از مشاوره‌ی متخصصان آماری است که دارای تجربه پژوهش نیز باشند.

پیش‌نویس تهیه شده از فهرست یافته‌های تحقیق

الف- سیمای جمعیت شناختی نمونه مورد بررسی	
الف-۱- جنسیت	جدول فراوانی + نمودار دایره‌ای
الف-۲- رشته تحصیلی	جدول فراوانی + نمودار میله‌ای
الف-۱-۲- جنسیت به تفکیک رشته تحصیلی	جدول فراوانی دو بعدی
الف-۳- نوع مدرسه	جدول فراوانی + نمودار میله‌ای
الف-۱-۳- جنسیت به تفکیک نوع مدرسه	جدول فراوانی دو بعدی
الف-۴- تحصیلات والدین	جدول فراوانی + نمودار میله‌ای
الف-۵- نوع شغل والدین	جدول فراوانی + نمودار میله‌ای
ب- توصیف نتایج تحقیق	
ب-۱- شاخص کمی روحیه همکاری دانش‌آموزان	آماره‌های مرکزی و پراکندگی + نمودار هیستوگرام + جدول فراوانی از شاخص کیفی
ب-۲- شاخص کمی انسجام درون گروهی دانش‌آموزان	آماره‌های مرکزی و پراکندگی + نمودار هیستوگرام + جدول فراوانی از شاخص کیفی
ب-۳- شاخص کمی افسردگی دانش‌آموزان	آماره‌های مرکزی و پراکندگی + نمودار هیستوگرام + جدول فراوانی از شاخص کیفی
ب-۴- شاخص کمی رغبت به تحصیل	آماره‌های مرکزی و پراکندگی + نمودار هیستوگرام + جدول فراوانی از شاخص کیفی
ج- تحلیل نتایج تحقیق	
ج-۱- رابطه جنسیت با روحیه همکاری دانش‌آموزان	آزمون نرمال بودن، T-Test یا یومان ویتنی + نمودار نمودار میله‌ای مقایسه‌ای
ج-۲- رابطه رشته تحصیلی با رغبت به تحصیل دانش‌آموزان	آزمون نرمال بودن، Anova یا کراسکال والیس + نمودار میله‌ای مقایسه‌ای
ج-۳- تأثیر سواد مادر بر افسردگی دانش‌آموزان	آزمون نرمال بودن، ضریب همبستگی اسپیرمن یا پیرسون + نمودار پراکنش

همانطور که ملاحظه می‌شود برخی از بندهای فهرست فوق مربوط به توصیف و برخی دیگر مربوط به تحلیل است. دقت در بندهای توصیف نشان می‌دهد که توصیف می‌تواند یک متغیره و یا دو متغیره باشد اما تحلیل نتایج، امری است که غالباً دو متغیره است و نیاز به آزمون فرضیه یا پاسخگویی به سؤالات دارد. نمودار زیر تقسیم‌بندی فوق را از تحلیل داده‌ها به اختصار نشان می‌دهد.

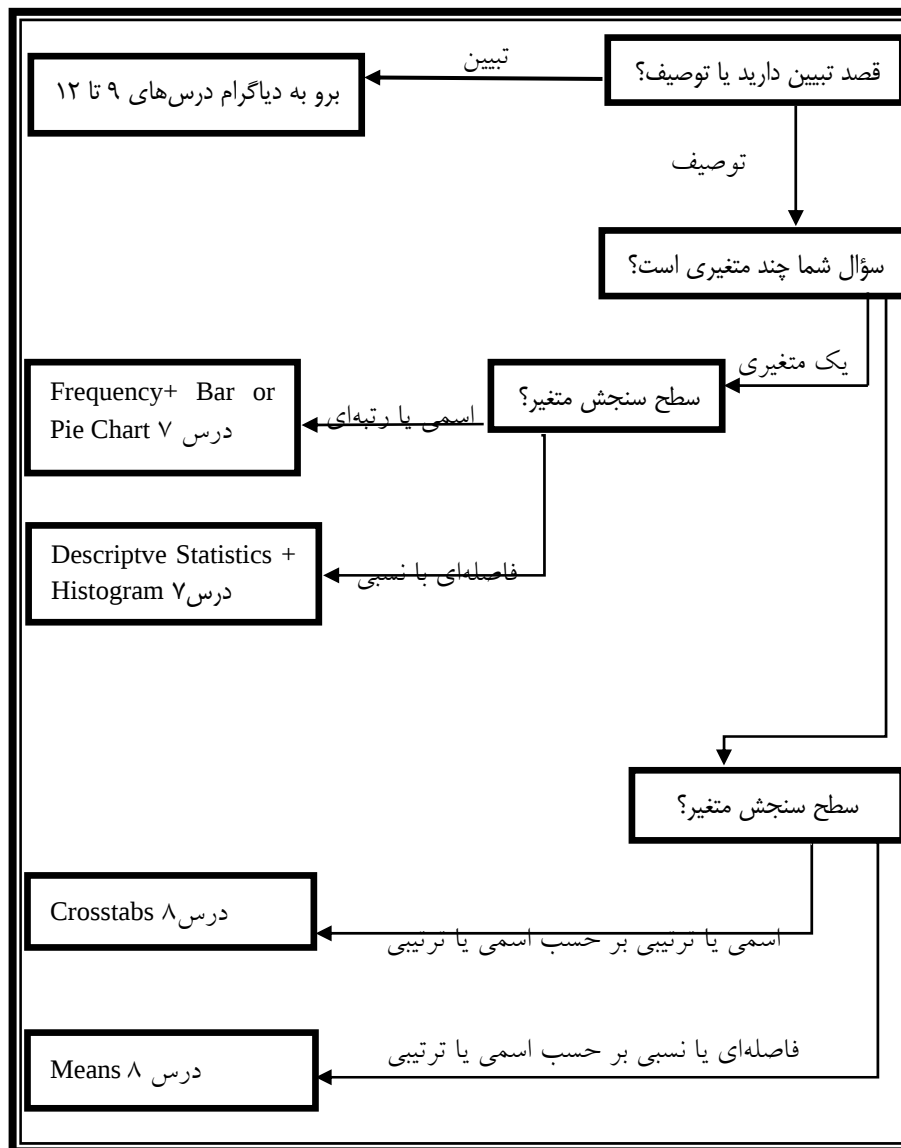
شکل شماره (۳۷) اشکال مختلف تحلیل داده‌ها در گزارش‌های پژوهشی



در اینجا برای ارائه یک تصویر کلی از روش‌های آماری متناسب با متغیرها و نیز دستورات متناظر با آنها در نرم‌افزار SPSS از یک دیاگرام استفاده می‌شود. پایه اصلی تصمیم‌گیری در انتخاب مسیر صحیح آشنایی با سطوح سنجش و هدف محقق است.

۱- این شکل از توصیف به دلیل آنکه کمتر رایج است، در کتاب حاضر مورد بحث قرار نمی‌گیرد.

شکل شماره (۳۸) دیاگرام انتخاب روش مناسب برای توصیف نتایج



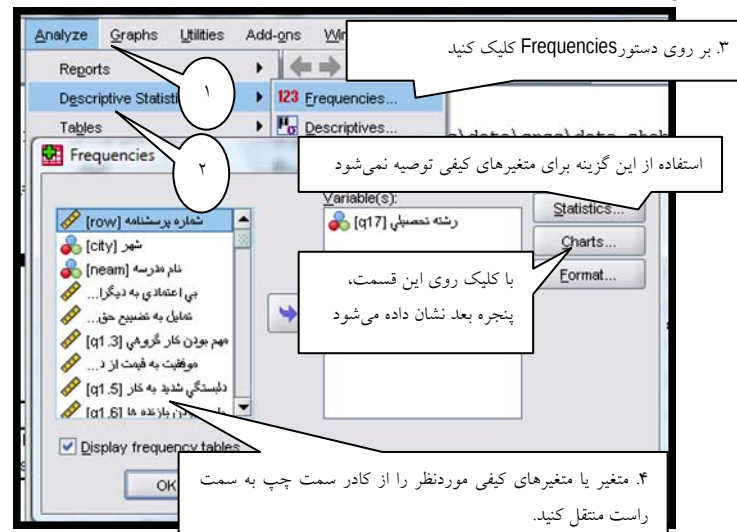
۲- توصیف یک متغیری

از آنجا که توصیف متغیرهای کمی با کیفی تفاوت دارد در این قسمت به توضیح هر یک از این دو نوع توصیف، به طور جداگانه می‌پردازیم.

• توصیف یک متغیر کیفی

هرگاه بخواهید متغیرهای کیفی (در سطح سنجش اسمی یا رتبه‌ای) را توصیف کنید یا به عبارت دیگر در صدد پاسخ دادن به سؤالی نظیر این باشید که "توزیع فراوانی مطلق، درصدی یا تجمعی متغیر X چگونه است؟"، باید طبق شکل زیر عمل کنید.

شکل شماره (۳۹) شیوه‌ی اجرای دستور frequencies برای توصیف متغیرهای کیفی



با کلیک بر روی Charts، پنجره‌ی زیر باز می‌شود که در آن باید نوع نمودار موردنیاز را از بین انواع سه‌گانه میله‌ای، دایره‌ای، یا هیستوگرام (با یا بدون خط نرمال) و نیز نوع فراوانی مورد استفاده (مطلق یا درصدی) را تعیین کرد. یادآوری می‌شود که از نمودارهای میله‌ای و دایره‌ای صرفاً برای نمایش متغیرهای کیفی استفاده می‌شود. همچنین توصیه می‌شود به دلیل سرعت بالای انتقال پیام به خواننده، به جای استفاده از فراوانی مطلق از درصد استفاده شود.

درس هفتم - توصیف یک متغیری نتایج □ ۱۰۱

شکل شماره (۴۰) انتخاب نمودار مناسب برای نمایش متغیر



اگر محقق بخواهد بداند "توزیع فراوانی رشته تحصیلی در بین دانش‌آموزان چگونه است؟" در صورت اجرای دستور frequencies با خروجی‌های زیر مواجه خواهد شد:

جدول شماره (۱۹) توزیع فراوانی رشته تحصیلی دانش‌آموزان

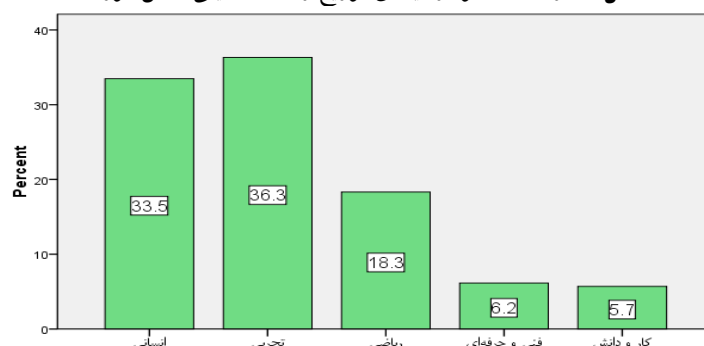
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	انسانی	223	18.4	33.5	33.5
	تجربی	242	19.9	36.3	69.8
	ریاضی	122	10.0	18.3	88.1
	فنی و حرفه‌ای	41	3.4	6.2	94.3
	کار و دانش	38	3.1	5.7	100.0
	Total	666	54.8	100.0	
Missing	فاقد رشته	549	45.2		
	Total	1215	100.0		

جدول فوق نشان می‌دهد از مجموع ۱۲۱۵ نفر دانش‌آموز دبیرستانی مورد بررسی، بیشترین درصد (۱۹/۹ درصد) در رشته تجربی و کمترین درصد (۳/۱ درصد) در شاخه کاردانش

۱۰۲ □ راهنمای کاربردی نرم‌افزار SPSS

مشغول تحصیل هستند، لازم به ذکر است که ۴۵/۲ درصد از اعضای نمونه به علت اشتغال به تحصیل در پایه‌ی اول دبیرستان، در موقعیت انتخاب رشته تحصیلی خاصی قرار نگرفته‌اند. همانگونه که ملاحظه می‌شود ستون "درصد معتبر" (Valid Percent) ناظر به همین واقعیت است. در این ستون مبنای درصدگیری، دانش‌آموزان دارای رشته تحصیلی است. مثلاً دانش‌آموزان رشته‌ی تجربی اگر چه ۱۹/۹ درصد کل نمونه را به خود اختصاص داده‌اند اما در بین دانش‌آموزانی که انتخاب رشته نموده‌اند، درصدشان بالغ بر ۳۶/۳ درصد است.

شکل شماره (۴۱) نمودار میله‌ای توزیع رشته تحصیلی دانش‌آموزان



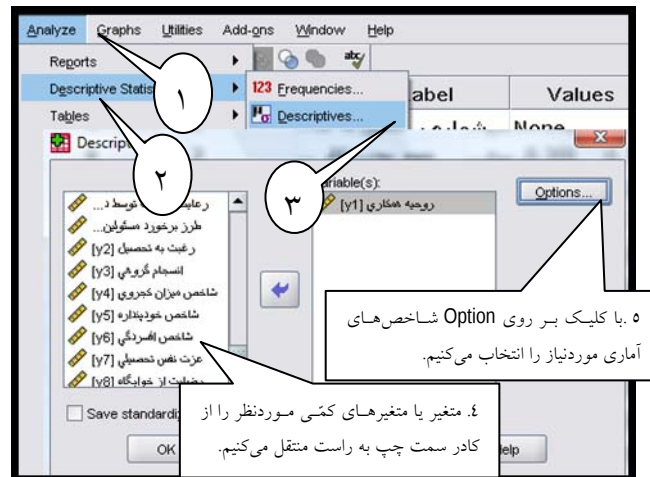
• توصیف یک متغیر کمی

اگر بخواهید متغیرهای کمی (در سطح سنجش فاصله‌ای یا نسبی) را توصیف کنید یا به عبارت دیگر درصدد پاسخ دادن به سؤالی نظیر این باشد که "میزان متغیر X چگونه است؟" باید به شرح تصویر زیر عمل کنید و به کمک شاخص‌های مرکزی، پراکندگی، کشیدگی و چولگی، متغیر یا متغیرهای کمی مورد نظر را از زوایای مختلف، توصیف کنید^۱.

۱- همانگونه که در بخش شاخص‌سازی نیز اشاره شد در بسیاری از موارد نمی‌توان متغیرهای کمی و یا شاخص‌ها را از طریق جداول فراوانی مطلق و درصدی توصیف کرد چرا که تعدد و تنوع بسیار زیاد مقادیر این متغیرها به اندازه‌ای است که منجر به تشکیل جداول طولانی با درصدهای پراکنده خواهد شد. برای یادآوری به درس چهارم "کدگذاری مجدد و دسته‌بندی" و مثالهای ارائه شده در این درس در مورد تبدیل متغیر کمی به کیفی مراجعه کنید.

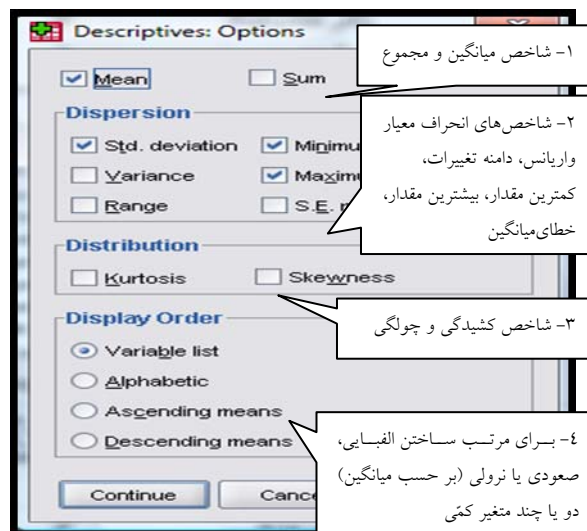
درس هفتم - توصیف یک متغیری نتایج □ ۱۰۳

شکل شماره (۴۲) شیوه اجرای دستور Descriptives برای توصیف متغیرهای کمی



با کلیک روی Option پنجره زیر باز می‌شود و امکانات زیر در دسترس محقق قرار می‌گیرد:

شکل شماره (۴۳) شیوه انتخاب شاخص‌های آماری مناسب برای متغیر



اگر بخواهید به این سؤال پاسخ دهید که "وضعیت روحیه همکاری دانش‌آموزان در نمونه‌ی مورد بررسی چگونه است؟" از آنجا که متغیر "روحیه همکاری دانش‌آموزان" براساس

۱۰۴ □ راهنمای کاربردی نرم افزار SPSS

توضیحات قبلی یک شاخص کمی است، با اجرای دستور فوق، با خروجی زیر روبرو خواهید شد.

جدول شماره (۲۰) توصیف شاخص های آماری "روحیه همکاری دانش آموزان"

		Statistic	Std. Error
روحیه همکاری	N	1215	
	Minimum	1.21	
	Maximum	5.00	
	Mean	3.7117	.01407
	Std. Deviation	.49027	
	Skewness	-.644	.070
	Kurtosis	1.272	.140
Valid N (listwise)		N	1215

جدول فوق نشان می دهد که حجم نمونه ی مورد بررسی برای سنجش روحیه همکاری ۱۲۱۵ نفر از دانش آموزان بوده است. نتایج مربوط به شاخص کمی روحیه همکاری نشان می دهد که کمترین میزان این متغیر ۱/۲۱ و بیشترین مقدار آن ۵ است. آماره ی میانگین نیز نشان می دهد که متوسط روحیه همکاری در بین نمونه ی مورد بررسی ۳/۷۱ (متوسط رو به بالا) است. انحراف معیار نیز بیانگر آن است که متوسط پراکندگی روحیه همکاری دانش آموزان از میانگین ۰/۴۹ است. ضریب چولگی^۱ بدست آمده منفی و شدید (۰/۶۴-) است که این امر بیانگر این واقعیت است که روحیه همکاری اکثریت اعضای نمونه بالاتر از میانگین یاد شده است. ضریب کشیدگی نیز مثبت و شدید (۱/۲۷) است که نشان می دهد شباهت نمونه ی مورد بررسی از لحاظ روحیه همکاری بیشتر از حدّ نرمال است.

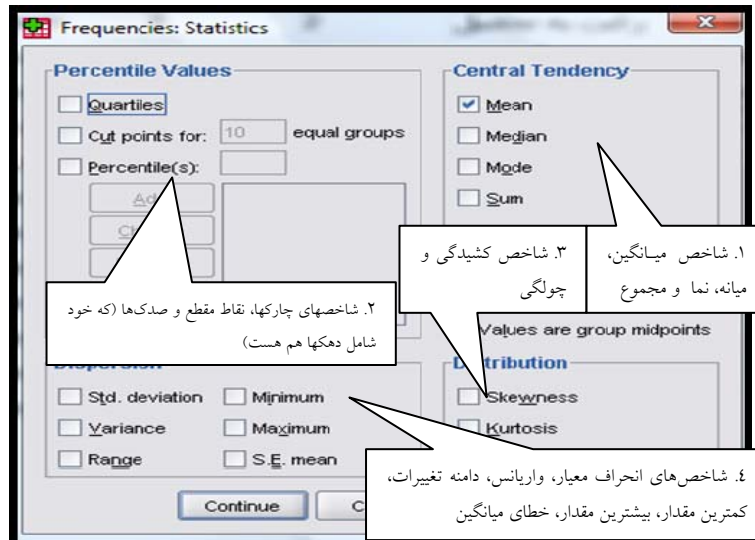
البته با استفاده از دستور Frequencies از طریق Analyze → Descriptive Statistics و کلیک بر روی Statistics می توان برای متغیرهای کمی علاوه بر شاخص های مرکزی، پراکندگی، چولگی و کشیدگی که در پنجره option در دسترس کاربر قرار دارد،

۱- قدر مطلق شاخص چولگی و کشیدگی براساس حد نصاب های زیر به سه دسته تقسیم شده است: کمتر از ۰.۱ ضعیف، بین ۰.۱ تا ۰.۵ متوسط و بیشتر از ۰.۵ شدید.

درس هفتم - توصیف یک متغیری نتایج □ ۱۰۵

شاخص‌های وضعی نظیر چارک‌ها، دهک‌ها، صدک‌ها و نقاط مقطع را نیز محاسبه کرد^۱. تصویر زیر امکانات پنجره‌ی Statistics را در دستور Frequencies نشان می‌دهد. شکل شماره (۴۴) شیوه انتخاب شاخص‌های آماری مناسب برای متغیر روحیه همکاری دانش‌آموزان در

پنجره Statistics



همانطور که ملاحظه می‌شود در این دستور امکان مرتب کردن الفبایی، نزولی یا صعودی متغیرها وجود ندارد.

پس از انجام دستورات فوق، جدول شاخص‌های آماری و نیز نمودار مستطیلی مربوط به متغیر "روحیه همکاری دانش‌آموزان" نیز در گزارش پژوهشی آورده می‌شود^۲.

۱- در صورت استفاده از این دستور برای متغیرهای کمی، محقق می‌تواند به جای آوردن جداول عریض و طویل متغیرهای کمی که عملاً هیچ چیزی را به خواننده منتقل نمی‌کند، از نمودار مستطیلی در کنار جدول شاخصهای آماری استفاده کند. برای یادآوری به درس چهارم "کدگذاری مجدد و دسته‌بندی" و مبحث تبدیل متغیرهای کمی به کیفی مراجعه شود.

۲- لازم به ذکر است که جدول فراوانی به دلیل طولانی بودن حذف شده است

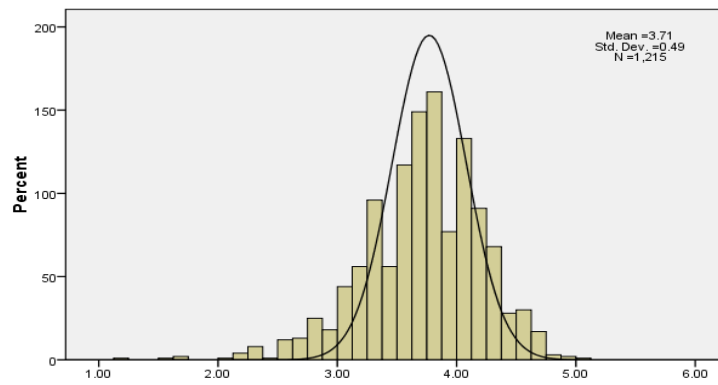
جدول شماره (۲۱) توصیف شاخص‌های آماری "روحیه همکاری دانش‌آموزان"

N	Valid	1215
	Missing	0
	Mean	3.7117
	Std. Error of Mean	.01407
	Median	3.7857
	Std. Deviation	.49027
	Variance	.240
	Skewness	-.644
	Std. Error of Skewness	.070
	Kurtosis	1.272
	Std. Error of Kurtosis	.140
	Range	3.79
	Minimum	1.21
	Maximum	5.00
Percentiles	10	3.0714
	25	3.4286
	30	3.5000
	50	3.7857
	60	3.8571
	75	4.0000
	80	4.0769

در اینجا نیز تفسیر همانند تفسیری است که در صفحات پیشین آمد برای مثال صدک دهم (Percentiles) نشان می‌دهد ده درصد دانش‌آموزان مورد بررسی از لحاظ روحیه همکاری، نمره‌ای کمتر یا مساوی ۳/۰۷ و ۹۰ درصد بیشتر یا مساوی با آن گرفته‌اند. نمودار زیر هم مؤید نتایج مندرج در جدول فوق است.

درس هفتم – توصیف یک متغیری نتایج □ ۱۰۷

شکل شماره (۴۵) شاخص کمی روحیه همکاری دانش‌آموزان



تمرین درس هفتم

براساس داده‌های فایل world95 تمرینات زیر را انجام دهید.

تمرین ۱- درصد شهرنشینی (Urban) در جهان را براساس شاخص‌های مرکزی توصیف و به سؤالات زیر پاسخ دهید:

- ۱- متوسط درصد شهرنشینی در جهان چقدر است؟
- ۲- بالاترین و پایین‌ترین درصد شهرنشینی در جهان چقدر است و متعلق به کدام کشورهاست؟
- ۳- میانه درصد شهرنشینی در جهان چقدر است؟ تفسیر آن را بنویسید.
- ۴- نقطه هفتاد درصد شهرنشینی جهان چند است؟ تفسیر آن را بنویسید.

تمرین ۲- در ارتباط با درصد باسوادان (Literacy) به سؤالات زیر پاسخ دهید.

۱- پراکندگی درصد باسوادان را براساس شاخص‌های واریانس، انحراف معیار و دامنه تغییرات تفسیر کنید.

- ۲- خطای میانگین درصد باسوادان چقدر است؟
- ۳- مقدار چولگی و کشیدگی درصد باسوادان در چه حدی است؟ توصیف کنید.
- ۴- به نظر شما در اکثریت کشورها، درصد باسوادان بیشتر از میانگین است یا کمتر از آن؟ دلایل خود را بنویسید.

تمرین ۳- وضعیت دین (Religion) کشورهای جهان را توصیف و به سؤالات زیر پاسخ دهید.

- ۱- تعداد کشورهای مسلمان چند مورد است؟
- ۲- درصد کشورهای ارتدکس چقدر است؟
- ۳- دین غالب در نمونه مورد بررسی کدام است؟
- ۴- درصد کشورهای غیرمسلمان چقدر است؟

درس هشتم

توصیف دو متغیری نتایج

در پایان این درس انتظار می‌رود دانشجو

- ۱- نحوه تهیه جدول فراوانی برای دو متغیر کیفی را بداند و بتواند آن را تفسیر کند.
- ۲- اهمیت درصدگیری ستونی، سطری یا کل را تشخیص بدهد و از آنها در موارد مناسب استفاده کند.
- ۳- بتواند آماره‌های مرکزی، پراکندگی، وضعی، چولگی و کشیدگی مربوط به متغیرها یا شاخص‌های کمی را برحسب طبقات یک متغیر کیفی، تهیه و تفسیر کند.

۱- مقدمه

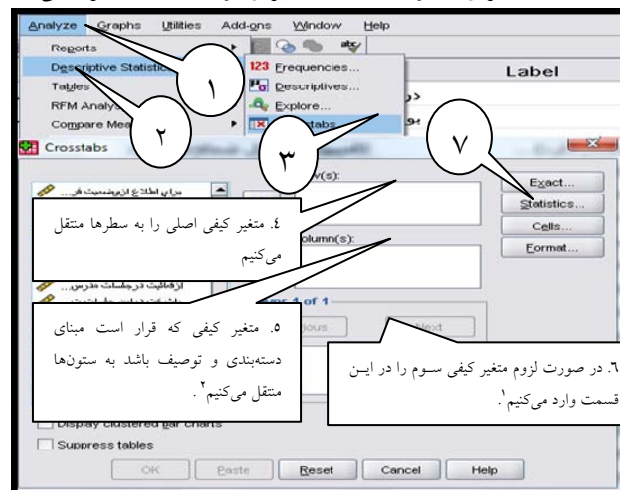
طبیعی است که در ارائه نتایج گزارش پژوهشی، علاوه بر توصیف یک متغیری، به توصیف دو متغیری هم نیاز پیدا می کنید. وقتی قصد داشته باشید یک متغیر را بر حسب متغیر دیگری توصیف کنید، در این حالت می گوییم توصیف از نوع دو متغیری است. توصیف دو متغیره بر حسب سطح سنجش متغیرها به سه دسته تقسیم می شود:

- **توصیف متغیر کیفی بر حسب متغیر کیفی:** مانند کسی که بخواهد درصد مشاغل مختلف را بر حسب جنسیت اعضای نمونه توصیف کند
 - **توصیف متغیر کمی بر حسب متغیر کیفی:** مانند فردی که بخواهد سن اعضای نمونه را به تفکیک وضعیت تأهل نشان دهد.
 - **توصیف متغیر کمی بر حسب متغیر کمی:** مانند کسی که بخواهد درآمد اعضای نمونه را بر حسب سابقه خدمت آنها نشان دهد.
- از آنجا که شکل سوم کمتر رایج است و فقط با نمودار پراکنش انجام می شود، از حیطه ی دروس مبانی spss خارج می شود لذا در اینجا به ارائه ی مباحث مربوط به دو نوع اول می پردازیم.

۲- توصیف یک متغیر کیفی بر حسب یک متغیر کیفی

برای تحقق چنین هدفی می توان از دستورهای مختلفی در نرم افزار spss استفاده کرد اما در اینجا دستور Crosstabs به سبب سهولت و فراوانی کاربرد آن پیشنهاد می شود. به این منظور باید مطابق تصویر زیر عمل کرد.

شکل شماره (۴۶) اجرای دستور crosstabs برای توصیف یک متغیر کیفی - کیفی



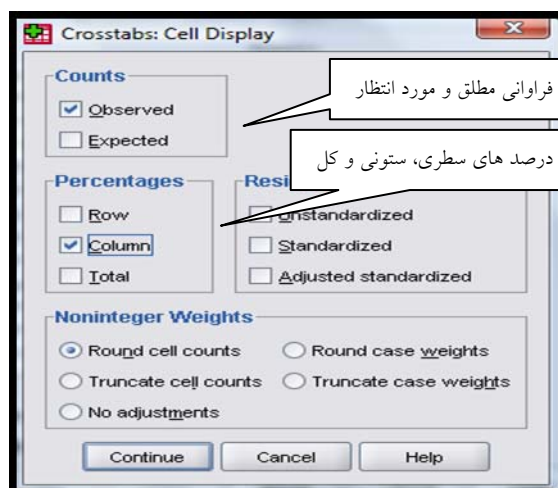
درس هشتم - توصیف دو متغیری نتایج □ ۱۱۱

■ ۱ در این حالت جداول دو بعدی که براساس متغیرهای سطر و ستون فوق‌الذکر تهیه می‌شوند به تفکیک طبقات متغیر کیفی سوم تهیه می‌شوند.

■ ۲ برای نمونه در مثال‌های ارائه شده در این صفحه برای هر یک از دسته‌های توصیف، به ترتیب متغیرهای جنسیت، وضعیت تأهل، سابقه خدمت، متغیرهایی مبنایی برای دسته‌بندی و توصیف هستند.

با کلیک بر روی Cells با پنجره زیر روبرو می‌شوید. برای توصیف نتایج فقط با بخش درصدها (Percentage) سروکار دارید. از این بخش متناسب با هدف توصیفی، درصد سطری یا ستونی یا کل را باید انتخاب کرد. پیشنهاد می‌شود اگر متغیر اصلی در ستون باشد، بهتر است درصد ستونی انتخاب شود. طبیعی است که اگر درصدها ستونی باشند، تفسیر باید به صورت سطری انجام شود. همچنین پیشنهاد دیگر آن است که متغیری را در ستون قرار دهید که دارای مقولات کمتری باشد تا در خروجی، جدول کوچک و شکیل‌تر داشته باشید.

شکل شماره (۴۷) استفاده از پنجره Cells برای نمایش اعداد در جدول خروجی



اگر بخواهید "کیفیت غذای مدرسه را از دیدگاه دانش‌آموزان دختر و پس" را توصیف کنید، برای این منظور متغیر جنسیت را که مبنای درصدگیری است به ستون‌ها و متغیر کیفیت غذا را که هدف توصیف آن است، در سطرها قرار دهید و در بخش درصدها، درصد ستونی را انتخاب کنید. در این حالت با اجرای دستور فوق با جداول زیر روبرو می‌شوید.

جدول شماره (۲۲) ارزیابی دانش آموزان از کیفیت غذای مدرسه بر حسب جنسیت

Case Processing Summary						
بدون پاسخها			کل نمونه مورد بررسی			
پاسخگویان	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
	607	50.0%	608	50.0%	1215	100.0%

جدول فوق نشان می دهد که در توصیف کیفیت غذا بر حسب جنسیت، تنها ۵۰ درصد از دانش آموزان (۶۰۷ نفر) مورد تجزیه و تحلیل قرار گرفته اند و پنجاه درصد باقی مانده به دلیل بی پاسخ گذاشتن سؤال جنسیت یا کیفیت غذا و یا هردو، از تجزیه و تحلیل خارج شده اند و یا به دلایل دیگری سؤال کیفیت غذا از آنها پرسیده نشده است (مثلاً در مدرسی که به دانش آموزان ناهار نمی دهند، این سؤال پرسیده نشده است).

جدول شماره (۲۳) توصیف کیفیت غذای مدرسه بر حسب جنسیت دانش آموزان

Crosstabulation کیفیت غذا * جنس					
		جنس		Total	
		دخترانه	پسرانه		
کیفیت غذا	بسیار نامناسب	Count	25	46	71
	جنس	% within	8.7%	14.4%	11.7%
	نامناسب	Count	51	58	109
	جنس	% within	17.8%	18.1%	18.0%
	بینابین	Count	78	78	156
	جنس	% within	27.2%	24.4%	25.7%
	مناسب	Count	106	105	211
	جنس	% within	36.9%	32.8%	34.8%
	بسیار مناسب	Count	27	33	60
	جنس	% within	9.4%	10.3%	9.9%
	Total	Count	287	320	607
	جنس	% within	100.0%	100.0%	100.0%

این جدول ارزیابی دانش آموزان دختر و پسر را از کیفیت غذای مدرسه نشان می دهد. همانگونه که ملاحظه شود، اکثریت دانش آموزان (۳۴/۸ درصد)، کیفیت غذای مدرسه را مناسب ارزیابی کرده اند که در این میان نسبت دانش آموزان دختر (۳۶/۹ درصد) اندکی بیشتر از

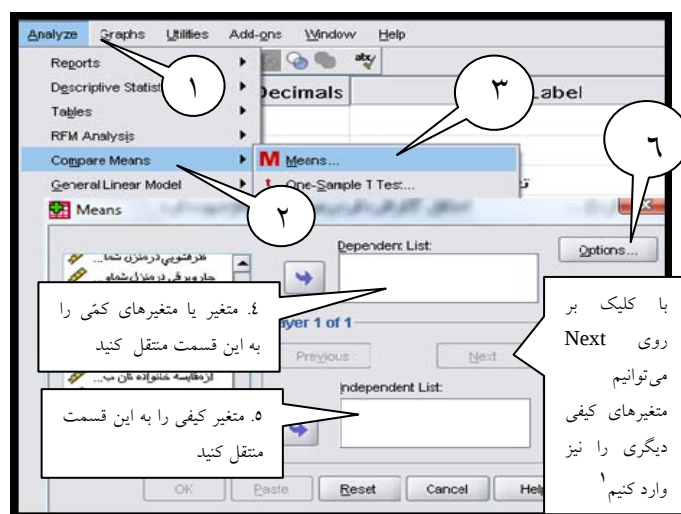
درس هشتم - توصیف دو متغیری نتایج □ ۱۱۳

پسرها (۳۲/۸ درصد) است. از سوی دیگر ۱۱/۷ درصد کیفیت غذا را نامناسب ارزیابی کرده‌اند که در میان آنها نسبت پسران (۱۴/۴) اندکی بیشتر از دختران (۸/۷) است.

۳- توصیف یک متغیر کمی بر حسب یک متغیر کیفی

برای این منظور نیز راه‌حل‌های مختلفی در نرم‌افزار SPSS وجود دارد که پیشنهاد می‌شود از دستور Means در منوی Compare Means استفاده شود.

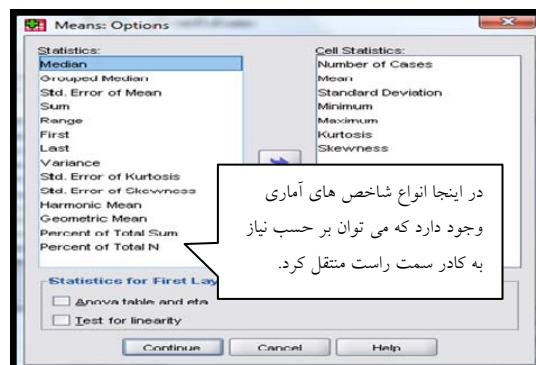
شکل شماره (۴۸) اجرای دستور Means برای توصیف یک متغیر کمی بر حسب متغیر کیفی



■ ۱- در این حالت نتایج به تفکیک طبقات متغیر کیفی ارائه می‌شود.

با کلیک بر روی Option پنجره‌ی زیر باز می‌شود و می‌توان شاخص‌های مرکزی یا پراکندگی را برای توصیف متغیر کمی، در این قسمت انتخاب کرد.

شکل شماره (۴۹) استفاده از پنجره option برای انتخاب شاخص‌های آماری



برای مثال محقق می خواهد بداند "میزان روحیه همکاری در بین دانش آموزان دختر و پسر چقدر است؟" او با اجرای دستور فوق، با خروجی زیر مواجه می شود.

جدول شماره (۲۴) وضعیت روحیه همکاری دانش آموزان بر حسب جنسیت

جنس	N	Minimum	Maximum	Mean	Std. Deviation	Skewness	Kurtosis
دختر	580	1.67	4.60	3.6307	.41940	-.572	.872
پسر	635	1.20	4.73	3.5832	.48845	-.649	1.293
Total	1215	1.20	4.73	3.6059	.45722	-.646	1.267

نتایج بدست آمده بیانگر آن است که:

- ۱- کل نمونه مورد بررسی (۱۲۱۵ دانش آموز) وارد تحلیل شده اند. همانگونه که ملاحظه می شود از این تعداد ۵۸۰ نفر دختر و ۶۳۵ نفر را دانش آموزان پسر تشکیل می دهند.
- ۲- حداقل روحیه همکاری در بین دانش آموزان دختر (۱/۶۷) و حداکثر (۴/۶۰) است. در مورد پسرها نیز حداقل روحیه همکاری ۱/۲۰ بوده و حداکثر آن ۴/۷۳ است.
- ۳- میانگین حاصل نشان می دهد اگرچه روحیه همکاری در بین دانش آموزان دختر و پسر در حد متوسطی است. اما این روحیه در بین دانش آموزان دختر (۳/۶۳) نسبت به پسرها (۳/۵۸) کمی بیشتر است.
- ۴- شاخص انحراف معیار نیز نشان می دهد که روحیه همکاری دانش آموزان دختر به طور متوسط ۰/۴۱ در حول میانگین قرار دارد در حالیکه در مورد پسرها کمی بیشتر (۰/۴۸) است.
- ۵- ضریب چولگی بدست آمده در مورد هر دو گروه دانش آموزان دختر (۰/۵۷-) و پسر (۰/۶۴-) منفی و شدید است. این وضعیت بیانگر این واقعیت است که اکثریت دانش آموزان دختر و پسر روحیه همکاری بالاتر از میانگین فوق الذکر دارند.
- ۶- ضریب کشیدگی نیز در مورد هر دو گروه مثبت و شدید است بدین معنی که اگر چه پراکندگی شاخص کمی روحیه همکاری در هر دو گروه بیشتر از حد نرمال است^۱، اما میزان این پراکندگی در بین پسران (۱/۲۹) بیشتر از دختران (۰/۸۷) است.

۱- مقادیر بالای کشیدگی نشانه پراکندگی کمتر یا شباهت بیشتر بین اعضای نمونه است.

تمرین درس هشتم

براساس فایل (GSS93 SUBSET VALUE) تمرینات زیر را انجام دهید.

تمرین ۱- گرایش اعضای نمونه به موسیقی کلاسیک به تفکیک جنسیت در چه حدی است؟

تمرین ۲- وضعیت سواد (Educ) و درآمد کل خانواده (Income) را در نژادهای مختلف (Race) توصیف کنید.

تمرین ۳- چند درصد مردان با اشتغال زنان در خارج از خانه (fewwork) موافق هستند؟ نسبت زنان مخالف با این موضوع چقدر است؟

تمرین ۴- میزان موافقت با آموزش مسائل جنسی (sexeduc) را بر حسب وضعیت تأهل توصیف کنید. چند درصد زنان مطلقه با این موضوع موافق هستند؟ نسبت مردان مجردی که با این موضوع مخالف هستند، در چه حدی است؟

تمرین ۵- سطح سواد (degree) در بین پایگاه‌های مختلف اجتماعی - اقتصادی (sei) در چه حدی است؟

بخش سوم

آمار استنباطی روابط همبستگی و علی

در پایان این درس انتظار می‌رود دانشجو

۱- بتواند تفاوت سؤالات و فرضیات را براساس سطح سنجش متغیرها و هدف محقق تشخیص دهد.

۲- روش‌های آماری متناسب با سؤالات و فرضیات تحقیق را شناسایی کند.

۱- مقدمه

از آنجا که هدف اغلب تحقیقاتی که به شیوه نمونه‌گیری انجام می‌شوند استنباط ویژگی‌های جامعه آماری و تعمیم نتایج از نمونه به آن می‌باشد، محقق نیاز به استفاده از آمار استنباطی دارد. آمار استنباطی مجموعه‌ای از روش‌های آماری است که محقق به کمک آنها می‌تواند به سؤالات یا فرضیات مربوط به جامعه آماری، از طریق نمونه مورد بررسی پاسخ می‌دهد.

انتخاب یک روش از مجموعه روش‌های آمار استنباطی، به دو معیار اصلی بستگی دارد.

۱- سطح سنجش متغیرهای مندرج در سؤالات یا فرضیات

۲- هدف محقق از بیان سؤال یا فرضیه

تنها پس از شناسایی این دو ملاک عمده است که می‌توان روش آماری مناسبی را برای پاسخ به سؤالات یا آزمون فرضیات تحقیق انتخاب کرد.

در ارتباط با معیار نخست لازم به ذکر است که متغیرها از حیث سطح سنجش به چهار دسته اسمی، رتبه‌ای، فاصله‌ای و نسبی تبدیل می‌شوند که دو مورد اول (اسمی و رتبه‌ای) در شمار متغیرهای کیفی قرار می‌گیرند و دو مورد آخر (فاصله‌ای و نسبی) جزء متغیرهای کمی محسوب می‌شوند. تجربه نشان داده است که روش‌های آماری مربوط به متغیرهای کمی (فاصله‌ای یا نسبی) عملاً یکسان هستند.

در ارتباط با معیار دوم نیز لازم است بدانید که محقق معمولاً در فرضیاتی که طراحی می‌کند در جستجوی کشف نوعی رابطه بین متغیرهاست. این رابطه می‌تواند یکی از حالت‌های همبستگی، علی و یا تفاوت باشد. به فرضیه‌های زیر که در بر گیرنده این سه حالت است، دقت کنید.

۱. میانگین نمره درس آمار در بین دختران بیشتر از پسران است (در اینجا رابطه از نوع تفاوت است).

۲. نوع شغل بر میزان پس‌انداز تأثیر می‌گذارد (در اینجا رابطه از نوع تأثیر است).

۳. بین مدرک تحصیلی و رضایت از شغل همبستگی وجود دارد (در اینجا رابطه از نوع همبستگی است).

جدول زیر روش‌های آماری مناسب را برای پاسخگویی به سؤالات یا فرضیات تحقیق با توجه به دو معیار فوق، به صورت فشرده نشان می‌دهد.

درس نهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی □ ۱۱۹

جدول شماره (۲۵) روش‌های آماری مناسب برای آزمون فرضیات یا سؤالات تحقیق

هدف محقق				
سنجش رابطه از نوع تفاوت	سنجش رابطه از نوع علی	سنجش رابطه از نوع همبستگی		
۱- برای گروه‌های همبسته: الف- ضریب کاپا ب- ضریب مک نمار ۲- برای گروه‌های مستقل: ضریب ریسک	کی دو + ضریب لامدا	کی دو + ضرایب فی و وی کرامر	اسمی-اسمی	روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی
۱- H کراسکال والیس برای گروه‌های مستقل ۲- آزمون فریدمن (برای گروه‌های همبسته)	d سامرز	۱- نای بی کندال (برای جداول مربع) ۲- نای سی کندال (برای جداول مستطیل) ۳- ضریب همبستگی اسپیرمن	رتبه‌ای - رتبه‌ای	
۱- در گروه‌های همبسته الف- برای مقایسه دو متغیر کمی از T-Test Paired ب- برای بیش از دو متغیر کمی تحلیل واریانس چند متغیری (MANOVA) ۲- در گروه‌های مستقل (ابتدا یکی از متغیرهای کمی را گروه‌بندی می‌کنیم و سپس): الف- T-Test برای دو گروه مستقل ب- تحلیل واریانس برای چند گروه مستقل	رگرسیون چندگانه	۱- ضریب همبستگی پیرسون (اگر توزیع هر دو متغیر کمی نرمال باشد). ۲- ضریب همبستگی اسپیرمن (اگر توزیع لافیل یکی از متغیرهای کمی غیر نرمال باشد).	کمی ^۱ - کمی ^۲	
۱- U مان ویتنی برای دو گروه مستقل ۲- H کراسکال والیس برای چند گروه مستقل	کمی دو + ضریب لامدا	کی دو + وی کرامر	اسمی - رتبه‌ای	
در صورت نرمال بودن توزیع متغیر کمی: ۱- T-Test برای دو گروه مستقل ۲- تحلیل واریانس برای چند گروه مستقل	مجه‌ذور ضریب اتا	ضریب اتا	اسمی - کمی	
در صورت نرمال بودن توزیع متغیر کمی: ۱- T-Test برای دو گروه مستقل ۲- تحلیل واریانس برای چند گروه مستقل	مجه‌ذور ضریب اتا	ضریب همبستگی اسپیرمن	رتبه‌ای - کمی	

- از آنجا که سؤالات و فرضیات غالباً دو متغیری هستند، غالباً ترکیبات مختلفی از حیث سطوح سنجش پدید می‌آید.
- چون روش‌های آماری مربوط به سطوح فاصله‌ای و نسبی عمدتاً یکسان هستند، به جای سطح سنجش فاصله‌ای و نسبی، مسامحتاً از اصطلاح کمی استفاده شده است.
- اگر توزیع متغیر کمی، نرمال نباشد، باید از U مان ویتنی (برای مقایسه در دو گروه) و نیز H کراسکال والیس (برای مقایسه بیش از دو گروه) استفاده کرد.

تمرین درس مقدمات روابط همبستگی و علی

برای آزمون هر یک از فرضیات زیر کدام روش آماری مناسب است؟ دلایل خود را توضیح دهید.

- الف- مدت ورزش زنان و مردان تفاوت معناداری با یکدیگر دارد.
- ب- مدت زمانی که شهروندان به ورزش پیاده‌روی اختصاص می‌دهند به طرز معناداری بیشتر از مدت زمانی است که به فوتبال اختصاص می‌دهند.
- ج- نوع شغل شهروندان بر مدت مشارکت آنان در ورزش همگانی تأثیر می‌گذارد.
- د- نوع شغل شهروندان بر نوع ورزش همگانی مورد علاقه آنان تأثیر می‌گذارد.
- ه- بین علاقه شهروندان به ورزش همگانی و رضایت آنان از زندگی همبستگی معناداری وجود دارد.
- و- هر چقدر پابندی شهروندان به عبادات بیشتر باشد، علاقه آنان به ورزش همگانی بیشتر خواهد شد.
- ز- سن و تحصیلات شهروندان بر مدت مشارکت آنان در ورزش همگانی تأثیر می‌گذارد.
- ح- علاقه به ورزش در بین شهروندان زن و مرد تفاوت معناداری با یکدیگر دارد.
- ط- رضایت از زندگی در بین ورزشکاران رشته‌های مختلف تفاوت معناداری با یکدیگر دارد.

درس نهم

روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی

در پایان این درس انتظار می‌رود دانشجو

- ۱- بتواند موارد کاربرد آزمون کی دو را شناسایی کند و مفروضات آن را بداند.
- ۲- توانایی تفسیر جداول دو متغیری را داشته باشد.
- ۳- بتواند نتایج آزمون کی دو را تفسیر کند.
- ۴- بتواند موارد کاربرد ضرایب همبستگی و علی متناسب با متغیرهای کیفی اسمی را شناسایی کند.
- ۵- بتواند شدت ضرایب همبستگی و علی را تفسیر کند.

۱- مقدمه

همانطور که در مقدمه این بخش ذکر شد، روابط بین متغیرهای اسمی را می‌توان در سه قالب همبستگی، علی و تفاوت تنظیم کرد.

جدول شماره (۲۶) روش‌های آماری مناسب برای آزمون فرضیات یا سؤالات تحقیق دارای

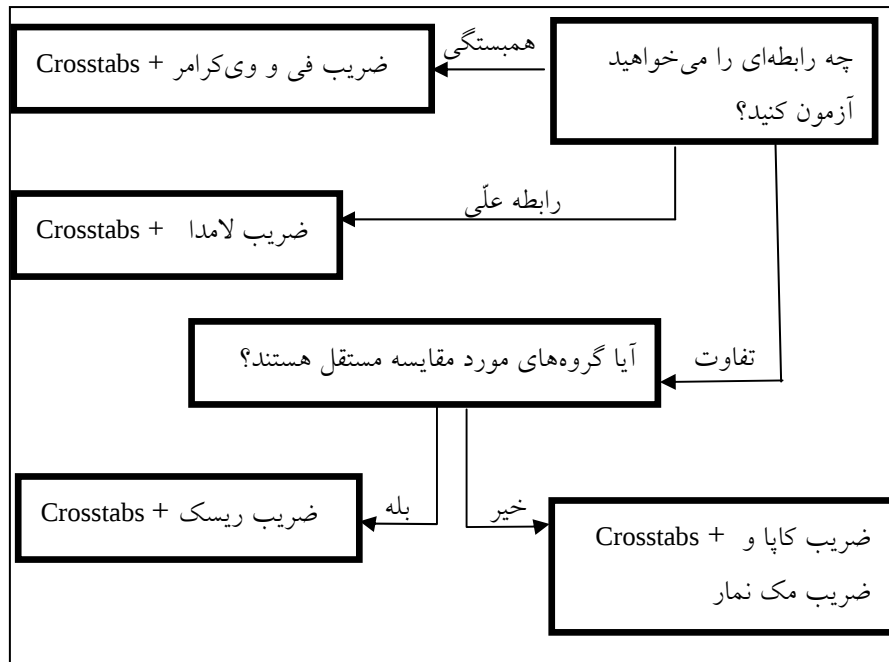
متغیرهای اسمی

سنجش رابطه از نوع تفاوت	سنجش رابطه از نوع علی	سنجش رابطه از نوع همبستگی	
۱- برای گروه‌های همبسته: الف- ضریب کاپا ب- ضریب مک نمار ۲- برای گروه‌های مستقل: ضریب ریسک	کی دو + ضریب لامبدا	کی دو + ضرایب فی و وی کرامر	اسمی - اسمی
۱- نسبت موافقین کاندیدای الف قبل و بعد از سخنرانی او تفاوت معناداری پیدا خواهد کرد. ۲- نسبت طرفداران کاندیدای الف بین زنان و مردان تفاوت معناداری با یکدیگر دارد.	رشته تحصیلی دانش‌آموزان بر ترجیحات شغلی آنان تأثیر می‌گذارد.	بین وضعیت اشتغال و تأهل جوانان همبستگی معناداری وجود دارد.	مثال

دیاگرام بعد تصویر روشنی را از نحوه تصمیم‌گیری در هنگام آزمون فرضیاتی که متغیرهایشان در سطح اسمی قرار دارد، در اختیار می‌گذارد.

درس نهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی □ ۱۲۳

شکل شماره (۵۰) نحوه تصمیم‌گیری درباره روش آماری آزمون فرضیات دارای متغیرهای سطح اسمی



اما نظر به کاربرد وسیع روابط همبستگی و علی در پژوهش‌های فرهنگی و اجتماعی، روابط بین متغیرهای اسمی را تنها به دو دسته همبستگی و علی محدود ساخته و مبحث تفاوت‌ها را به کتاب‌های مرتبط^۱ ارجاع می‌دهیم.

۲- آزمون و ضریب مناسب برای سنجش رابطه علی بین متغیرهای اسمی

• آزمون کی دو و ضریب تأثیر لامدا

فرض کنید محقق در نظر دارد این فرضیه را مورد سنجش قرار دهد که:

۱- برای مطالعه بیشتر در این زمینه به کتاب‌های نگهبان (۱۳۸۲)، حسینی (۱۳۸۲) و حبیب‌پور و صفری (۱۳۸۸) مراجعه کنید.

۱۲۴ □ راهنمای کاربردی نرم افزار SPSS

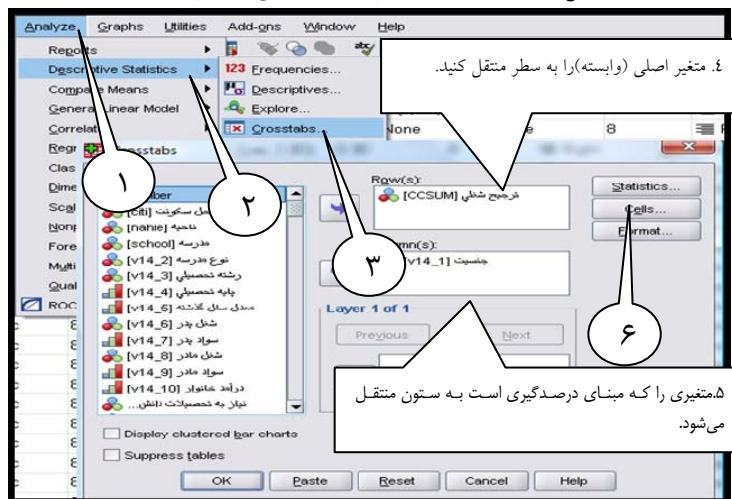
"جنسیت بر نوع ترجیحات شغلی دانش آموزان تأثیر می گذارد."^۱ او باید به دلایل زیر از آزمون کی دو^۲ و ضریب لامدا استفاده کند.

الف) سطح سنجش: در این فرضیه دو متغیر "جنسیت" و "نوع ترجیحات شغلی" وجود دارد که هر دو در سطح سنجش اسمی هستند.

ب) هدف محقق: در این فرضیه محقق قصد سنجش تأثیر یک متغیر بر متغیر دیگر را دارد. لذا رابطه موردنظر از نوع علی است.

در چنین شرایطی باید دستور Crosstabs را از منوی Analyze اجرا کرد.

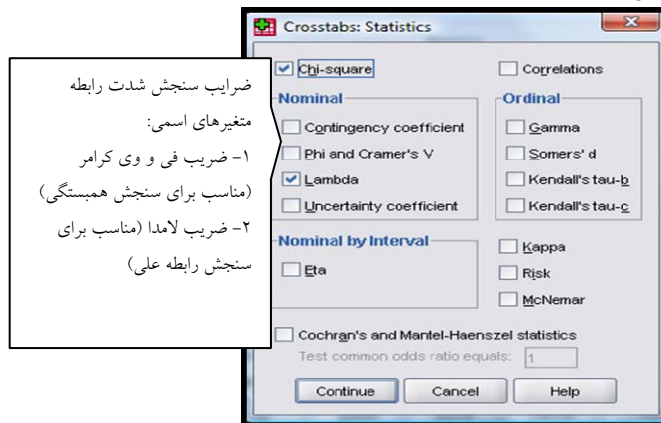
شکل شماره (۵۱) شیوه اجرای آزمون کی دو (مرحله ۱)



- ۱- و یا قصد داشته باشد به این سؤال «آیا جنسیت بر ترجیحات شغلی دانش آموزان تأثیر می گذارد؟» پاسخ دهد.
- ۲- مفروضات آزمون کی دو: استفاده از آزمون کی دو تحت تأثیر محدودیت های خاصی است. این فرمول کلی را فقط در مواردی می توان استفاده کرد که سه شرط زیر تحقق یابد:
 - یکی از متغیرها بیش از دو مقوله داشته باشد.
 - کمتر از ۲۰ درصد خانه ها، فراوانی مورد انتظار کمتر از ۵ داشته باشند
 - هیچ خانه ای فراوانی مورد انتظار کمتر از ۱ نداشته باشد. (کرامر، ۱۳۸۶: ۷۹) در غیر اینصورت باید از تصحیح یتس یا آزمون دقیق فیشر (Fisher Exact Test) (برای جداول ۲در۲) استفاده کرد.

درس نهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی □ ۱۲۵

همانگونه که در شکل بالا مشاهده می‌شود، پس از انتقال متغیرهای اسمی به سطرها و ستون‌های مربوطه^۱ با کلیک بر روی گزینه Statistics، با تصویر زیر مواجه می‌شویم. شکل شماره (۵۲) شیوه اجرای آزمون کی دو (مرحله ۲)



همانطور که ملاحظه می‌شود چون هدف محقق سنجش رابطه علی است، علاوه بر انتخاب آزمون کی دو، از میان ضرایب مختلف تصویر فوق، فقط ضریب لامدا انتخاب شده است. بدیهی است اگر فرضیه از نوع همبستگی می‌بود، باید به جای لامدا، ضرایب فی و وی کرامر انتخاب می‌شد.^۲

در مرحله بعد پس از انتخاب ضریب مناسب بر روی Continue کلیک و گزینه Cells را انتخاب کرد و بعد از آن، به شرح تصویر زیر عمل می‌شود.

۱- پیشنهاد می‌شود در روابط همبستگی، متغیرهای اسمی کم مقوله و در روابط علی، متغیرهای مستقل در قسمت ستون قرار داده شود.

۲- ضریب فی مخصوص جداول ۲ در ۲ است اما وی کرامر محدودیت ندارد.

شکل شماره (۵۳) شیوه اجرای آزمون کی دو (مرحله ۳)



همانگونه که مشاهده می‌شود، در مرحله اول آزمون کی دو، متغیر جنسیت به ستون منتقل شد تا ترجیحات شغلی برحسب جنسیت توصیف شود، به همین دلیل در اینجا فقط درصد ستونی^۱ را باید انتخاب نمود، سپس گزینه Continue را کلیک کرد. حال به تفسیر خروجی های این دستور توجه کنید.

جدول شماره (۲۷) خلاصه پردازش اطلاعات

	Valid		Missing		Total	
	Cases		Cases		Cases	
	N	Percent	N	Percent	N	Percent
ترجیح شغلی * جنسیت	1442	95.8%	64	4.2%	1506	100.0%

این جدول نشان می‌دهد که از مجموع ۱۵۰۶ نفر نمونه مورد بررسی ۹۵/۸ درصد (۱۴۴۲ نفر) به سؤالات مربوط به جنسیت و ترجیح شغلی پاسخ داده و در محاسبات آزمون کی دو مورد تجزیه و تحلیل قرار گرفته‌اند. تنها ۴/۲ درصد (۶۴ نفر) هستند که حداقل به یکی از این سؤالات پاسخی نداده‌اند و از تحلیل حاضر خارج شده‌اند.

۱- در ارتباط با نحوه درصدگیری پیشنهاد می‌شود در روابط علی از درصدگیری براساس متغیر مستقل استفاده شود اما در روابط همبستگی می‌شود براساس یکی از دو متغیر یا هر دو و یا درصدگیری کل عمل کرد.

درس نهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی □ ۱۲۷

جدول شماره (۲۸) جدول توافقی بین دو متغیر ترجیح شغلی و جنسیت

Crosstabulation ترجیح شغلی * جنسیت				
		جنسیت		Total
		پسر	دختر	
ترجیح شغلی	فنی	Count 221 % within جنسیت 30.9%	Count 112 % within جنسیت 15.4%	333 23.1%
	علمی	Count 188 % within جنسیت 26.3%	Count 199 % within جنسیت 27.4%	387 26.8%
	هنری	Count 99 % within جنسیت 13.8%	Count 215 % within جنسیت 29.6%	314 21.8%
	آموزشی	Count 80 % within جنسیت 11.2%	Count 145 % within جنسیت 19.9%	225 15.6%
	مدیریتی	Count 77 % within جنسیت 10.8%	Count 26 % within جنسیت 3.6%	103 7.1%
	دفتری	Count 50 % within جنسیت 7.0%	Count 30 % within جنسیت 4.1%	80 5.5%
Total		Count 715 % within جنسیت 100.0%	Count 727 % within جنسیت 100.0%	1442 100.0%

مقایسه دختران و پسران از لحاظ علایق شغلی بیانگر تفاوت‌های قابل توجهی است. به گونه‌ای که اولویت اول تا چهارم ترجیحات شغلی دختران به ترتیب مشاغل هنری (۲۹/۶ درصد)، مشاغل علمی (۲۷/۴ درصد)، مشاغل آموزشی (۱۹/۹ درصد) و مشاغل فنی (۱۵/۴ درصد) می‌باشد در حالیکه ترجیحات شغلی پسران تقریباً عکس ترجیحات شغلی دختران است بدین معنی که آنها در اولویت نخست مشاغل فنی (۳۰/۹ درصد) سپس مشاغل علمی (۲۶/۳ درصد) و در مرحله سوم مشاغل هنری (۱۳/۸ درصد) را ترجیح می‌دهند. همانطور که

ملاحظه می‌شود وجه شباهت ترجیحات شغلی دختران و پسران در عدم ترجیح مشاغل دفتری و مدیریتی بر سایر مشاغل است.

جدول شماره (۲۹) تعیین سطح معناداری^۱ متغیرهای مورد سنجش

Chi-Square Tests			
	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	127.784 ^a	5	.000
Likelihood Ratio	130.946	5	.000
Linear-by-Linear Association	4.303	1	.038
N of Valid Cases	1442		

a. 0 cells (.0%) have expected count less than 5. The minimum expected count is 39.67.

این جدول نتیجه آزمون کی دو را نشان می‌دهد همانطور که ملاحظه می‌شود با توجه به میزان خطا ($\text{sig}=0.000$) می‌توان با ۰/۹۵ اطمینان، ادعا کرد که بین جنسیت و ترجیح شغلی رابطه معناداری وجود دارد ($\text{Pearson Chi-Square}=127.784, \text{df}=5, \text{Sig.}=.000$)

از نتایج جدول فوق نباید فوراً چنین استنباط شود که لزوماً تأثیر جنسیت بر ترجیح شغلی است که معنادار شده است، زیرا ممکن است حالت‌های دیگر رابطه بین این دو متغیر (مثلاً تأثیر ترجیح شغلی بر جنسیت یا همبستگی بین این دو متغیر) معنادار شده باشد^۲. لذا محقق باید به سایر نتایج مربوطه نیز توجه کند و در یک جمع‌بندی کلی نسبت با تأیید یا رد فرضیه خود اقدام کند.

از سوی دیگر باید توجه داشت که منظور از معناداری یک رابطه (خواه همبستگی، خواه علی و خواه تفاوت) آن است که نتایج حاصله از نمونه مورد بررسی، به جامعه آماری قابل تعمیم است. لذا معنادار بودن روابط بین دو متغیر، مترادف با قوی یا ضعیف بودن رابطه بین آنها نیست. بلکه برای این منظور باید ضرایب مربوط به هر رابطه را مورد بررسی قرار داد. لذا چنانچه نتیجه آزمون کی دوی مربوط به همبستگی دو متغیر اسمی، همانند مثال فوق-

۱- حداکثر خطای مجاز در تحقیقات علوم انسانی ۵ درصد است که با نماد Sig در خروجی‌های این نرم‌افزار نشان داده می‌شود.

۲- نرم‌افزار spss حالت‌های سه‌گانه رابطه را (اعم از تأثیر علی، همبستگی، تفاوت) در نظر می‌گیرد. این وظیفه محقق است که باید با توجه به فرضیه یا سؤال تحقیق خود یکی از آنها را ملاک تفسیر قرار دهد.

درس نهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی □ ۱۲۹

معنادار شد، برای تعیین شدت میزان دو متغیر، باید خروجی‌های ذیل نیز تفسیر شود. بدیهی است که در صورت معنادار نبودن رابطه، نیازی به تفسیر جدول ذیل نیست.

جدول شماره (۳۰) تعیین شدت^۱ تأثیر جنسیت بر ترجیحات شغلی

اینجا شدت رابطه علی را در وضعیت‌های مختلف نشان می‌دهد		مقدار خطای استاندارد	مقدار آماره t	میزان خطا	
		Value	Asymp. Std. Error ^a	Approx. T ^b	Approx. Sig.
Nominal by Lambda Nominal	Symmetric	.129	.022	5.483	.000
	Dependent ترجیح شغلی	.046	.027	1.710	.087
	Dependent جنسیت	.252	.027	8.102	.000
Goodman and Kruskal tau	Dependent ترجیح شغلی	.020	.004		.000 ^c
	Dependent جنسیت	.089	.015		.000 ^c

a. Not assuming the null hypothesis.
b. Using the asymptotic standard error assuming the null hypothesis.
c. Based on chi-square approximation

این جدول نتایج ضریب لامدای مربوط به رابطه جنسیت و ترجیح شغلی را در سه حالت نشان می‌دهد:

الف- حالتی که همبستگی این دو متغیر مدّ نظر باشد (Lambda Symmetric)

ب- حالتی که ترجیح شغلی، متغیر وابسته باشد

ج- حالتی که جنسیت، متغیر وابسته باشد.

نتایج ضریب لامدا در حالت (ب) تأثیر جنسیت را بر ترجیح شغلی دانش‌آموزان نشان می‌دهد. همانطور که ملاحظه می‌شود شدت تأثیر جنسیت بر ترجیح شغلی بسیار ضعیف بوده و در عین حال معنادار نیست (Lambda=0.046, Sig=0.087). لذا فرضیه محقق دال بر تأثیر جنسیت بر ترجیح شغلی رد می‌شود به عبارتی دقیق‌تر نتایج جدول قبل در زمینه تفاوت ترجیحات شغلی دختران و پسران قابلیت تعمیم به جامعه آماری را ندارد.^۲

۱- شدت رابطه براساس قدر مطلق ضرایب به شکل زیر دسته‌بندی و تفسیر می‌شود: ۰ تا ۰.۲ خیلی ضعیف، ۰.۲ تا ۰.۴ ضعیف، ۰.۴ تا ۰.۶ متوسط، ۰.۶ تا ۰.۸ قوی و ۰.۸ تا ۱ بسیار قوی.

۲- خطای مجاز در آزمون فرضیات حداکثر ۵ درصد است.

۱- شاید در اینجا این سؤال پیش بیاید اگر تأثیر جنسیت بر علایق شغلی در این جدول معنادار نیست پس چرا در جدول قبلی، آزمون کی دو حاکی از وجود رابطه معناداری بین جنسیت و علایق شغلی بود. در پاسخ باید به میزان خطای همبستگی این دو متغیر و نیز تأثیر علایق شغلی بر جنسیت اشاره کرد که هر دو معنادار هستند. به

لازم به ذکر است که در جدول فوق علاوه بر ضریب لامدا، ضریب گودمن کراسکال نیز بر حسب اینکه متغیر وابسته ترجیح شغلی یا جنسیت باشد، در خروجی وجود دارد.

۳- آزمون و ضریب مناسب برای سنجش رابطه همبستگی بین متغیرهای اسمی

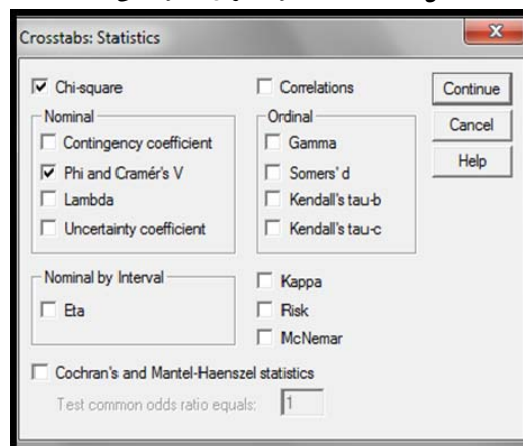
• آزمون کی دو و ضرایب همبستگی فی و وی کرامر

حال اگر محققى بخواهد این فرضیه را بسنجد که "بین ترجیح شغلی و رشته تحصیلی دانش‌آموزان همبستگی وجود دارد، باید به دلایل زیر از آزمون کی دو و ضریب فی و وی کرامر استفاده کند.

الف) سطح سنجش: در این فرضیه دو متغیر ترجیح شغلی و رشته تحصیلی هر دو سطح سنجش اسمی هستند.

ب) هدف محقق: در این فرضیه محقق قصد سنجش رابطه همبستگی بین دو متغیر ترجیح شغلی و رشته تحصیلی را دارد.

شکل شماره (۵۴) شیوه اجرای آزمون کی دو



عبارتی دیگر معناداری کی دو در جدول قبل ناظر به همین دو حالت بود نه حالتی که محقق در فرضیه خودش قصد بررسی آن را داشت.

درس نهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی □ ۱۳۱

جدول شماره (۳۱) جدول توافقی بین دو متغیر ترجیح شغلی و رشته تحصیلی

Crosstabulation * رشته تحصیلی

within % رشته تحصیلی

		رشته تحصیلی					Total
		ریاضی	تجربی	انسانی	کار و دانش	فنی و حرفه‌ای	
فنی	ترجیح	37.8%	11.9%	15.7%	40.0%	50.0%	24.1%
علمی	شغلی	26.8%	48.6%	8.3%	17.1%	16.7%	28.2%
هنری		12.4%	22.6%	25.5%	20.0%	16.7%	19.4%
آموزشی		10.7%	8.2%	30.4%	11.4%		15.0%
مدیریتی		5.5%	4.1%	12.7%	5.7%		6.9%
دفتری		6.9%	4.5%	7.4%	5.7%	16.7%	6.3%
Total		100.0%	100.0%	100.0%	100.0%	100.0%	100.0%

مقایسه ترجیحات شغلی دانش‌آموزان در رشته‌های تحصیلی مختلف بیانگر تفاوت قابل توجهی است. بدین معنا که ترجیح شغلی اکثریت قابل‌توجهی از دانش‌آموزان رشته‌های فنی حرفه‌ای و کار و دانش در حوزه فنی قرار دارد اما ترجیح شغلی دانش‌آموزان رشته انسانی در دو حوزه آموزشی و هنری بیشتر است. از سوی دیگر غالب دانش‌آموزان رشته تجربی گرایش به سمت مشاغل علمی دارند اما در رشته ریاضی، حوزه‌های اصلی علایق شغلی دانش‌آموزان فنی و علمی است.

جدول شماره (۳۲) تعیین سطح معناداری متغیرهای مورد سنجش

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	188.991 ^a	20	.000
Likelihood Ratio	188.845	20	.000
Linear-by-Linear Association	27.207	1	.000
N of Valid Cases	779		

a. 8 cells (26.7%) have expected count less than 5. The minimum expected count is .38.

این جدول نتیجه آزمون کی دو را نشان می‌دهد. همانطور که ملاحظه می‌شود با توجه به میزان خطا $\alpha = 0.000$ (sig=) می‌توان با 0.95 اطمینان، ادعا کرد که بین ترجیح شغلی و رشته

۱۳۲ □ راهنمای کاربردی نرم‌افزار SPSS

تحصیلی رابطه همبستگی معناداری وجود دارد. ($\text{sig} = ۰/۰۰۰$) و $\text{pearson chi-square} = ۱۸۸/۹۹۱$

جدول شماره (۳۳) تعیین جهت و شدت تأثیر جنسیت بر ترجیحات شغلی

Symmetric Measures		
	Value	Approx. Sig.
Nominal by Nominal Phi	.493	.000
Cramer's V	.246	.000
N of Valid Cases	779	

این جدول ضرایب فی و وی کرامر را در مورد شدت رابطه همبستگی بین دو متغیر ترجیح شغلی و رشته تحصیلی نشان می‌دهد. همانطور که ملاحظه می‌شود شدت همبستگی در حد ضعیف اما معنادار است ($\text{cramers V} = ۰/۲۴۶$)^۱.

۱- ضریب فی برای جداول ۲ در ۲ و ضریب وی کرامر برای کلیه جداول (خواه ۲ در ۲ و خواه بزرگتر) استفاده می‌شود. لذا برای سادگی انجام کار بهتر است همواره به تفسیر وی کرامر بپردازید.

درس نهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای اسمی □ ۱۳۳

تمرینات درس نهم

۱- براساس نتایج مندرج در جداول زیر که رابطه بین وضعیت تأهل و انجام ورزش را نشان می‌دهند، به سؤالات بعد پاسخ دهید.

ورزش * وضعیت تأهل Crosstabulation

	وضعیت تأهل		
	مجرد	متأهل	Total
ورزش خیر Count	201	399	600
% within وضعیت تأهل	41.6%	68.6%	56.3%
ورزش بله Count	282	183	465
% within وضعیت تأهل	58.4%	31.4%	43.7%
Total Count	483	582	1065
% within وضعیت تأهل	100.0%	100.0%	100.0%

Chi-Square Tests				
	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)
Pearson Chi-Square	77.888 ^a	1	.000	
Continuity Correction ^b	76.796	1	.000	
Likelihood Ratio	78.604	1	.000	
Fisher's Exact Test				.000
Linear-by-Linear Association	77.815	1	.000	
N of Valid Cases	1065			
a. 0 cells (.0%) have expected count less than 5. The minimum expected count is 210.89.				
b. Computed only for a 2x2 table				

Directional Measures

			Value	Asymp. Std. Error ^a	Approx. T ^b	Approx. Sig.
Nominal by Nominal	Lambda	Symmetric	.190	.038	4.676	.000
		Dependent ورزش	.174	.043	3.709	.000
		وضعیت تاهل Dependent	.205	.040	4.637	.000
	Goodman and Kruskal tau	Dependent ورزش	.073	.016		.000 ^c
		وضعیت تاهل Dependent	.073	.016		.000 ^c

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

c. Based on chi-square approximation

الف- چند درصد از کل نمونه ورزش می‌کنند؟ نسبت ورزشکاران مجرد بیشتر است یا ورزشکاران متأهل؟ چقدر؟

ب- آیا رابطه بین وضعیت تاهل و انجام ورزش معنادار است؟ شدت رابطه مزبور چقدر است؟

۲- به توجه به فایل satisfy.sav آیا جنسیت (gender) مشتری بر شیوه پرداخت (payment) تأثیر معناداری دارد؟
نتایج مربوطه را تفسیر کنید

۳- به توجه به فایل satisfy.sav آیا بین شیوه پرداخت (payment) و انگیزه خرید (reason1) همبستگی وجود دارد؟
نتایج مربوطه را تفسیر کنید

درس دهم

روش‌های آماری مناسب برای سنجش رابطه متغیرهای رتبه‌ای

در پایان این درس انتظار می‌رود دانشجو

- ۱- بتواند موارد کاربرد ضرایب همبستگی و علی رتبه‌ای را شناسایی کند.
- ۲- توانایی تفسیر جداول دو متغیری رتبه‌ای را داشته باشد.
- ۳- بتواند جهت و شدت ضرایب همبستگی و علی رتبه‌ای را تفسیر کند.

۱- مقدمه

همانطور که در مقدمه این بخش توضیح داده شد، متغیرهای رتبه‌ای براساس هدفی که محقق از طراحی فرضیه یا سؤال دارد، می‌توانند در سه وضعیت متفاوت قرار بگیرند. جدول زیر وضعیت‌های مذکور را همراه با ضرایب یا آزمون‌های مناسب و مثال‌های مرتبط نشان می‌دهد.^۱

جدول شماره (۳۴) روش‌های آماری مناسب برای آزمون فرضیات یا سؤالات تحقیق دارای متغیرهای رتبه‌ای

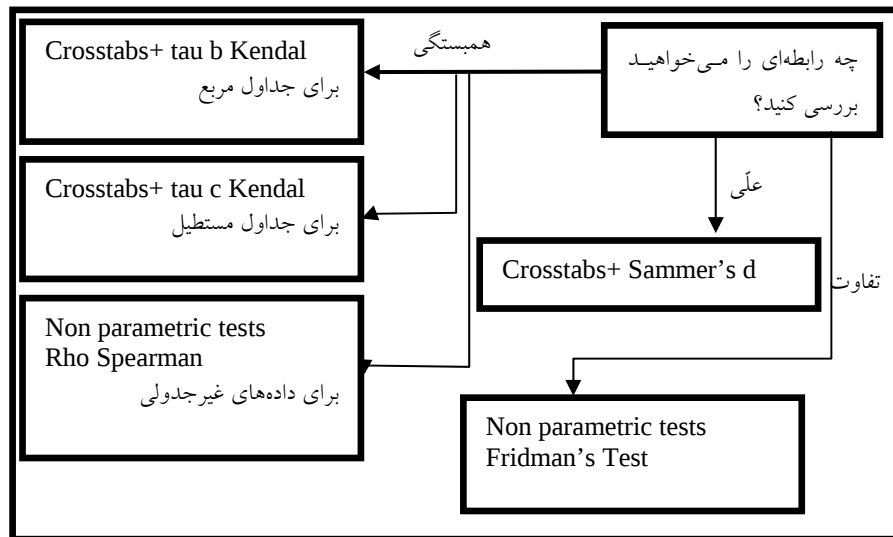
سطح سنجش	ضرایب مناسب برای سنجش شدت همبستگی	ضرایب مناسب برای سنجش شدت رابطه علی	آزمون‌های مناسب برای سنجش معناداری تفاوت
رتبه‌ای-رتبه‌ای	۲- تای بی کندال (برای جداول مربع) ۲-تای سی کندال (برای جداول مستطیل) ۳- ضریب همبستگی اسپیرمن	d سامرز	۲- H کراسکال والیس برای گروه‌های مستقل ۲- آزمون فریدمن (برای گروه‌های همبسته)
مثال	بین رضایت از شغل و رضایت از زندگی همبستگی مستقیم و معناداری وجود دارد.	تحصیلات افراد تأثیر مستقیم بر نیاز آنان به پیشرفت می‌گذارد.	۱- رغبت به تحصیل رشته‌های مختلف تحصیلی تفاوت معناداری با یکدیگر دارد. ۲- بین رضایت افراد از ارتقاء، درآمد و همکاران تفاوت معناداری وجود دارد.

برای تصمیم‌گیری در زمینه آزمون فرضیاتی که هر دو متغیر آنها در سطح رتبه‌ای قرار دارد می‌توان طبق شکل زیر عمل کرد:

۱- لازم به ذکر است که اگر چه ضریب همبستگی اسپیرمن برای متغیرهای رتبه‌ای است اما از این ضریب وقتی استفاده می‌شود که تنوع رتبه‌ها در متغیرهای مورد بررسی بالا و تکرار آن پایین باشد. ضریب همبستگی اسپیرمن به دلیل شباهت ساختاری در خروجی و تفسیر، همراه با ضریب همبستگی پیرسون در درس بعد توضیح داده خواهد شد.

درس دهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای رتبه‌ای □ ۱۳۷

شکل شماره (۵۵) راهنمای انتخاب آزمون برای رابطه متغیرهای کیفی



۲- ضرایب مناسب برای سنجش رابطه علّی بین متغیرهای رتبه‌ای

• ضریب تأثیر d سامرز

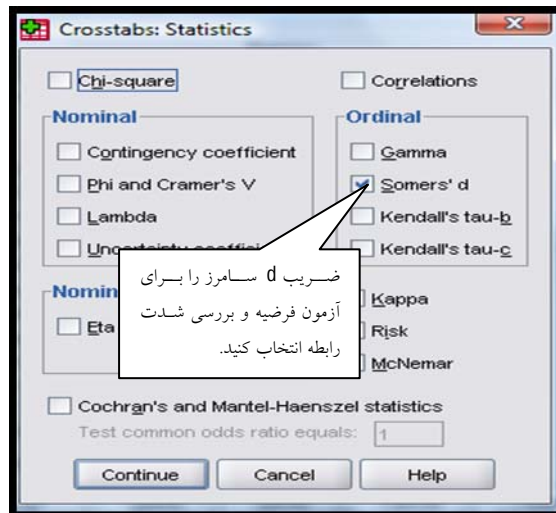
اگر قصد داشته باشید این فرضیه را به محک آزمون بگذارید که "وضعیت اقتصادی خانواده بر تقید به روزه‌داری دانش‌آموزان تأثیر دارد" به دلایل زیر باید از آزمون d سامرز استفاده کنید.

الف) سطح سنجش: در این فرضیه دو متغیر "وضعیت اقتصادی" و "تقید به روزه‌داری" وجود دارد که هر دو در سطح سنجش رتبه‌ای هستند.

ب) هدف محقق: در این فرضیه، یک متغیر بر متغیر دیگر تأثیر دارد. لذا رابطه از نوع علّی است.

برای این منظور همانند متغیرهای اسمی در درس قبل، دستور **Crosstabs** را از منوی **Analyze** انتخاب کرده و پس از انتقال متغیرهای اسمی به سطرها و ستون‌های مربوطه بر روی **Statistics** کلیک کرده و به شرح تصویر زیر عمل کنید.

شکل شماره (۵۶) انتخاب آزمون‌ها و ضرایب مناسب



چنانچه بخواهید به بررسی رابطه همبستگی متغیرهای رتبه‌ای پردازید، باید به جای d سامرز، تایی^۱ یا تایی سی^۲ کندال را انتخاب کنید. نتایج آزمون فوق، در جداول بعد منعکس شده است.

جدول شماره (۳۵) خلاصه پردازش اطلاعات

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
تقید به روزه داری *شاخص کیفی وضعیت اقتصادی	1206	99.3%	9	.7%	1215	100.0%

۱- برای جداول مربع (جداولی که تعداد سطر و ستون‌های آن با هم برابرند).

۲- برای جداول مستطیل (جداولی که تعداد سطر و ستون‌های آن با هم برابر نیستند).

درس دهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای رتبه‌ای □ ۱۳۹

این جدول نشان می‌دهد که از مجموع ۱۲۱۵ دانش‌آموز در نمونه مورد بررسی، ۹۹/۳ درصد (۱۲۰۶ نفر) به سؤالات مربوط به تقید به روزه‌داری و وضعیت اقتصادی پاسخ داده و مورد تجزیه و تحلیل قرار گرفته‌اند و تنها در حدود ۰/۷ درصد (۹ نفر) از کل نمونه مورد بررسی به دلیل عدم پاسخگویی به حداقل یکی از سؤالات، از تحلیل حاضر خارج شده‌اند. جدول شماره (۳۶) جدول توافقی بین دو متغیر وضعیت اقتصادی خانواده و تقید به روزه‌داری دانش‌آموزان

Crosstabulation تقدیر به روزه داری * شاخص کیفی وضعیت اقتصادی						
			شاخص کیفی وضعیت اقتصادی			Total
			پایین	متوسط	بالا	
تقدیر به روزه داری	هیچ	Count	2	0	2	4
		شاخص % within کیفی وضعیت اقتصادی	4%	0%	7%	3%
	کمتر از ده روز	Count	28	31	40	99
		شاخص % within کیفی وضعیت اقتصادی	5.6%	7.6%	13.5%	8.2%
	از ۱۰ تا ۱۴ روز	Count	12	12	20	44
		شاخص % within کیفی وضعیت اقتصادی	2.4%	2.9%	6.8%	3.6%
	از ۱۵ تا ۱۹ روز	Count	19	22	19	60
		شاخص % within کیفی وضعیت اقتصادی	3.8%	5.4%	6.4%	5.0%
	از ۲۰ تا ۲۴ روز	Count	72	87	60	219
		شاخص % within کیفی وضعیت اقتصادی	14.3%	21.3%	20.3%	18.2%
	بیشتر از ۲۵ روز	Count	369	256	155	780
		شاخص % within کیفی وضعیت اقتصادی	73.5%	62.7%	52.4%	64.7%
Total		Count	502	408	296	1206
		شاخص % within کیفی وضعیت اقتصادی	100.0%	100.0%	100.0%	100.0%

همانگونه در جدول مشاهده می‌شود از کل نمونه مورد بررسی، اکثریت دانش‌آموزان (۶۴/۷ درصد) اظهار داشته‌اند که بیشتر از ۲۵ روز در ماه رمضان روزه می‌گیرند که این خود نشان از تقید بالای آنها به روزه‌داری است. اما مقایسه نتایج بدست آمده بیانگر آن است که این نسبت در طبقات پایین (۷۳/۵ درصد) بیشتر از طبقات بالای اقتصادی (۵۲/۴ درصد) است، از سوی دیگر نتایج نشان می‌دهد تنها ۸/۲ درصد از اعضای نمونه تقید ضعیف (کمتر از ده روز) به روزه دارند. با اینحال مقایسه این نسبت در وضعیت‌های مختلف اقتصادی نشان می‌دهد تقید ضعیف در طبقات اقتصادی بالا (۱۳/۵ درصد) بیشتر از طبقات پایین اقتصادی (۵/۶ درصد) است.

جدول شماره (۳۷) تعیین سطح معناداری، شدت و جهت تأثیر وضعیت اقتصادی بر تقید به روزه‌داری دانش‌آموزان

		خطای استاندارد		مقدار آماره	میزان خطا
		Directional Measures			
		جهت و شدت	Value	Asymp. Std. Error ^a	Approx. Sig. ^b
Ordinal by Ordinal	Somers' d	Symmetric	-.165	.026	-6.343
		Dependent تقید به روزه داری	-.150	.024	-6.343
		شاخص کیفی وضعیت اقتصادی	-.182	.028	-6.343
		Dependent			.000

a. Not assuming the null hypothesis.
b. Using the asymptotic standard error assuming the null hypothesis.

جدول فوق نتایج مربوط به رابطه تقید به روزه داری دانش‌آموزان و وضعیت اقتصادی آنها را در سه حالت نشان می‌دهد:

الف- حالتی که همبستگی این دو متغیر مدنظر باشد. (Somers'd Symmetric)

ب- حالتی که تقید به روزه‌داری، متغیر وابسته باشد.

ج- حالتی که وضعیت اقتصادی، متغیر وابسته باشد.

نتایج آزمون d سامرز در ارتباط با فرضیه (ب)، نشان می‌دهد که می‌توان با بیش از ۹۹ درصد اطمینان ادعا نمود وضعیت اقتصادی خانواده تأثیر ضعیف معکوس و اما معناداری بر تقید به روزه‌داری می‌گذارد. (Somers' d = -0.15, Sig = 0.000) به این معنی که با بهبود وضعیت اقتصادی خانواده، تقید دانش‌آموزان به روزه‌داری تا حدی کم می‌شود و برعکس با افت وضعیت اقتصادی تقید دانش‌آموزان به روزه گرفتن، نسبتاً بیشتر می‌شود.

۳- ضرایب مناسب برای سنجش رابطه همبستگی بین متغیرهای رتبه‌ای

• ضرایب همبستگی تایبی و تایسی کندال

چنانچه بخواهید این فرضیه را به محک آزمون بگذارید که "بین تصویر مذهبی دانش‌آموزان از خود و تقید به روزه‌داری آنها همبستگی وجود دارد" به دلایل زیر باید از آزمون تایبی کندال یا تایسی کندال استفاده کنید.

الف) سطح سنجش: در این فرضیه دو متغیر "تصویر مذهبی از خود" و "تقید به روزه‌داری" وجود دارد که هر دو دارای سطح سنجش رتبه‌ای هستند.

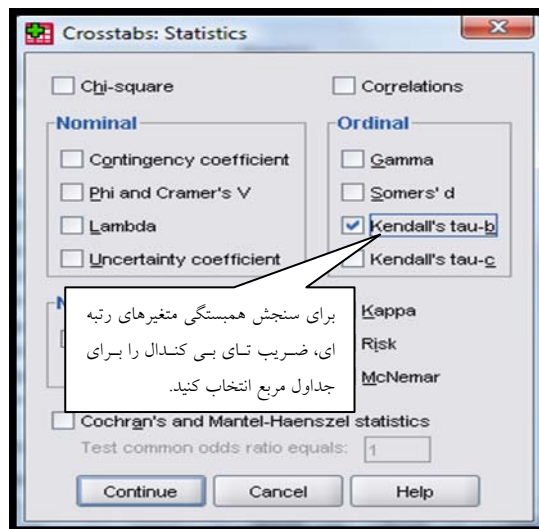
درس دهم – روش‌های آماری مناسب برای سنجش رابطه متغیرهای رتبه‌ای □ ۱۴۱

الف) سطح سنجش: در این فرضیه دو متغیر "تصویر مذهبی از خود" و "تقید به روزه‌داری" وجود دارد که هر دو دارای سطح سنجش رتبه‌ای هستند.

ب) هدف محقق: در این فرضیه بین این دو متغیر تأثیر و تأثر متقابل (همبستگی) وجود دارد.

برای این منظور نیز همانند متغیرهای اسمی در درس قبل، دستور Crosstabs را از منوی Analyze انتخاب کرده، پس از انتقال متغیرهای اسمی به سطرها و ستون‌های مربوطه بر روی Statistics کلیک کرده و به شرح تصویر زیر عمل کنید.

شکل شماره (۵۷) انتخاب آزمون‌ها و ضرایب مناسب



از آنجا که سطر و ستون جدول توافقی این دو متغیر با یکدیگر برابر است (۶ در ۶)، ضریب تای بی کندال را از میان ضرایب رتبه‌ای انتخاب کنید^۱. با اجرای این آزمون جداول خروجی زیر مشاهده می‌شود.

۱- در صورتیکه شکل جدول توافقی به صورت مستطیل باشد، در این حالت باید از آزمون تای سی کندال استفاده کرد.

جدول شماره (۳۸) خلاصه پردازش اطلاعات

Case Processing Summary					
نمونه مورد تحلیل		نمونه خارج شده از تحلیل		کل نمونه مورد بررسی	
		Cases Missing			
	N	Percent	N	Percent	N
* تقید به روزه داری تصویر مذهبی از خود	1204	99.1%	11	.9%	1215
					Percent
					100.0%

این جدول نشان می‌دهد که از مجموع ۱۲۱۵ دانش‌آموز در نمونه مورد بررسی، ۹۹/۱ درصد (۱۲۰۴ نفر) به سؤالات مربوط به تقید به روزه‌داری و تصویر مذهبی از خود، پاسخ داده‌اند و وارد این آزمون شده‌اند. تنها در حدود ۹ درصد (۱۱ نفر) از کل نمونه مورد بررسی از تحلیل حاضر خارج شده‌اند.

جدول شماره (۳۹) جدول توافقی بین دو متغیر تصویر مذهبی از خود و تقید به روزه‌داری دانش‌آموزان

Crosstabulation تقید به روزه داری * تصویر مذهبی از خود							
		تصویر مذهبی از خود					Total
		هیچ	خیلی کم	کم	متوسط	زیاد	
تقید به روزه داری	هیچ	Count 2	0	0	1	1	4
	% within خود	10.0%	.0%	.0%	.2%	.3%	.3%
	کمتر از ده روز	Count 9	9	9	41	21	98
	% within خود	45.0%	23.1%	10.3%	6.8%	6.5%	8.1%
	۱۰ تا ۱۴ روز	Count 1	2	6	21	8	44
	% within خود	5.0%	5.1%	6.9%	3.5%	2.5%	3.7%
	۱۵ تا ۱۹ روز	Count 0	6	7	26	19	60
	% within خود	.0%	15.4%	8.0%	4.3%	5.8%	5.0%
	۲۰ تا ۲۴ روز	Count 3	4	19	122	51	219
	% within خود	15.0%	10.3%	21.8%	20.3%	15.7%	18.2%
	بیشتر از ۲۵ روز	Count 5	18	46	390	225	779
	% within خود	25.0%	46.2%	52.9%	64.9%	69.2%	64.7%
Total		Count 20	39	87	601	325	1204
% within خود		100.0%	100.0%	100.0%	100.0%	100.0%	100.0%

نتایج مندرج در جدول فوق نشان می‌دهد که اکثریت نمونه مورد بررسی (۶۴/۷ درصد) اظهار داشته‌اند بیش از ۲۵ روز از ماه رمضان را مقید به روزه گرفتن هستند، اما این نسبت با بهبود تصویر مذهبی از خود، از ۲۵ درصد به ۷۲ درصد ارتقا می‌یابد (به جهت فلش در قسمت پایین جدول توجه کنید). در مقابل اگر چه حدوداً ۸ درصد از اعضای نمونه تقید ضعیف (کمتر از ده روز) به روزه دارند اما این نسبت با منفی شدن تصویر خود، از ۶/۸ به ۴۵ درصد افزایش می‌یابد (به جهت فلش در قسمت بالای جدول توجه کنید).

درس دهم – روش‌های آماری مناسب برای سنجش رابطه متغیرهای رتبه‌ای □ ۱۴۳

جدول شماره (۴۰) تعیین سطح معناداری و شدت و جهت‌بیین متغیرهای مورد سنجش

Symmetric		خطای استاندارد		مقدار آماره	میزان خطا
جهت و شدت	Value	Std. Error ^a	Approx. T	App. k	Sig.
Ordinal by Ordinal Kendall's tau-b	.121	.026	4.551		.000
N of Valid Cases	1204				

a. Not assuming the null hypothesis.
b. Using the asymptotic standard error assuming the null hypothesis.

همان‌گونه که ملاحظه می‌شود با بیش از ۹۹ درصد اطمینان می‌توان ادعا نمود که بین تصویر مذهبی از خود و تقید به روزه‌داری دانش‌آموزان همبستگی ضعیف اما مثبت و معنی داری وجود دارد (Kendall's tau-b= 0.121, Sig= 0.000).

تمرین درس دهم

۱- آیا نتایج مندرج در جداول زیر این فرضیه را که «رابطه همبستگی مستقیم و معناداری بین مدت فراغت و میزان تماشای تلویزیون وجود دارد» را تأیید می کند؟ نتایج حاصله را تفسیر کنید.

	مدت فراغت						
	هیچ	کمتر از یک ساعت	از یک تا دو ساعت	دو تا سه ساعت	سه تا چهار ساعت	بیشتر از چهار ساعت	Total
هیچ میزان تماشای تلویزیون	7.1%	19.4%	14.5%	6.1%	9.5%	10.9%	11.1%
کمتر از 30 دقیقه	14.3%	11.1%	5.8%	13.6%	5.4%	3.3%	7.4%
30 تا 60 دقیقه	14.3%	27.8%	30.4%	13.6%	12.2%	7.6%	16.5%
1 تا 2 ساعت	42.9%	36.1%	27.5%	36.4%	28.4%	19.6%	28.8%
2 تا 3 ساعت	14.3%	2.8%	15.9%	24.2%	29.7%	31.5%	23.1%
بیشتر	7.1%	2.8%	5.8%	6.1%	14.9%	27.2%	13.1%
Total	100.0%	100.0%	100.0%	100.0%	100.0%	100.0%	100.0%

Symmetric Measures

		Value	Asymp. Std. Error ^a	Approx. T ^b	Approx. Sig.
Ordinal by Ordinal	Kendall's tau-b	.266	.024	11.193	.000
	Kendall's tau-c	.255	.023	11.193	.000
	N of Valid Cases	1053			

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

۲- براساس داده های فایل satisf.sav فرضیات زیر را آزمون و نتایج حاصله را تفسیر کنید.

الف- بین رضایت از قیمت (price) و رضایت از کیفیت کالا (quality) همبستگی مستقیم و معناداری وجود دارد.

ب- سن (agecat) افراد تأثیر مستقیم بر رضایت آنان از کیفیت کالا (quality) می گذارد.

درس یازدهم

روش‌های آماری مناسب برای

سنجش رابطه متغیرهای کمی

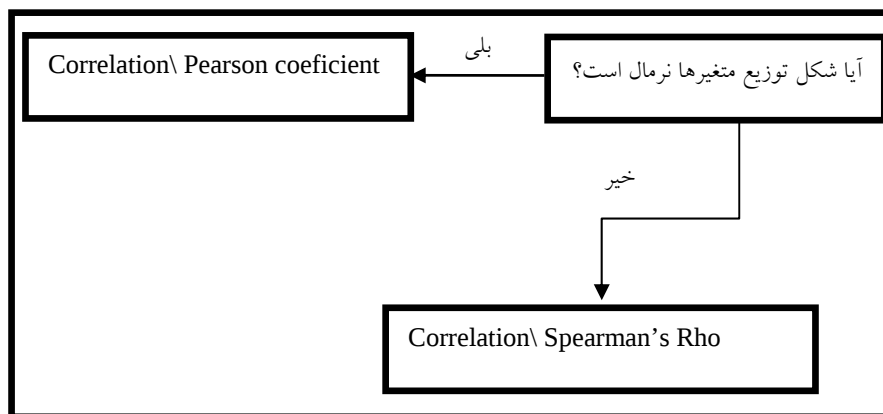
در پایان این درس انتظار می‌رود دانشجو

- ۱- موارد کاربرد آزمون فرض غیرنرمال بودن را بداند.
- ۲- توانایی آزمون فرض غیرنرمال بودن و تفسیر نتایج آن را داشته باشد.
- ۳- موارد کاربرد ضرایب همبستگی پیرسون و اسپیرمن را بداند.
- ۴- توانایی تفسیر جهت و شدت ضرایب همبستگی پیرسون و اسپیرمن را داشته باشد.
- ۵- موارد کاربرد ضریب همبستگی تفکیکی را بداند و قادر به محاسبه آن باشد.
- ۶- بتواند نتایج ضریب همبستگی تفکیکی را تفسیر کند.

۱- مقدمه

در آزمون رابطه بین متغیرهای کمی (فاصله‌ای یا نسبی)، محقق علاوه بر در نظر داشتن دو ملاک اصلی (سطح سنجش و هدف از سنجش)، باید شکل توزیع متغیرهای کمی را از حیث نرمال بودن بررسی کند. در نرم‌افزار SPSS روش رایج برای بررسی نرمال بودن توزیع متغیرهای کمی استفاده از آزمون کولموگروف اسمیرنف^۱ است. این آزمون هم از طریق منوی Explore قابل اجراست. ناپارامتریک^۲ و هم از طریق منوی Explore قابل اجراست.

شکل شماره (۵۸) راهنمای انتخاب آزمون مناسب برای رابطه متغیرهای کمی



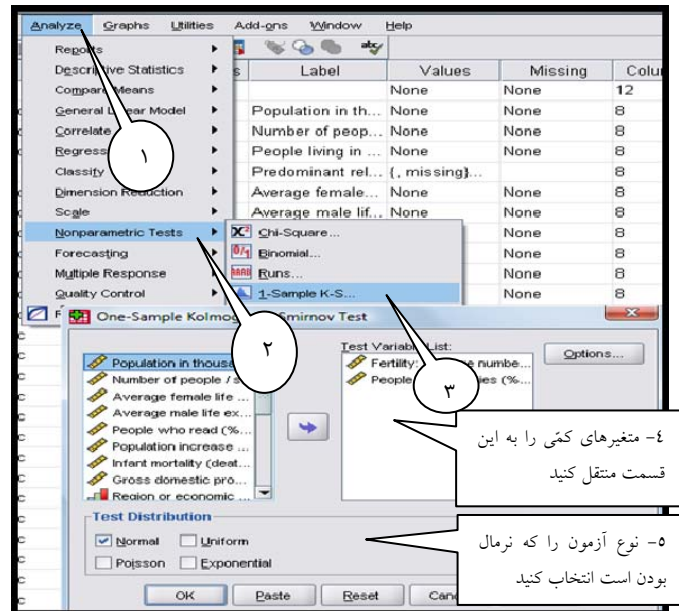
۲- آزمون فرض نرمال بودن متغیرهای کمی

اگر بخواهید این فرضیه را مورد سنجش قرار دهید که "بین نرخ باروری و درصد شهرنشینی رابطه معکوس وجود دارد." از آنجا که هر دو متغیر مزبور از نوع کمی هستند، ابتدا باید وضعیت نرمال بودن توزیع متغیرها، مورد بررسی قرار گیرد. برای این منظور باید آزمون کولموگروف اسمیرنف را از منوی آزمونهای ناپارامتریک به شرح تصویر ذیل اجرا کنید.

1- Kolmogorov-Smirnov Test
2- Nonparametric tests

درس یازدهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای کمی □ ۱۴۷

شکل شماره (۵۹) شیوه اجرای آزمون کولموگروف اسمیرنوف



با اجرای این دستور با خروجی ذیل مواجه می‌شوید.

جدول شماره (۴۱) بررسی وضعیت توزیع متغیرهای "درصد شهرنشینی و نرخ باروری"

One-Sample Kolmogorov-Smirnov Test			Fertility: average number of kids	People living in cities (%)
N			107	108
Normal Parameters ^{a,b}	Mean		3.563	56.53
	Std. Deviation		1.9025	24.203
Most Extreme Differences	Absolute		.161	.090
	Positive		.161	.063
	Negative		-.118	-.090
Kolmogorov-Smirnov Z			1.662	.932
Asymp. Sig. (2-tailed)			.008	.350

a. Test distribution is Normal.

b. Calculated from data.

همانگونه که ملاحظه می‌شود خطای مربوط به متغیر درصد شهرنشینی بیشتر از پنج صدم است ($\text{sig}=0.350$)، لذا فرض صفر (فرض نرمال بودن توزیع) رد نمی‌شود. به عبارتی دیگر توزیع متغیر درصد شهرنشینی دارای حالت نرمال است. اما تغییر نرخ باروری از خطای کمتر از پنج صدم ($\text{sig}=0.008$) برخوردار بوده و بیانگر این واقعیت است که توزیع این متغیر دارای حالت نرمال نیست به عبارتی دیگر فرض صفر (فرض نرمال بودن توزیع) در مورد این متغیر رد می‌شود.

در مثال فوق، از آنجا که لااقل یکی از دو متغیر کمی مورد بررسی دارای توزیع غیرنرمال (متغیر نرخ باروری) است، لذا برای سنجش همبستگی این دو متغیر از ضریب همبستگی اسپیرمن استفاده می‌شود.

۳- ضرایب مناسب برای سنجش رابطه همبستگی بین متغیرهای کمی

• ضریب همبستگی پیرسون

اگر بخواهید این فرضیه را مورد سنجش قرار دهید که: "بین نرخ رشد جمعیت و درصد شهرنشینی همبستگی وجود دارد"، باید به دلایل زیر از ضریب همبستگی پیرسون^۱ استفاده کنید.

• **سطح سنجش:** در این فرضیه دو متغیر "نرخ رشد جمعیت" و "درصد شهرنشینی" وجود دارد که هر دو دارای سطح سنجش فاصله‌ای (کمی) هستند.

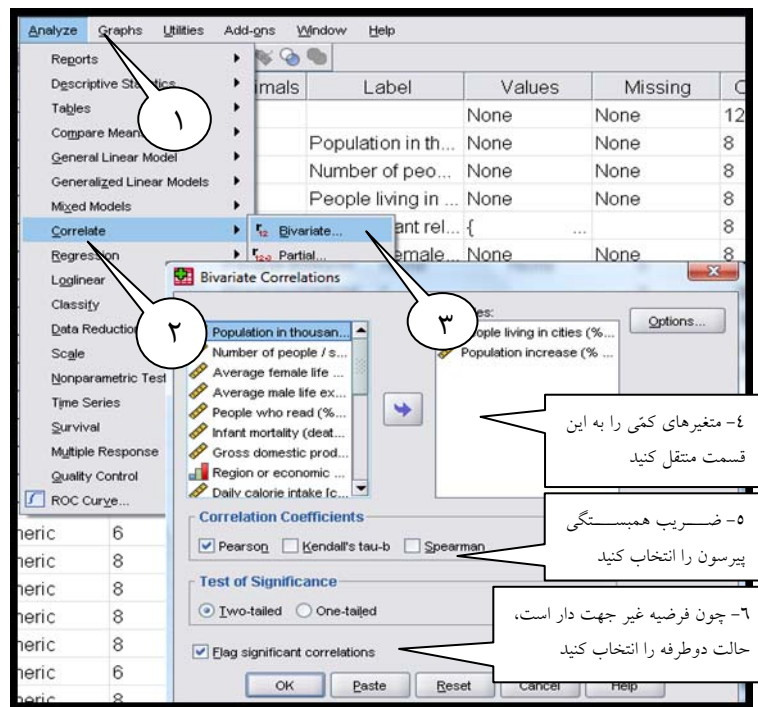
• **هدف محقق:** محقق فرض می‌کند که بین این دو متغیر تأثیر و تأثیر متقابل (همبستگی) وجود دارد.

• **شکل توزیع متغیر:** با توجه به نتیجه بدست آمده از آزمون کولمو گروف اسمیرنف شکل توزیع هر دو متغیر نرمال است. برای انجام این آزمون باید به شرح تصویر زیر عمل شود.

۱- مفروضات آزمون پیرسون عبارتست از: الف) متغیرهای مورد محاسبه از نوع کمی (فاصله‌ای یا نسبی) باشند، ب) توزیع متغیرهای کمی حالت نرمال داشته باشد. ج) رابطه متغیرهای کمی از نوع خطی باشد.

درس یازدهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای کمی □ ۱۴۹

شکل شماره (۶۰) شیوه اجرای آزمون پیرسون



جدول شماره (۴۲) نتایج سنجش همبستگی متغیرهای نرخ رشد جمعیت و درصد شهرنشینی

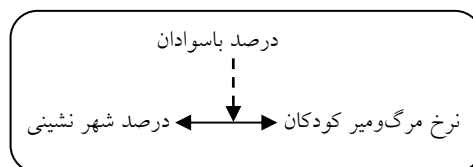
Correlations			
		People living in cities (%)	Population increase (% per year)
People living in cities (%)	Pearson Correlation	1	-.375*
	Sig. (1-tailed)		.000
	N	108	108
Population increase (% per year)	Pearson Correlation	-.375**	1
	Sig. (1-tailed)	.000	
	N	108	109

**. Correlation is significant at the 0.01 level (1-tailed).

نتایج حاصل از ضریب همبستگی پیرسون نشان می‌دهد که بین درصد شهرنشینی و نرخ رشد جمعیت در کشورهای جهان همبستگی معکوس و معناداری وجود دارد (Pearson $Correlation = -0.375, Sig = 0.000$) به این معنی که با افزایش نرخ رشد جمعیت در کشورهای جهان، درصد شهرنشینی تا حدی پایین می‌آید و از آن طرف، با افزایش درصد شهرنشینی در کشورهای جهان، نرخ رشد جمعیت تا حدی پایین می‌آید.

• ضریب همبستگی تفکیکی

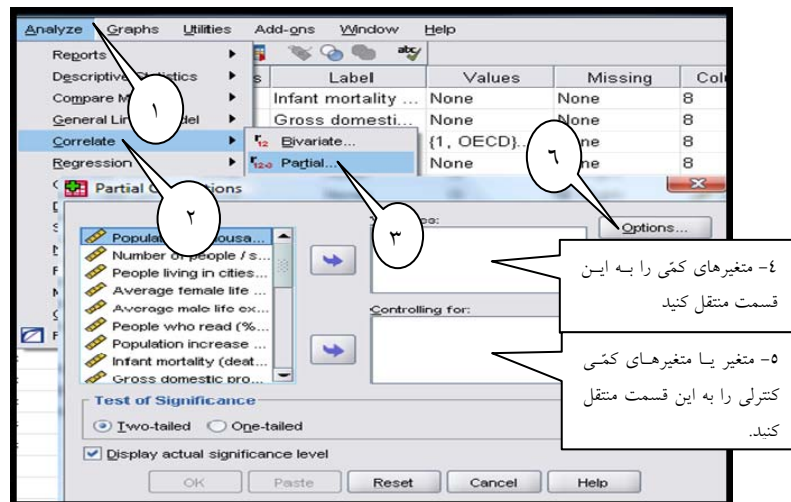
در صورتیکه بخواهید ضمن کنترل یک متغیر کمی، رابطه همبستگی بین دو متغیر کمی دیگر را اندازه‌گیری کنید، در این حالت باید از ضریب همبستگی تفکیکی استفاده کنید. مثلاً اگر به رابطه معکوس و معنادار بدست آمده بین نرخ مرگ‌ومیر کودکان و درصد شهرنشینی در کشورهای جهان تردید کردید، بدین معنا که همبستگی این دو متغیر را ناشی از متغیر سومی (درصد باسوادان کشورها) بدانید، باید برای رفع این تردید، تأثیر این متغیر سوم (درصد باسوادان کشورها) را کنترل کنید و بعد از آن به سنجش رابطه همبستگی واقعی بین نرخ مرگ‌ومیر کودکان و درصد شهرنشینی بپردازید.



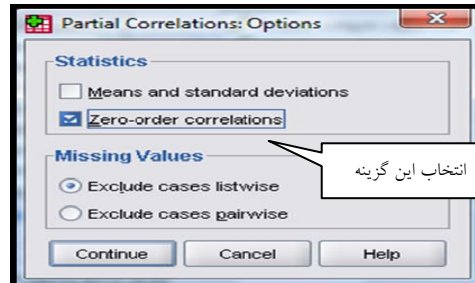
برای این منظور باید دستور Partial Correlation را از منوی Analyze به شرح تصویر زیر اجرا کند.

درس یازدهم - روش‌های آماری مناسب برای سنجش رابطه متغیرهای کمی □ ۱۵۱

شکل شماره (۶۱) شیوه اجرای ضریب همبستگی تفکیکی (مرحله ۱)



شکل شماره (۶۲) شیوه اجرای ضریب همبستگی تفکیکی (مرحله ۲)



پیشنهاد می‌شود برای تفسیر دقیق‌تر نتایج بر روی گزینه Option کلیک کرده و همبستگی مرتبه صفر^۱ این دو متغیر را نیز انتخاب کنید تا نتایج همبستگی دو متغیر را در حالت بدون کنترل نیز در خروجی‌ها نشان دهد. در این حالت می‌توان از طریق مقایسه نتایج، به اثر کنترلی متغیر یا متغیرهای کمی (قبل و بعد از کنترل) بیشتر پی برد.

1- Zero order correlation

جدول شماره (۴۳) همبستگی تفکیکی "نرخ مرگ‌ومیر کودکان" و "درصد شهر نشینی" با کنترل "درصد باسوادان"

Correlations					
ضرایب همبستگی پیرسون در مرتبه صفر (کنترل نشده)			People living in cities (%)	Infant mortality (deaths per 1000 live births)	People who read (%)
Control Variables					
-none ^a	People living in cities (%)	Correlation	1.000	-.733	.650
		Significance (2-tailed)	.	.000	.000
		df	0	105	105
	Infant mortality (deaths per 1000 live births)	Correlation	-.733	1.000	-.900
		Significance (2-tailed)	.000	.	.000
		df	105	0	105
متغیر کنترل	People who read (%)	Correlation	.650	-.900	1.000
		Significance (2-tailed)	.000	.000	.
		df	105	105	0
People who read (%)	People living in cities (%)	Correlation	1.000	-.449	
		Significance (2-tailed)	.	.000	
		df	0	104	
	Infant mortality (deaths per 1000 live births)	Correlation	-.449	1.000	
		Significance (2-tailed)	.000	.	
		df	104	0	

a. Cells contain zero-order (Pearson) correlations.

نتایج حاصله بیانگر آن است که ضریب همبستگی درصد شهرنشینی و نرخ مرگ‌ومیر کودکان در مرتبه صفر (بدون کنترل متغیر کمّی) معکوس و قوی ($-0/733$) بود، اما پس از کنترل متغیر کمّی درصد باسوادان به ($-0/449$) تقلیل پیدا می‌کند به عبارتی دیگر با کنترل این متغیر سوم (درصد باسوادان)، شدت همبستگی دو متغیر مزبور به اندازه قابل توجهی کاهش پیدا کرده است. با این حال رابطه همبستگی بین درصد شهرنشینی و نرخ مرگ‌ومیر کودکان همچنان منفی و معنادار است.

• ضریب همبستگی اسپیرمن

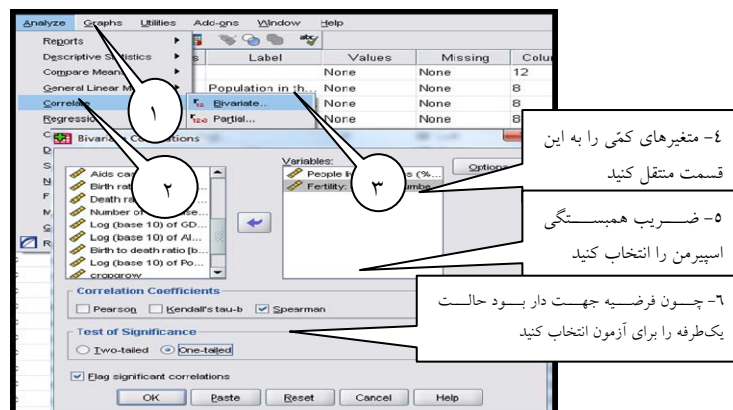
مجدداً به فرضیه‌ای که طی آن به سنجش رابطه معکوس بین نرخ باروری و درصد شهرنشینی پرداختیم^۱، برمی‌گردیم. در اینجا برای سنجش رابطه همبستگی این دو متغیر کمی به دلایل زیر باید از ضریب همبستگی اسپیرمن استفاده شود.

- **سطح سنجش:** در این فرضیه دو متغیر "نرخ باروری" و "درصد شهرنشینی" وجود دارد که یکی از آن دو (نرخ باروری)، علیرغم کمی بودن، از توزیع غیرنرمال برخوردار هستند و به همین دلیل نمی‌توان ضریب همبستگی پیرسون را برای سنجش همبستگی بین آنها به کار برد. لذا باید از ضریب همبستگی اسپیرمن استفاده کرد.

- **هدف محقق:** بین این دو متغیر تأثیر و تأثر متقابل (همبستگی) وجود دارد.

برای این منظور باید، دستور correlate را از منوی Analyze انتخاب کرده و به شرح تصویر زیر عمل کنید.

شکل شماره (۶۳) شیوه اجرای آزمون اسپیرمن



نتایج حاصله به شرح جدول بعد می‌باشد. لازم به ذکر است که در تفسیر نتایج ضرایب همبستگی پیرسون و اسپیرمن علاوه بر شدت و معناداری که در ضرایب فی و وی کرامر به آنها اشاره شد، به جهت ضریب همبستگی (مستقیم یا معکوس بودن رابطه) هم باید توجه کرد.

۱- فرضیه محقق، جهت‌دار (معکوس) و بنابر این یک‌طرفه است. در مواردی که جهت مستقیم یا معکوس رابطه همبستگی یا علی مشخص نباشد، فرضیه از نوع غیرجهت‌دار و دوطرفه است.

جدول شماره (۴۴) نتایج سنجش همبستگی متغیرهای "درصد شهرنشینی" و "نرخ باروری"

Correlations				
			People living in cities (%)	Fertility: average number of kids
Spearman's rho	People living in cities (%)	Correlation Coefficient	1.000	<u>-.570**</u>
		Sig. (1-tailed)	.	.000
		N	108	106
	Fertility: average number of kids	Correlation Coefficient	-.570**	1.000
		Sig. (1-tailed)	.000	.
		N	106	107

**. Correlation is significant at the 0.01 level (1-tailed).

نتایج حاصل از آزمون همبستگی اسپیرمن بین درصد شهرنشینی و نرخ باروری نشان می‌دهد ضریب همبستگی این دو متغیر معکوس و معنادار اما در حد متوسط است (Spearman's rho=-.570, Sig=.000). بدین معنا که با افزایش درصد شهرنشینی در کشورهای جهان، نرخ باروری تا حدی پایین می‌آید و متقابلاً با افزایش نرخ باروری در کشورهای جهان، درصد شهرنشینی تا حدی پایین می‌آید.

۱- چون رابطه از نوع همبستگی است به همین دلیل باید تفسیر حالت دو سویه داشته باشد.

درس یازدهم – روش‌های آماری مناسب برای سنجش رابطه متغیرهای کمی □ ۱۵۵

تمرینات درس یازدهم

۱- با توجه به نتایج مندرج در جدول به سوالات بعد پاسخ دهید.

Correlations

Control Variables		مدت فراغت	مدت ورزش	وضعیت درآمد خانواده
مدت فراغت -none ^a	Correlation	1.000	.254	.292
	Significance (2-tailed)	.	.000	.000
	df	0	1144	1144
مدت ورزش	Correlation	.254	1.000	.070
	Significance (2-tailed)	.000	.	.018
	df	1144	0	1144
وضعیت درآمد خانواده	Correlation	.292	.070	1.000
	Significance (2-tailed)	.000	.018	.
	df	1144	1144	0
وضعیت درآمد خانواده	Correlation	1.000	.245	
	Significance (2-tailed)	.	.000	
	df	0	1143	
مدت ورزش	Correlation	.245	1.000	
	Significance (2-tailed)	.000	.	
	df	1143	0	

a. Cells contain zero-order (Pearson) correlations.

- الف- نتایج مربوط به همبستگی مدت فراغت و مدت پرداختن به ورزش را تفسیر کنید.
 ب- چه تفاوتی در رابطه همبستگی این دو متغیر در زمان قبل و بعد از کنترل وضعیت درآمد ایجاد می‌شود؟

۲- با توجه به فایل Employee data

- الف- همبستگی حقوق اولیه و فعلی را محاسبه و نتایج مربوط را تفسیر کنید.
 ب- چه تفاوتی در رابطه همبستگی این دو متغیر در زمان قبل و بعد از کنترل تحصیلات می‌شود؟

بخش چهارم

آمار استنباطی
مقایسه‌ها

درس دوازدهم

آزمون‌های مناسب برای

مقایسه گروه‌ها

در پایان این درس انتظار می‌رود دانشجو

- ۱- با وضعیت‌های مختلفی که در مقایسه بین متغیرهای سطوح مختلف سنجش ایجاد می‌شود آشنا شده و تفاوت‌های آنها را توضیح دهد.
- ۲- وضعیت‌های سه‌گانه‌ی آزمون برابری میانگین‌ها را شناسایی و کاربرد آنها را بداند.
- ۳- بتواند فرضیات متناسب با آزمون‌های برابری میانگین‌ها را آزمون و خروجی‌های مربوطه را تفسیر کند.

۱- مقدمه

ممکن است در انجام تحقیق، علاوه بر بررسی رابطه همبستگی و یا علی موجود در فرضیه‌ها و یا سؤالات تحقیق، قصد مقایسه نیز داشته باشید. اگر متغیرها را به لحاظ سطح سنجش در دو دسته کیفی (اسمی و رتبه‌ای) و کمی (فاصله‌ای و نسبی) قرار دهیم، به لحاظ نظری می‌توان ۶ حالت برای مقایسه در نظر گرفت. به جدول زیر دقت کنید.^۱

جدول شماره (۴۵) حالت‌های مختلف مقایسه در فرضیه‌ها یا سؤالات تحقیق از حیث سطوح

سنجش متغیرها

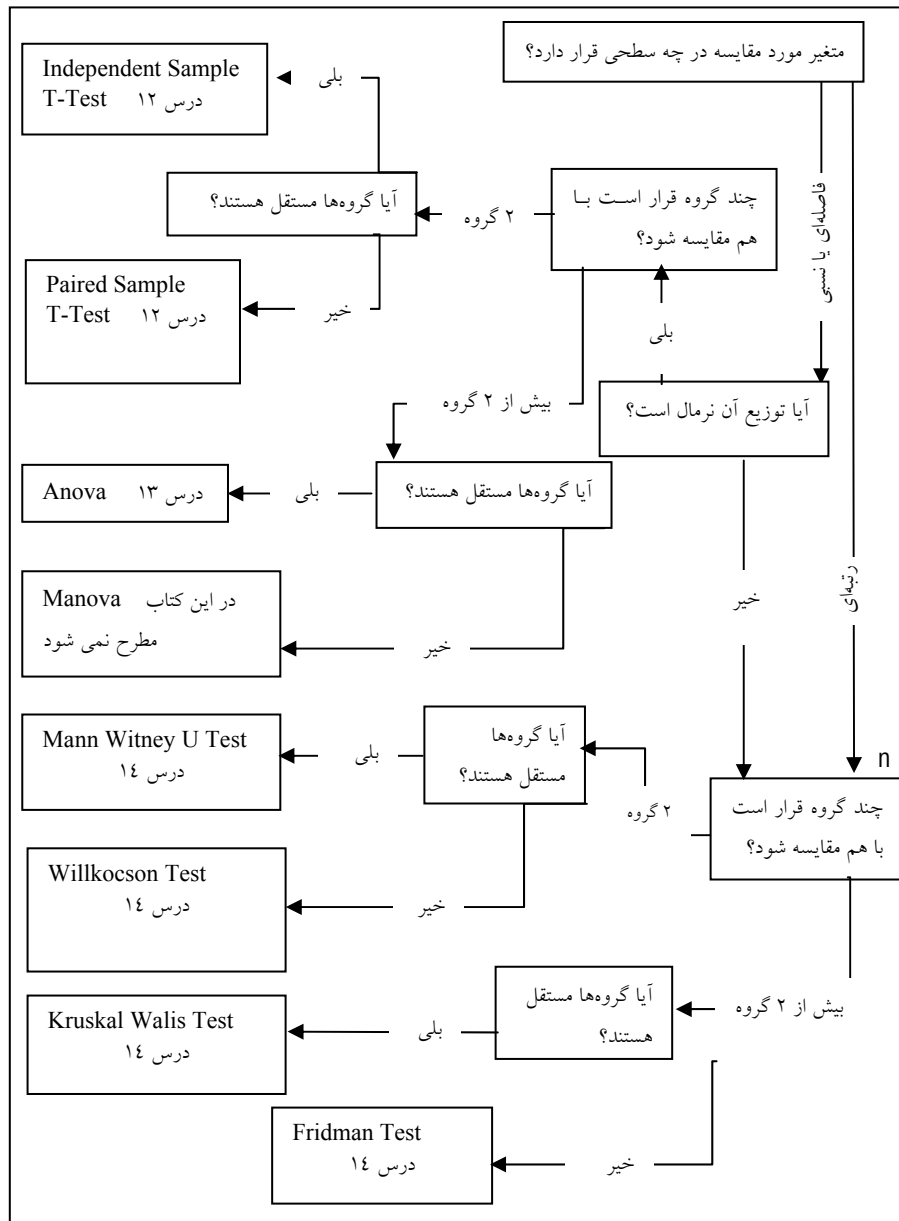
اسمی	رتبه‌ای	کمی
اسمی	- مقایسه رضایت از تحصیل در بین دانشجویان دختر و پسر - مقایسه رضایت از تحصیل در بین دانشجویان علوم انسانی، علوم پایه و فنی-مهندسی	- مقایسه میانگین نمرات درس زبان بین دانشجویان دختر و پسر - مقایسه میانگین نمرات درس زبان بین دانشجویان علوم انسانی، علوم پایه و فنی-مهندسی
رتبه‌ای	- مقایسه رضایت از شغل و رضایت از زندگی در بین زنان کارمند	- مقایسه میانگین نمرات درس زبان بر حسب نگرش (مثبت، خنثی و منفی) نسبت به استاد
کمی		- مقایسه میانگین نمرات پیش آزمون و پس آزمون درس زبان در بین دانشجویان - مقایسه میانگین نمرات زبان، آمار و روش تحقیق در بین دانشجویان - مقایسه میانگین نمرات زبان با یک عدد فرضی (مثلاً ۱۷) ^۲

۱- در اینجا به لحاظ سرفصل‌های درس آمار در دوره کارشناسی، تنها به برخی از مقایسه‌های جدول فوق می‌پردازیم.

۲- در این حالت که همان آزمون برابری میانگین‌ها به صورت یک گروهی برای آن مناسب است. مقدار عدد ثابت می‌تواند براساس تحقیقات پیشین یا میانگین جامعه آماری و یا یک تئوری خاص تعیین شود.

درس دوازدهم - آزمون‌های مناسب برای مقایسه گروه‌ها □ ۱۶۱

شکل شماره (۶۴) راهنمای انتخاب آزمون‌های آماری مناسب برای مقایسه گروه‌ها

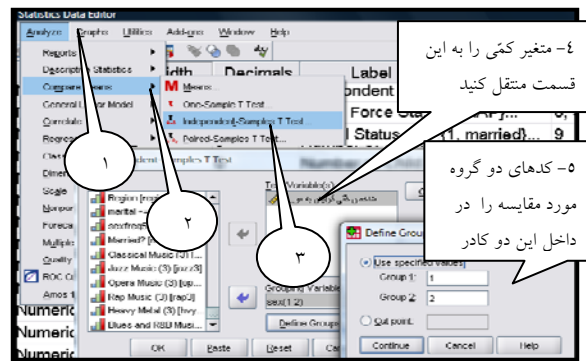


۲- مقایسه میانگین یک متغیر کمی در دو گروه

• آزمون برابری میانگین‌ها برای نمونه‌های مستقل

اگر بخواهید میانگین یک متغیر کمی را در دو گروه جداگانه^۱ مقایسه کنید، باید از آزمون "برابری میانگین‌ها برای نمونه‌های مستقل"^۲ استفاده کنید.^۳

شکل شماره (۶۵) شیوه اجرای آزمون برابری میانگین‌ها برای نمونه‌های مستقل (T-test)



۱- در اینجا لازم است درباره مفهوم گروه‌های مستقل و همبسته توضیح مختصری بدهیم. هرگاه اعضای که قرار است میانگین نمراتشان مورد مقایسه قرار گیرد در گروه‌های مجزایی (مثلاً زن و مرد یا مشاغل مختلف) عضویت داشته باشند، اصطلاحاً گروه‌های مستقل نامیده می‌شود. اما مواردی نیز هست که میانگین نمرات یک متغیر که در دو نوبت در یک گروه اندازه‌گیری شده است با یکدیگر مقایسه می‌شود مثلاً میانگین نمرات درس زبان میان ترم و پایان ترم یک گروه از دانشجویان. در اینجا اصطلاحاً با گروه‌های همبسته مواجهیم. همانطور که ملاحظه می‌شود در گروه‌های مستقل هر فرد عضو نمونه، یک نمره مجزا دارد اما در گروه‌های همبسته هر یک از اعضای نمونه دو نمره مجزا دارد.

2. Independent Samples Test

۳- مفروضات آزمون برابری میانگین‌ها برای نمونه‌های مستقل عبارتست از :

۱. حجم نمونه در دو گروه حتی‌الامکان برابر باشد.
۲. در طرح‌های آزمایشی انتساب اعضا به گروه‌ها به صورت تصادفی باشد.
۳. توزیع متغیر کمی در دو گروه باید نرمال باشد. مخصوصاً وقتی حجم نمونه در گروه‌ها کمتر از ۱۵ باشد. البته برای نمونه‌های بیشتر از ۴۰ تا حدی از این معیار اغماض می‌شود.
۴. واریانس متغیر کمی در دو گروه باید هم‌اندازه باشد. آزمون لون در نرم‌افزار SPSS به همین منظور اعمال می‌شود. لازم به ذکر است که اگر حجم نمونه نابرابر و واریانس درون گروهی متفاوت باشد نتایج قابل اعتماد نخواهد بود.
۵. نمونه‌های دو گروه باید مستقل از یکدیگر باشند به عبارتی دقیقتر نمرات افراد در یک گروه متأثر از نمرات اعضای گروه دیگر نباشد.
۶. در صورتی که داده‌ها فاقد مفروضات فوق باشند باید از آزمون غیرپارامتریک U مان ویتنی استفاده نمود.

درس دوازدهم - آزمون‌های مناسب برای مقایسه گروه‌ها □ ۱۶۳

مثلاً محققی را در نظر بگیرید که می‌خواهد به آزمون این فرضیه بپردازد که "شاخص کمی گرایش به موسیقی^۱ در بین زنان میانگین پایین‌تری در مقایسه با مردان دارد". به این ترتیب در اولین گام پس از اطمینان از نرمال بودن توزیع متغیر گرایش به موسیقی^۲، با اجرای دستور فوق با خروجی ذیل روبرو می‌شود.

جدول شماره (۴۶) مقایسه میانگین شاخص کمی گرایش به موسیقی در بین زنان و مردان

Group Statistics				
جنسیت	N	Mean	Std. Deviation	Std. Error Mean
MUSIC مردان	636	2.9429	.55038	.02182
زنان	851	2.9052	.56209	.01927

این جدول نشان می‌دهد که نمونه مورد بررسی مرکب از ۸۵۱ نفر زن و ۶۳۶ نفر مرد بوده است. همانگونه که ملاحظه می‌شود در اینجا میانگین گرایش به موسیقی در بین زنان اندکی (حدود ۰/۰۴) از مردان کمتر است (میانگین ۲/۹۰ در برابر ۲/۹۴). شاخص انحراف معیار نیز حاکی از آن است که پراکندگی گرایش به موسیقی در دو گروه تقریباً شبیه هم است (انحراف معیار ۰/۵۶ در برابر ۰/۵۵). مقایسه خطای میانگین نیز حاکی از آن است که مقدار این شاخص در بین زنان اندکی کمتر از مردان است (۰/۱۹ در برابر ۰/۲۱).

نتایج فوق توصیف متغیر کمی در سطح نمونه مورد بررسی بود. اما برای پاسخ به این سؤال که آیا تفاوت‌های موجود بین دو گروه زنان و مردان از حیث میانگین و انحراف معیار مشاهده شده، قابلیت تعمیم به جامعه آماری را دارد یا خیر، باید به نتایج آزمون معناداری توجه کرده و آنها را با یکدیگر مقایسه کنید.

۱- گرایش به موسیقی از ترکیب گویه‌های متعددی طراحی شده است، لذا می‌توان آن را در شمار متغیرهای کمی محسوب نمود.

۲- شرح مراحل مربوط به آزمون نرمال بودن توزیع در درس دهم آمده است. لازم به ذکر است که در چنین مواقعی اگر توزیع متغیر کمی نرمال نباشد، باید از آزمون ناپارامتریک لمان ویتنی استفاده کنیم. این آزمون در درس سیزدهم توضیح داده خواهد شد.

جدول شماره (۴۷) آزمون مقایسه میانگین شاخص کمی گرایش به موسیقی در بین زنان و مردان

Independent Samples Test									
		Levene's Test for Equality of Variances		t-test for Equality of Means					
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference Lower Upper
MUSIC	Equal variances assumed	.043	.835	1.289	1485	.198	.0376	.02920	-.01963 .09493
	Equal variances not assumed			1.293	1383.006	.196	.0376	.02911	-.01946 .09476

نتایج آزمون لون (Leven's test) حاکی از آن است که تفاوت ناچیزی که در پراکندگی گرایش به موسیقی در بین زنان و مردان دیده شد، قابلیت تعمیم به جامعه آماری را ندارد ($F=0.43$, $Sig=0.83$). بنابراین چون فرض برابری واریانس رد نشده است، برای قضاوت در مورد برابری میانگین‌ها به آماره T در سطر اول مراجعه می‌کنیم.^۱ همانطور که ملاحظه می‌شود در اینجا نیز نتایج آزمون برابری میانگین‌ها بیانگر این واقعیت است که تفاوت اندکی که بین میانگین گرایش به موسیقی زنان و مردان در سطح نمونه مشاهده شد، قابلیت تعمیم به جامعه آماری را ندارد ($T=1.289$, $df=1485$, $Sig=0.19$). لذا می‌توان ادعا کرد که زنان و مردان در جامعه آماری، دارای گرایش‌های مشابه به موسیقی هستند.

در مثال دوم محقق را در نظر بگیرید که قصد آزمون این فرضیه را دارد "میانگین حقوق فعلی زنان از مردان پایین تر است" به این ترتیب در اولین گام پس از اطمینان از نرمال بودن توزیع متغیر حقوق، با اجرای دستور فوق با خروجی ذیل روبرو می‌شوید.

۱- همانطور که در جدول شماره (۴۷) مشاهده می‌شود، دو سطر مجزا برای تفسیر نتایج آزمون برابری میانگین‌ها وجود دارد، سطر اول در صورتی که تفاوت واریانس‌ها معنادار نباشد، مورد استناد قرار گیرد و سطر دوم در صورت معنادار بودن تفاوت واریانس‌ها.

درس دوازدهم - آزمون‌های مناسب برای مقایسه گروه‌ها □ ۱۶۵

جدول شماره (۴۸) مقایسه میانگین حقوق فعلی در بین زنان و مردان

Group Statistics					
جنسیت	N	Mean	Std. Deviation	Std. Error Mean	
Current Salary زنان	216	\$26,031.92	\$7,558.021	\$514.258	
مردان	258	\$41,441.78	\$19,499.214	\$1,213.968	

این جدول نشان می‌دهد که نمونه مورد بررسی مرکب از ۲۱۶ نفر زن و ۲۵۸ نفر مرد بوده است. همانگونه که ملاحظه می‌شود در اینجا میانگین حقوق در بین زنان در حدود ۱۵ دلار از مردان کمتر است.

(میانگین حقوق 26,031.92 در برابر 41,441.78). شاخص انحراف معیار نیز حاکی از آن است که پراکندگی میزان حقوق در بین زنان و مردان بایکدیگر متفاوت است (انحراف معیار 7,558 در برابر 19,499). مقایسه خطای میانگین نیز حاکی از آن است که مقدار این شاخص در بین زنان بسیار کمتر از مردان است (514.258 در برابر 1,213.968).

جدول شماره (۴۹) آزمون مقایسه میانگین حقوق فعلی در بین زنان و مردان

Independent Samples Test									
	Levene's Test for Equality of Variances		t-test for Equality of Means						
								95% Confidence Interval of the Difference	
	F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	Lower	Upper
CurrentSalary Equal variances assumed	119.669	.000	-10.945	472	.000	\$-15,409.8	\$1,407.9	\$-18,176.4	\$-12,643.3
Equal variances not assumed			-11.688	344.262	.000	\$-15,409.8	\$1,318.4	\$-18,002.9	\$-12,816.7

نتایج آزمون لون (Leven's test) حاکی از آن است که تفاوت‌های مشاهده شده در خصوص پراکندگی حقوق زنان و مردان، قابل تعمیم به جامعه آماری نیز هست (F=119.669, Sig=.000). بنابراین چون دو گروه زنان و مردان از حیث میزان حقوق دارای واریانس‌های متفاوتی هستند (معنادار بودن تفاوت واریانس‌ها)، برای قضاوت در مورد برابری میانگین‌ها به آماره T در سطر دوم مراجعه می‌کنیم. همانطور که ملاحظه می‌شود نتایج آزمون برابری میانگین‌ها نیز نشان می‌دهد که اختلاف مشاهده شده در میانگین حقوق

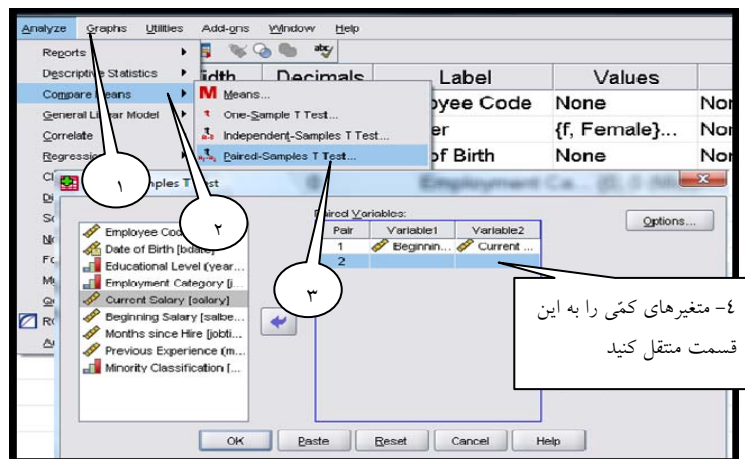
زنان و مردان، قابل تعمیم به جامعه آماری را دارد (T=-11.688. df=344.262, Sig=.000).

۳- مقایسه میانگین دو متغیر کمی در یک گروه

• آزمون برابری میانگین‌ها برای نمونه‌های همبسته

اگر بخواهید میانگین دو متغیر کمی را در یک گروه واحد با یکدیگر مقایسه کنید، برای این کار باید از آزمون "برابری میانگین‌ها برای نمونه‌های همبسته" استفاده کنید.^۲

شکل شماره (۶۶) شیوه اجرای آزمون برابری میانگین‌ها برای نمونه‌های همبسته (paired-sample)



۱-Paired sample T-Test

۲- مفروضات آزمون برابری میانگین‌ها برای نمونه‌های همبسته عبارتست از:

- متغیرهایی که مورد مقایسه قرار می‌گیرد باید از نوع کمی باشد.
- وزیع متغیرهای کمی باید نرمال باشد. مخصوصاً وقتی حجم نمونه کمتر از ۱۵ باشد. البته برای نمونه‌های بیشتر از ۳۰ تا حدی از این معیار اغماض می‌شود.
- واریانس متغیرهای کمی باید هم‌اندازه باشد.
- در صورتی که داده‌ها فاقد مفروضات فوق باشند باید از آزمون غیرپارامتریک ویلکاکسون استفاده نمود.

درس دوازدهم – آزمون‌های مناسب برای مقایسه گروه‌ها □ ۱۶۷

اگر بخواهید این فرضیه را آزمون کنید که "میانگین حقوق بدو استخدام کارمندان یک شرکت از میانگین حقوق فعلی آنها پایین تر است"^۱ پس از اطمینان از رعایت مفروضات دراولین گام، با اجرای دستور فوق با خروجی ذیل روبرو می‌شوید.

جدول شماره (۵۰) مقایسه وضعیت آماره‌های دو متغیر "حقوق فعلی" و "حقوق بدو استخدام"

Paired Samples Statistics					
		Mean	N	Std. Deviation	Std. Error Mean
Pair	حقوق فعلی	\$34,419.57	474	\$17,075.661	\$784.311
1	حقوق بدو استخدام	\$17,016.09	474	\$7,870.638	\$361.510

این جدول نشان می‌دهد که تحقیق بر روی ۴۷۴ نفر از کارمندان شرکت انجام شده است. همانگونه که ملاحظه می‌شود میانگین حقوق فعلی این کارمندان (۳۴/۴۱۹ دلار) در مقایسه با میانگین حقوق آنها در زمان بدو استخدام (۱۷,۰۱۶ دلار)، افزایش یافته است. انحراف معیار نیز نشان می‌دهد که پراکندگی حقوق فعلی کارمندان حول میانگین ۱۷,۰۷۵ دلار است اما این شاخص در مورد حقوق بدو استخدام (۷,۸۷۰ دلار) می‌باشد. لذا می‌توان این گونه استنباط کرد که حقوق فعلی کارمندان در مقایسه با حقوق آنها در زمان بدو استخدام، شباهت بیشتری با یکدیگر (یا پراکندگی کمتری از یکدیگر) دارد. شاخص اشتباه استاندارد نیز حاکی از این است که خطای برآورد میانگین حقوق فعلی بیش از دو برابر خطای برآورد حقوق بدو استخدام است.^۲

۱- همانطور که ملاحظه می‌شود در اینجا با یک گروه مواجه هستیم که قرار است مقادیر دو متغیر (حقوق فعلی و حقوق اولیه) آنها با یکدیگر مقایسه شود.

۲- اگر از رابطه $\bar{x} \pm t.S_e$ استفاده کنیم می‌توان با ۹۵ درصد اطمینان ادعا کرد متوسط حقوق فعلی جامعه آماری بین ۳۲۸۸۲ و ۳۵۹۵۶ دلار قرار دارد. اما حدود اطمینان برای حقوق بدو استخدام جامعه آماری ۱۶۳۰۷ و ۱۷۷۲۴ دلار می‌باشد.

جدول شماره (۵۱) وضعیت همبستگی بین حقوق فعلی و حقوق بدو استخدام

Paired Samples Correlation				شدت و جهت همبستگی بین دو متغیر
	N	Correlation	Sig.	
Pair 1 حقوق فعلی & حقوق بدو استخدام	474	.880	.000	

نتایج حاصل از ضریب همبستگی نشان می‌دهد که بین حقوق فعلی و حقوق بدو استخدام کارکنان این شرکت، همبستگی مثبت و قوی و همچنین معناداری وجود دارد. ($\text{Correlation}=.880, \text{Sig}=.000$) به عبارتی دیگر آن دسته از کارمندانی که حقوقشان در بدو استخدام بالا بوده است، حقوق فعلی‌شان نیز به میزان زیادی بالاست و کسانی که در بدو استخدام حقوق پایینی داشته‌اند، حقوق فعلی پایینی دارند.

حال می‌خواهیم با توجه به خروجی بعد، مشخص کنیم آیا تفاوت‌های مشاهده شده بین حقوق فعلی و بدو استخدام نمونه مورد بررسی قابلیت تعمیم به جامعه آماری را دارد یا خیر.

جدول شماره (۵۲) نتیجه آزمون معنی‌داری برابری میانگین‌ها برای دو متغیر (حقوق فعلی و حقوق بدو استخدام)

Paired Samples T										تفاوت زوجها (حقوق بدو استخدام و حقوق فعلی)	
		Paired Differences					t	df	Sig. (2-tailed)		
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference						
					Lower	Upper					
Pair 1	فی فعلی - حقوق بدو استخدام	\$17,403.48	\$10,814.620	\$496.732	\$16,427.41	\$18,379.56	35.036	473	.000		

جدول فوق نشان می‌دهد متوسط تفاضل حقوق بدو استخدام و فعلی نمونه مورد بررسی ۱۷۴۰۳ دلار و خطای آن ۴۹۶ دلار است. شاخص انحراف معیار بیانگر این واقعیت است که پراکندگی تفاضل حقوق فعلی و بدو استخدام ۱۰۸۱۴ دلار می‌باشد. حدود اطمینان بدست آمده نیز نشان می‌دهد که با ۹۵ درصد اطمینان می‌توان ادعا نمود که تفاوت حقوق فعلی و حقوق بدو استخدام جامعه آماری بین ۱۶۴۲۷ تا ۱۸۳۷۹ دلار است. مقدار خطای محاسبه شده نشان

درس دوازدهم - آزمون‌های مناسب برای مقایسه گروه‌ها □ ۱۶۹

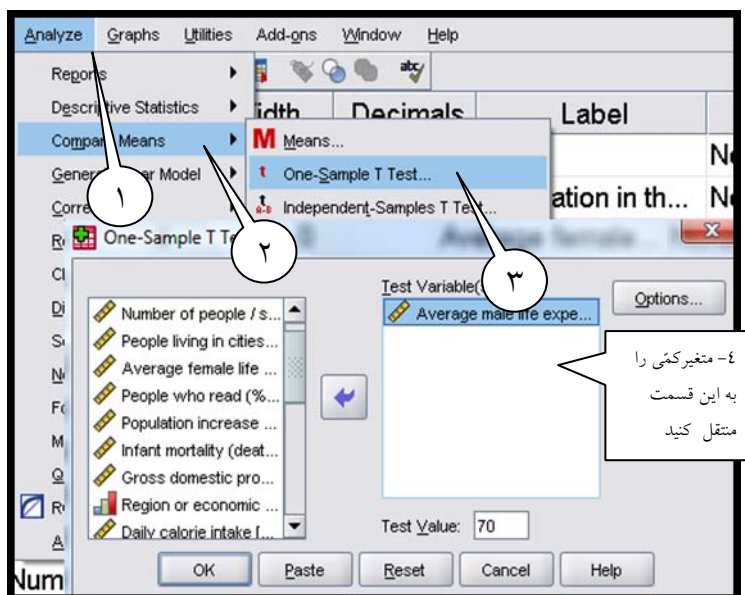
می‌دهد که تفاوت مشاهده شده بین حقوق فعلی و بدو استخدام نمونه مورد بررسی، قابلیت تعمیم به جامعه آماری را دارد.

۴- مقایسه میانگین یک متغیر کمی با یک مقدار ثابت

• آزمون برابری میانگین‌ها برای یک نمونه

اگر قصد مقایسه میانگین یک متغیر کمی را با یک مقدار ثابت داشته باشید، در این صورت باید از آزمون "برابری میانگین‌ها برای یک نمونه" استفاده کنید.^۲

شکل شماره (۶۷) شیوه اجرای آزمون برابری میانگین‌ها با یک مقدار ثابت (one-sample)



مثلاً اگر قرار باشد این فرضیه را آزمون کنید که "امید به زندگی مردان در کل کشورهای جهان، به طرز معناداری کمتر از ۷۰ سال است"، در اولین گام پس از اطمینان از رعایت مفروضات، با اجرای دستور فوق با خروجی ذیل روبرو می‌شوید.

1- One sample T-Test

۲- مفروضات آزمون برابری میانگین‌ها برای یک گروه واحد عبارتست از:

- متغیری که مورد مقایسه قرار می‌گیرد باید از نوع کمی باشد.

- توزیع متغیر کمی باید نرمال باشد. مخصوصاً وقتی حجم نمونه کمتر از ۱۵ باشد. البته برای نمونه‌های بیشتر از ۴۰ تا حدی این معیار قابل چشم‌پوشی است.

جدول شماره (۵۳) توصیف متغیر امید به زندگی مردان

One-Sample Statistics				
	N	Mean	Std. Deviation	Std. Error Mean
امید به زندگی مردان	109	64.92	9.273	.888

جدول فوق نشان می‌دهد که این مطالعه بر روی ۱۰۹ کشور جهان انجام شده است. همانگونه که ملاحظه می‌شود متوسط امید به زندگی مردان در این کشورها ۶۴ سال است. شاخص انحراف معیار نیز بیانگر آن است که متوسط فاصله امید به زندگی مردان از میانگین در حدود ۹ سال است حال برای تعمیم نتایج بدست آمده از نمونه مورد بررسی (۱۰۹ کشور جهان)، به جامعه آماری (کل کشورهای جهان) به خروجی ذیل توجه کنید.

جدول شماره (۵۴) نتیجه آزمون معنی‌داری برابر ی میانگی‌ها با یک مقدار ثابت

One-Sample Test						
	Test Value = 70					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
امید به زندگی مردان	-5.723	108	.000	-5.08	-6.84	-3.32

این جدول نتایج حاصل از آزمون "برابری میانگین‌ها برای یک نمونه" را گزارش می‌دهد. همانگونه که ملاحظه می‌شود مقدار خطای محاسبه شده بیانگر آن است که فرضیه محقق رد نمی‌شود. به این معنی که امید به زندگی مردان در در جامعه آماری به طرز معناداری کمتر از ۷۰ سال است. به گونه‌ای که با ۹۵ درصد اطمینان می‌توان ادعا نمود که امید به زندگی مردان در کشورهای جهان بین ۳/۳ تا ۶/۸ سال از ۷۰ کمتر است.

درس دوازدهم – آزمون‌های مناسب برای مقایسه گروه‌ها □ ۱۷۱

تمرینات درس دوازدهم

۱- آیا با توجه به نتایج مندرج در جدول زیر می‌توان ادعا نمود تفاوت معناداری بین مدت ورزش در بین زنان و مردان وجود دارد؟ نتایج را تفسیر کنید.

Group Statistics

جنسیت	N	Mean	Std. Deviation	Std. Error Mean
زن مدت ورزش	612	156.6961	239.00920	9.66137
مرد	585	253.4769	351.42426	14.52960

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						
										95% Confidence Interval of the Difference
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference		
مدت ورزش	Equal variances assumed	4.29	.000	-5.59	1195	.00	-96.78	17.30	-130.73	-62.83
	Equal variances not assumed			-5.54	1.0E3	.00	-96.78	17.44	-131.01	-62.54

۲- آیا براساس نتایج جدول زیر تفاوت معناداری بین مدت پرداختن به فوتبال و پیاده‌روی وجود دارد؟

Paired Samples Statistics

	Mean	N	Std. Deviation	Std. Error Mean
Pair 1 پیاده روی	60.4318	1188	128.24252	3.72069
فوتبال	31.3005	1188	114.17416	3.31253

Paired Samples Correlations

	N	Correlation	Sig.
Pair 1 پیاده روی & فوتبال	1188	-.090	.002

Paired Samples Test

	Paired Differences							
	Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference		t	df	Sig. (2-tailed)
				Lower	Upper			
Pair 1 پیاده روی - فوتبال	29.13131	179.17463	5.19838	18.93227	39.33036	5.604	1187	.000

۳- براساس داده‌های فایل Employee data.sav آیا تفاوت معناداری بین حقوق بدو استخدام (salbegin) کارکنان زن و مرد وجود دارد؟

۴- براساس داده‌های فایل World95.sav آیا بین امید به زندگی زنان کشورهای مختلف جهان با امید به زندگی مردان همان کشورها تفاوت معناداری وجود دارد؟

۵- براساس داده‌های فایل Employee data.sav آیا حقوق بدو استخدام (salbegin) کارکنان بطرز معناداری با ۳۰.۰۰۰ دلار تفاوت دارد؟

درس سیزدهم

تحلیل واریانس یک طرفه

در این درس انتظار می‌رود دانشجو

- ۱- موارد کاربرد تحلیل واریانس یک طرفه را بداند.
- ۲- بتواند فرضیات متناسب با تحلیل واریانس را آزمون کند.
- ۳- بتواند موارد استفاده از آزمون‌های تعقیبی را تشخیص دهد.
- ۴- بتواند نتایج تحلیل واریانس یک طرفه را تفسیر کند.

۱- مقدمه

اگر بخواهید میانگین یک متغیر کمی را در بیش از دو گروه^۱ مقایسه کنید، باید از آزمون تحلیل واریانس یک‌طرفه^۲ استفاده کنید^۳. گروه‌ها در اینجا همان طبقات متغیر کیفی هستند. به عنوان مثال برای مقایسه درآمد گروه‌های مختلف شغلی باید از آزمون تحلیل واریانس یک‌طرفه استفاده کنید.

۲- مقایسه میانگین یک متغیر کمی در بیش از دو گروه

آزمون تحلیل واریانس یک‌طرفه

این آزمون اگرچه به مقایسه واریانس بین گروهی و درون گروهی متغیر کمی در دو یا چند گروه می‌پردازد، اما واریانس بین گروهی در این آزمون براساس تفاوت میانگین‌های هر گروه از میانگین کل محاسبه می‌شود. توجه داشته باشید که تنها بررسی وجود یا عدم تفاوت معنادار در این آزمون کافی نیست، بلکه لازم است از طریق انجام آزمونهای تعقیبی این مورد بررسی شود که کدامیک از گروه‌های مورد مقایسه با یکدیگر تفاوت دارند.

۱- این آزمون معمولاً برای متغیرهای کیفی که بیش از دو مقوله دارند استفاده می‌شود اما منطقاً و به لحاظ آماری استفاده از آن برای مقایسه دو گروه ایرادی ندارد. از این آزمون به عنوان آزمون F نیز یاد شده است.

2- one-way ANOVA

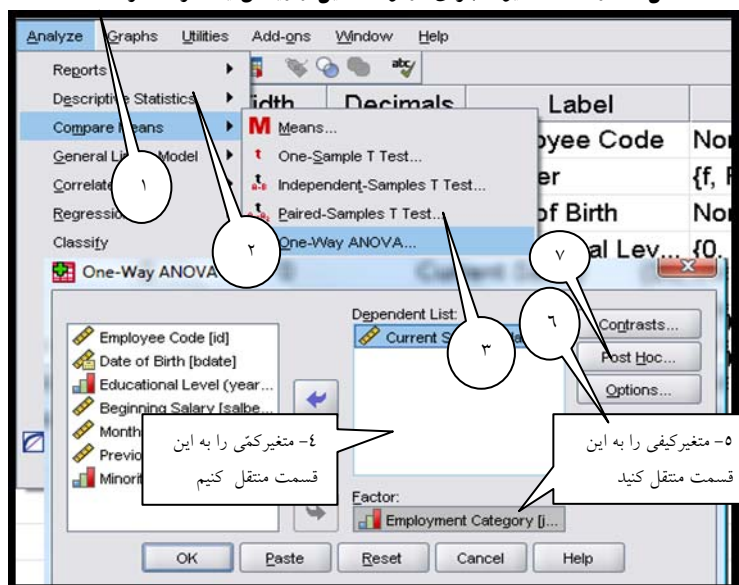
۳- مفروضات آزمون تحلیل واریانس عبارتست از:

۱. متغیری که در گروه‌ها مورد مقایسه قرار می‌گیرد باید از نوع کمی باشد.
۲. حجم نمونه در گروه‌ها حتی‌الامکان برابر باشد.
۳. در طرح‌های آزمایشی انتساب اعضا به گروه‌ها به صورت تصادفی باشد.
- توزیع متغیر کمی در گروه باید نرمال باشد. البته برای نمونه‌های بزرگ تا حدی این معیار قابل چشم‌پوشی است. لازم به ذکر است که اگر حجم نمونه نابرابر و واریانس درون گروهی متفاوت باشد نتایج قابل اعتماد نخواهد بود.
۴. واریانس متغیر کمی در گروه‌ها باید هم‌اندازه باشد. آزمون لون در نرم‌افزار SPSS به همین منظور اعمال می‌شود.
۵. نمونه‌های گروه‌ها باید مستقل از یکدیگر باشند به عبارتی دقیق‌تر نمرات افراد در یک گروه متاثر از نمرات اعضای گروه دیگر نباشد.
۶. در صورتی که داده‌ها فاقد مفروضات فوق باشند باید از آزمون غیرپارامتریک H کراسکال والیس استفاده نمود.

درس سیزدهم - تحلیل واریانس یک طرفه □ ۱۷۵

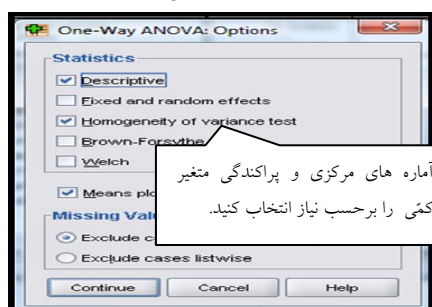
برای انجام این آزمون باید به شرح تصویر ذیل عمل شود.

شکل شماره (۶۸) شیوه اجرای آزمون تحلیل واریانس یک طرفه (مرحله ۱)



سپس با انتخاب گزینه option به شرح تصویر زیر عمل کنید. با انتخاب گزینه Descriptive شاخص‌های مرکزی و پراکندگی متغیر کمی در هر یک از طبقات محاسبه می‌شود و گزینه Homogeneity of variance test هم نتایج مربوط به پیش‌فرض برابری واریانس درون گروهی را در اختیار قرار می‌دهد.

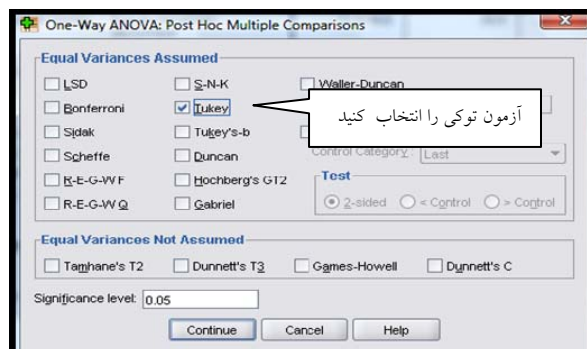
شکل شماره (۶۹) شیوه اجرای آزمون تحلیل واریانس یک طرفه (مرحله ۲)



اجرای دستور تا این مرحله، صرفاً نشان می‌دهد که آیا میانگین متغیر مورد نظر در گروه‌های مورد سنجش تفاوتی دارد یا خیر، حال در صورتی که نتیجه آزمون معنی‌دار باشد

(قابل تعمیم به جامعه آماری) و اگر بخواهیم دریابیم که تفاوت‌های معنادار مربوط به کدامیک از گروه‌هاست، باید از آزمون‌های تعقیبی که در تصویر آمده است، استفاده کرد.

شکل شماره (۷۰) شیوه اجرای آزمون تحلیل واریانس یک‌طرفه (مرحله ۳)



آزمون‌های تعقیبی فوق را می‌توان براساس هدف آماری به سه دسته تقسیم نمود:

دسته نخست شامل آزمون‌هایی است که هدفشان مقایسه دو به دوی میانگین تمامی گروه‌هاست. دسته دوم شامل آزمون‌هایی است که گروه‌های مورد بررسی را از حیث متغیرهای کمی مورد مقایسه، به زیرمجموعه‌های همگن تقسیم می‌کنند. ملاک همگنی در این بخش از آزمون‌ها آن است که بین گروه‌های یک زیر مجموعه تفاوت معناداری وجود نداشته باشد. دسته سوم نیز آزمون‌هایی است که هر دو کار فوق را انجام می‌دهند. یعنی هم به مقایسه دو به دو گروه‌ها می‌پردازند و هم آنها را به زیرمجموعه‌های همگن تقسیم می‌کنند. لازم به ذکر است که پیش‌فرض مهم تمامی آزمون‌های تعقیبی، برابری واریانس‌های درون گروهی گروه‌های مورد مقایسه است.

درس سیزدهم – تحلیل واریانس یک طرفه □ ۱۷۷

جدول زیر دسته‌بندی فوق را همراه با برخی ویژگی‌های این آزمون‌ها نشان می‌دهد.
جدول شماره (۵۵) ویژگی انواع آزمون‌های تعقیبی

آزمون تعقیبی	هدف آماری	محدودیت
LSD	مقایسه دو به دو گروه‌ها	تعداد گروه‌ها بیشتر از ۳ نباشد.
SIDAK	مقایسه دو به دو گروه‌ها	ندارد
BONFERRONI	مقایسه دو به دو گروه‌ها	توان این آزمون در تمایز گروه‌ها کم
R-E-G-W Q	شناسایی زیرمجموعه‌های همگن	ندارد
SNK	شناسایی زیرمجموعه‌های همگن	ندارد
BTUKEY	شناسایی زیرمجموعه‌های همگن	ندارد
DUNCAN	شناسایی زیرمجموعه‌های همگن	ندارد
WALLER	شناسایی زیرمجموعه‌های همگن	حجم نمونه گروه‌ها باید برابر باشد
TUKEY	مقایسه دو به دو و شناسایی زیرمجموعه‌های همگن	حجم نمونه برابر و توان تمایز کم
SCHEFFE	مقایسه دو به دو و شناسایی زیرمجموعه‌های همگن	ندارد
GT2	مقایسه دو به دو و شناسایی زیرمجموعه‌های همگن	حجم نمونه نابرابر باشد.
GABRIEL	مقایسه دو به دو و شناسایی زیرمجموعه‌های همگن	ندارد

اگر قصد دارید به این سؤال پاسخ دهد که "آیا حقوق فعلی کارکنان طبقات مختلف شغلی (کادر دفتری، نگهبانان و مدیران) با یکدیگر تفاوت معناداری دارد؟"، با اجرای دستور One-way Anova با خروجی های ذیل روبرو می شوید.

جدول شماره (۵۶) توصیف متغیر "حقوق فعلی کارکنان در طبقات شغلی مختلف"

Descriptives								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
دفتری	363	\$27,838.54	\$7,567.995	\$397.217	\$27,057.40	\$28,619.68	\$15,750	\$80,000
نگهبان	27	\$30,938.89	\$2,114.616	\$406.958	\$30,102.37	\$31,775.40	\$24,300	\$35,250
مدیر	84	\$63,977.80	\$18,244.776	\$1,990.668	\$60,018.44	\$67,937.16	\$34,410	\$135,000
Total	474	\$34,419.57	\$17,075.661	\$784.311	\$32,878.40	\$35,960.73	\$15,750	\$135,000

همانگونه که ملاحظه می‌شود میانگین حقوق کارمندان در طبقات مختلف شغلی بایکدیگر متفاوت است. به طوریکه بالاترین میانگین حقوق مربوط به مدیران بوده (63,977 دلار) و کارمندان دفتری نیز پایین‌ترین میانگین حقوق را دارا هستند (27,838 دلار). شاخص انحراف معیار نیز نشان می‌دهد که حقوق فعلی مدیران به طور متوسط ۱۸ دلار از میانگین فاصله دارد، در حالیکه در رده کارکنان دفتری این میزان بسیار کمتر (۷ دلار) است. حداقل و حداکثر حقوق این کارمندان نیز موید نتایج فوق است. ملاحظه می‌شود که حداقل حقوق در بین کارمندان دفتری ۱۵ دلار و حداکثر ۸۰ دلار است. در حالیکه در مورد کارمندان رده مدیریتی حداقل حقوق ۳۴ دلار و حداکثر آن ۱۳۵ دلار است.

جدول شماره (۵۷) آزمون مقایسه واریانس حقوق کارمندان در رده‌های شغلی سه‌گانه (دفتری، نگهبان، مدیر)

آزمون برابری واریانس حقوق سه طبقه شغلی (دفتری، نگهبان، مدیر)			
Test of Homogeneity of Variances			
Current Salary			
Levene Statistic	df1	df2	Sig.
59.733	2	471	.000

درس سیزدهم - تحلیل واریانس یک طرفه □ ۱۷۹

نتایج آزمون لون (Leven's test) حاکی از آن است که واریانس درون گروهی طبقات شغلی سه‌گانه تفاوت معناداری با یکدیگر دارد (Levene Statistic =59.733, Sig=.000). لذا استفاده از آزمون تحلیل واریانس مجاز نیست و باید از آزمون‌های ناپارامتریک مرتبط (نظیر تحلیل واریانس رتبه‌ای کراسکال والیس^۱) استفاده شود. با این حال نظر به آموزشی بودن این مطالب، با اندکی مسامحه، بقیه نتایج را تفسیر می‌کنیم.

جدول شماره (۵۸) آزمون مقایسه میانگین حقوق کارمندان در رده‌های شغلی سه‌گانه (دفتری، نگهبان، مدیر)

ANOVA					
Current Salary					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	8.944E10	2	4.472E10	434.481	.000
Within Groups	4.848E10	471	1.029E8		
Total	1.379E11	473			

نتایج مندرج در جدول فوق نشان می‌دهد تفاوت‌های مشاهده شده در میانگین حقوق کارمندان طبقات شغلی مختلف (مدیر، نگهبان، دفتری) قابل تعمیم به جامعه آماری است (F=434.481, Sig=.000).

حال در اینجا برای تشخیص این امر که تفاوت بین کدام یک از طبقات شغلی (دفتری، نگهبان، مدیر) به لحاظ آماری معنادار است، به نتیجه حاصل از آزمون تعقیبی توکی می‌پردازیم.

جدول شماره (۵۹) آزمون تعقیبی توکی برای بررسی تفاوت های زوجی طبقات سه گانه شغلی

Multiple Comparisons						
(I) Employment Category	(J) Employment Category	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Tukey HSD	
					Lower Bound	Upper Bound
Clerical	Custodial دفتری	\$-3,100.349	\$2,023.760	.277	\$-7,858.50	\$1,657.80
	Manager	\$-36,139.258*	\$1,228.352	.000	\$-39,027.29	\$-33,251.22
Custodial	Clerical نگهبان	\$3,100.349	\$2,023.760	.277	\$-1,657.80	\$7,858.50
	Manager	\$-33,038.909*	\$2,244.409	.000	\$-38,315.84	\$-27,761.98
Manager	Clerical مدیر	\$36,139.258*	\$1,228.352	.000	\$33,251.22	\$39,027.29
	Custodial	\$33,038.909*	\$2,244.409	.000	\$27,761.98	\$38,315.84

*. The mean difference is significant at the 0.05 level.

همانگونه که ملاحظه می شود تفاوت موجود در میانگین حقوق سه رده شغلی (دفتری، نگهبان، مدیر) به صورت دو به دو مورد مقایسه قرار گرفته شده است تا بدین وسیله تفاوت ها را به صورت خاص شناسایی کند. نتایج بدست آمده بیانگر آن است که تفاوت موجود در میانگین حقوق، بین کارمندان رده مدیریت با کارمندان دفتری و نگهبان است. به طوریکه تفاوت میانگین حقوق مدیران با کارکنان دفتری در حدود ۳۶ دلار و با نگهبانان در حدود ۳۳ دلار است (Sig=.000).

جدول شماره (۶۰) دسته بندی گروه های مورد مقایسه براساس آزمون توکی

Current Salary

Tukey HSD^{a, b}

Employment Category	N	Subset for alpha = 0.05	
		1	2
Clerical	363	\$27,838.54	\$63,977.80
Custodial	27	\$30,938.89	
Manager	84		
Sig.		.227	1.000

Means for groups in homogeneous subsets are displayed.

a. Uses Harmonic Mean Sample Size = 58.031.

Current Salary			
Tukey HSD ^{a,b}			
Employment Category	N	Subset for alpha = 0.05	
		1	2
Clerical	363	\$27,838.54	\$63,977.80
Custodial	27	\$30,938.89	
Manager	84		
Sig.		.227	1.000

Means for groups in homogeneous subsets are displayed.

a. Uses Harmonic Mean Sample Size = 58.031.

b. The group sizes are unequal. The harmonic mean of the group sizes is used. Type I error levels are not guaranteed.

این خروجی که ادامه آزمون تعقیبی توکی است بیانگر آن است که گروه‌های سه‌گانه مورد بررسی از حیث درآمد فعلی، دو دسته متمایز را تشکیل می‌دهند. دسته نخست که شامل کارمندان دفتری و نگهبانان است از لحاظ درآمد فعلی تفاوت معناداری با یکدیگر ندارند و یک طبقه واحد را تشکیل می‌دهند و تنها مدیران هستند که در یک طبقه جداگانه قرار می‌گیرند.

تمرین درس سیزدهم

- ۱- جداول زیر نتایج مربوط به مقایسه مدت ورزش (بر حسب دقیقه در هفته) در مناطق مختلف شهر تهران را نشان می دهد. براساس محتوای جداول به سؤالات زیر پاسخ دهید.
- الف- متوسط مدت ورزش در کدام منطقه بیشتر است؟ با توجه به میانگین و انحراف معیار توصیف کنید.
- ب- متوسط مدت ورزش در کدام منطقه کمتر است؟ با توجه به میانگین و انحراف معیار توصیف کنید.
- ج- آیا تفاوت های مشاهده شده قابلیت تعمیم به جامعه آماری را دارد؟
- د- آیا واریانس درون گروهی مناطق با یکدیگر برابر است؟
- ه- آیا براساس نتایج آزمون های تعقیبی، مناطق پنج گانه قابل تقلیل به دسته های کمتری نیز هست؟

Descriptives

مدت ورزش

	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
شمال	240	225.0250	267.21821	17.24886	191.0458	259.0042	.00	1050.00
شرق	171	160.9474	291.15881	22.26547	116.9950	204.8998	.00	1680.00
غرب	126	218.6667	307.18404	27.36613	164.5057	272.8276	.00	1770.00
مرکز	312	185.7212	223.21353	12.63698	160.8564	210.5859	.00	960.00
جنوب	357	216.1261	377.85684	19.99830	176.7964	255.4557	.00	2280.00
Total	1206	202.4726	302.41117	8.70812	185.3879	219.5574	.00	2280.00

Test of Homogeneity of Variances

مدت ورزش

Levene Statistic	df1	df2	Sig.
3.936	4	1201	.004

ANOVA

مدت ورزش

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	604074.153	4	151018.538	1.655	.158
Within Groups	1.096E8	1201	91254.125		
Total	1.102E8	1205			

درس سیزدهم – تحلیل واریانس یک طرفه □ ۱۸۳

Multiple Comparisons

مدت ورزش

Tukey HSD

(I) مناطق شهری	(J) مناطق شهری	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
شمال	شرق	64.07763	30.23037	.212	-18.5095	146.6647
	غرب	6.35833	33.23349	1.000	-84.4331	97.1497
	مرکز	39.30385	25.93659	.553	-31.5530	110.1607
	جنوب	8.89895	25.21585	.997	-59.9889	77.7868
شرق	شمال	-64.07763	30.23037	.212	-146.6647	18.5095
	غرب	-57.71930	35.46673	.480	-154.6118	39.1732
	مرکز	-24.77379	28.74250	.911	-103.2961	53.7486
	جنوب	-55.17868	28.09384	.284	-131.9289	21.5716
غرب	شمال	-6.35833	33.23349	1.000	-97.1497	84.4331
	شرق	57.71930	35.46673	.480	-39.1732	154.6118
	مرکز	32.94551	31.88605	.840	-54.1648	120.0558
	جنوب	2.54062	31.30260	1.000	-82.9757	88.0570
مرکز	شمال	-39.30385	25.93659	.553	-110.1607	31.5530
	شرق	24.77379	28.74250	.911	-53.7486	103.2961
	غرب	-32.94551	31.88605	.840	-120.0558	54.1648
	جنوب	-30.40490	23.41143	.692	-94.3632	33.5534
جنوب	شمال	-8.89895	25.21585	.997	-77.7868	59.9889
	شرق	55.17868	28.09384	.284	-21.5716	131.9289
	غرب	-2.54062	31.30260	1.000	-88.0570	82.9757
	مرکز	30.40490	23.41143	.692	-33.5534	94.3632

مدت ورزش

Tukey HSD

مناطق شهری	N	Subset for alpha = 0.05
		1
شرق	171	160.9474
مرکز	312	185.7212
جنوب	357	216.1261
غرب	126	218.6667
شمال	240	225.0250
Sig.		.193

Multiple Comparisons

مدت ورزش

Tukey HSD

(I) مناطق شهری	(J) مناطق شهری	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
شمال	شرق	64.07763	30.23037	.212	-18.5095	146.6647
	غرب	6.35833	33.23349	1.000	-84.4331	97.1497
	مرکز	39.30385	25.93659	.553	-31.5530	110.1607
	جنوب	8.89895	25.21585	.997	-59.9889	77.7868
شرق	شمال	-64.07763	30.23037	.212	-146.6647	18.5095
	غرب	-57.71930	35.46673	.480	-154.6118	39.1732
	مرکز	-24.77379	28.74250	.911	-103.2961	53.7486
	جنوب	-55.17868	28.09384	.284	-131.9289	21.5716
غرب	شمال	-6.35833	33.23349	1.000	-97.1497	84.4331
	شرق	57.71930	35.46673	.480	-39.1732	154.6118
	مرکز	32.94551	31.88605	.840	-54.1648	120.0558
	جنوب	2.54062	31.30260	1.000	-82.9757	88.0570
مرکز	شمال	-39.30385	25.93659	.553	-110.1607	31.5530
	شرق	24.77379	28.74250	.911	-53.7486	103.2961
	غرب	-32.94551	31.88605	.840	-120.0558	54.1648
	جنوب	-30.40490	23.41143	.692	-94.3632	33.5534
جنوب	شمال	-8.89895	25.21585	.997	-77.7868	59.9889
	شرق	55.17868	28.09384	.284	-21.5716	131.9289
	غرب	-2.54062	31.30260	1.000	-88.0570	82.9757

Means for groups in homogeneous subsets are displayed.

۲- براساس داده‌های فایل world95.sav آیا تفاوت معناداری بین درصد باسوادی (literacy) کشورهای خاورمیانه، آفریقا و آمریکای لاتین (متغیر منطقه اقتصادی Region) وجود دارد؟

۳- براساس داده‌های فایل world95.sav آیا تفاوت معناداری بین درصد شهرنشینی کشورهای متعلق به مناطق آب و هوایی مختلف (متغیر Climate) وجود دارد؟

درس چهاردهم

آزمون‌های مناسب برای

مقایسه متغیرهای رتبه‌ای

در پایان این درس انتظار می‌رود دانشجو

- ۱- با وضعیت‌های مختلف مقایسه متغیرهای رتبه‌ای آشنا باشد.
- ۲- کاربرد آزمون یومان ویتنی و کراسکال والیس را بداند.
- ۳- کاربرد آزمون فریدمن و ویلکاکسون را بداند.
- ۴- توانایی آزمون فرضیات متناسب با یومان ویتنی یا کراسکال والیس را داشته باشد.
- ۵- توانایی آزمون فرضیات متناسب با فریدمن و ویلکاکسون را داشته باشد.
- ۶- بتواند نتایج هر یک از آزمون‌های فوق را تفسیر کند.

۱- مقدمه

همانطور که در درس‌های گذشته اشاره شد، اگر در انجام یک تحقیق، قصد داشته باشید متغیرهای کیفی (از نوع رتبه‌ای) و یا متغیرهای کمی غیرنرمال را به شیوه‌های مختلفی مورد مقایسه قرار دهید، باید از مجموعه آزمون‌های مرتبط با مقایسه متغیرهای رتبه‌ای برحسب نیاز استفاده کنید. این آزمون‌ها عبارتست از:

جدول شماره (۶۱) آزمون‌های مربوط به مقایسه متغیرهای رتبه‌ای از حیث هدف تحقیق

ردیف	نام آزمون	هدف تحقیق	مثال
۱	ل‌مان ویتنی	مقایسه یک متغیر کیفی رتبه‌ای یا کمی غیرنرمال در دو گروه	آیا رضایت مشتریان (زنان و مردان) از کیفیت خدمات ارائه شده توسط بانک متفاوت است؟
۲	ویلکاکسون	مقایسه دو متغیر کیفی رتبه‌ای یا کمی غیرنرمال با یکدیگر	آیا گرایش جوانان به موسیقی کلاسیک و رپ با یکدیگر متفاوت است؟
۳	کراسکال والیس	مقایسه یک متغیر کیفی رتبه‌ای یا کمی غیرنرمال در بیش از دو گروه	آیا ارزیابی افراد گروه‌های سنی مختلف از موسیقی کلاسیک متفاوت است؟
۴	فریدمن	مقایسه چند متغیر کیفی رتبه‌ای یا کمی غیرنرمال در یک گروه	آیا رضایت مشتریان از قیمت کالا و تنوع محصولات و خدمات موجود در فروشگاه مواد غذایی متفاوت است؟

به این ترتیب باید با توجه به فرضیه تحقیق، از موارد چهارگانه فوق استفاده کرد. در اینجا به تفکیک هر یک از مثال‌های فوق مورد آزمون قرار گرفته و توضیحات لازم در پی می‌آید.

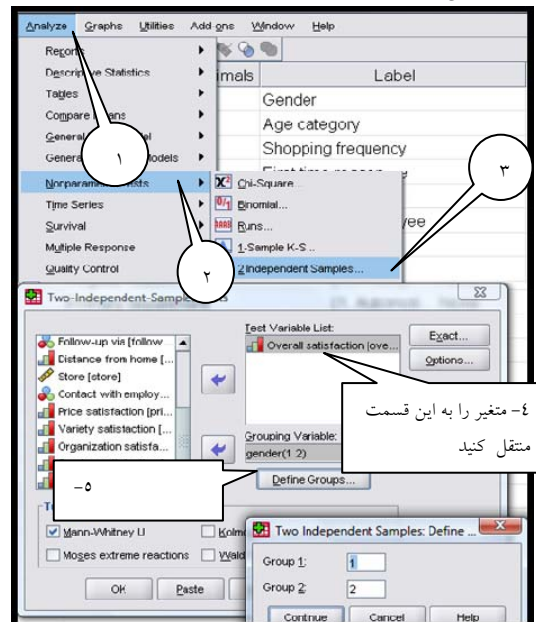
۲- مقایسه یک متغیر کیفی رتبه‌ای یا کمی غیرنرمال در دو گروه

• آزمون یومان ویتنی

اگر بخواهید یک متغیر کیفی رتبه‌ای یا کمی غیرنرمال را در دو گروه مقایسه کنید، باید از آزمون یومان ویتنی استفاده کنید. در این آزمون افراد بر حسب نمره‌ای که در متغیر رتبه‌ای دارند به شیوه صعودی مرتب می‌شوند، سپس آنها را جدا ساخته و متوسط رتبه‌های دو گروه از لحاظ متغیر مورد بررسی با یکدیگر مقایسه می‌شود. برای انجام این آزمون باید به شرح تصویر ذیل عمل شود.

درس چهاردهم - روشهای مناسب برای مقایسه متغیرهای رتبه‌ای □ ۱۸۷

شکل شماره (۷۱) شیوه اجرای آزمون یومان ویتنی



مثلاً اگر بخواهید این فرضیه را آزمون کنید که "آیا رضایت مشتریان (زنان و مردان) از کیفیت خدمات ارائه شده توسط بانک متفاوت است؟" با انجام این آزمون به شیوه فوق، با جداول خروجی زیر مواجه خواهید شد.

جدول شماره (۶۲) نتایج آزمون یومان ویتنی درباره مقایسه رضایت مشتریان به تفکیک جنس

Ranks			میانگین رتبه‌های دو گروه	
جنسیت	N	Mean Rank	Sum of Ranks	
رضایت مردان	209	295.84	61831.50	
رضایت زنان	373	289.07	107821.50	
Total	582			

همانطور که ملاحظه می‌شود از کل نمونه ۵۸۲ نفری، ۳۷۳ نفر را زنان و ۲۰۹ نفر را مردان تشکیل می‌دهند. مقایسه متوسط رتبه‌های زنان و مردان نشان می‌دهد که رضایت مردان ($\text{Mean Rank} = 295/84$) از زنان ($\text{Mean Rank} = 289/07$) بیشتر است.

Test Statistics ^a	
	REZAYAT
Mann-Whitney U	38070.500
Wilcoxon W	107821.5
Z	-.467
Asymp. Sig. (2-tailed)	.640

a. Grouping Variable: جنسیت

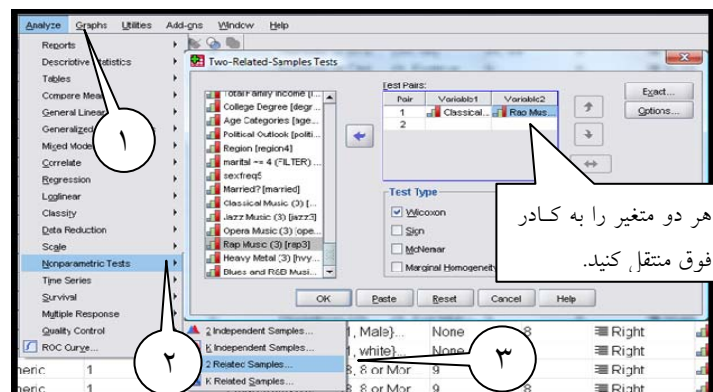
با توجه به میزان خطای بدست آمده ($z = -0.467$, $\text{sig} = 0.640$) می‌توان استنباط کرد که مشتریان (اعم از زن و مرد) از خدمات ارائه شده توسط بانک رضایت یکسانی دارند و تفاوت محسوسی بین آنها دیده نمی‌شود.

۳- مقایسه دو متغیر کیفی رتبه‌ای یا کمی غیرنرمال با یکدیگر

• آزمون ویلکاکسون

اگر قصد مقایسه دو متغیر کیفی رتبه‌ای یا کمی غیرنرمال را با یکدیگر داشته باشید، باید از آزمون ویلکاکسون استفاده کنید. برای انجام این آزمون باید به شرح تصویر ذیل عمل شود.

شکل شماره (۷۲) شیوه اجرای آزمون ویلکاکسون



درس چهاردهم - روشهای مناسب برای مقایسه متغیرهای رتبه‌ای □ ۱۸۹

در صورتیکه قصد سنجش این سؤال را داشته باشید که "آیا گرایش جوانان به موسیقی کلاسیک و رپ بایکدیگر متفاوت است؟" با انجام این آزمون جداول خروجی زیر را مشاهده خواهید کرد.

جدول شماره (۶۳) نتایج آزمون ویلکاکسون در زمینه مقایسه گرایش جوانان به موسیقی کلاسیک و رپ

Ranks				
		N	Mean Rank	Sum of Ranks
رپ - کلاسیک	Negative Ranks	164 ^a	405.06	66430.50
	Positive Ranks	926 ^b	570.37	528164.50
	Ties	293 ^c		
	Total	1383		

a. Rap Music < Classical Music
b. Rap Music > Classical Music
c. Rap Music = Classical Music

تعیین جهت گرایش جوانان به دو نوع موسیقی رپ و کلاسیک

همانگونه که ملاحظه می‌شود از ۱۳۸۳ نفر نمونه مورد بررسی، ۱۶۴ نفر به موسیقی کلاسیک بیشتر از موسیقی رپ گرایش دارند همچنین بالغ بر ۹۲۶ نفر گرایش‌شان به موسیقی رپ بیشتر از کلاسیک است. ضمناً از کل نمونه مورد بررسی ۲۹۳ نفر نیز هستند که گرایش آنها به هر دو نوع موسیقی (رپ و کلاسیک) مساوی است.

Test Statistics ^b	
	رپ - کلاسیک
Z	-22.447 ^a
Asymp. Sig. (2-tailed)	.000

a. Based on negative ranks.
b. Wilcoxon Signed Ranks Test

همانگونه که ملاحظه می‌شود با توجه به میزان خطای بدست آمده ($\text{sig} = ۰/۰۰۰$)، می‌توان ادعا کرد که در سطح جامعه آماری نیز با اطمینان بیش از ۹۹ درصد جوانان گرایش‌های متفاوتی نسبت به موسیقی کلاسیک و رپ دارند.

۱- همانطور که ملاحظه می‌شود در اینجا رتبه‌های افراد یک گروه در دو متغیر مقایسه می‌شود. به عبارتی دیگر هر فرد دارای دو رتبه است: رتبه در موسیقی رپ، رتبه در موسیقی کلاسیک.

۴- مقایسه یک متغیر کیفی رتبه‌ای یا کمّی غیرنرمال در بیش از دو گروه

• آزمون کراسکال والیس

اگر به یاد بیاورید در درس‌های قبل گفته شد که اگر محققى بخواهد میانگین یک متغیر کمّی را در بیش از دو گروه مقایسه کند باید از آزمون تحلیل واریانس یک طرفه استفاده کند، حال اگر این متغیر رتبه‌ای یا کمّی غیرنرمال باشد و بخواهد آن را در بیش از دو گروه مورد مقایسه قرار دهد، باید از تحلیل واریانس رتبه‌ای که عنوان آماری آن آزمون کراسکال والیس است استفاده شود. برای انجام این آزمون باید به شرح تصویر ذیل عمل شود.

شکل شماره (۷۳) شیوه اجرای آزمون کراسکال والیس



اگر قصد پاسخ به این سؤال را داشته باشید که "آیا ارزیابی مردم در سنین مختلف از موسیقی کلاسیک متفاوت است؟ با انجام این آزمون، جداول خروجی زیر مشاهده خواهد شد. جدول شماره (۶۴) نتایج آزمون کراسکال- والیس در زمینه مقایسه ارزیابی گروه‌های سنی مختلف از موسیقی کلاسیک

Ranks			
	گروه‌های سنی	N	Mean Rank
موسیقی کلاسیک	18-29	267	766.39
	30-39	334	728.39
	40-49	295	684.35
	50+	529	692.32
	Total	1425	

درس چهاردهم - روشهای مناسب برای مقایسه متغیرهای رتبه‌ای □ ۱۹۱

همانگونه که در جدول فوق مشاهده می‌شود در بین طبقات سنی مختلف، گروه سنی ۱۸ تا ۲۹ سال، ارزیابی مثبت‌تری از موسیقی کلاسیک دارند ($\text{Mean Rank} = ۷۶۶/۳۹$) منفی‌ترین ارزیابی نیز به گروه سنی ۴۰ تا ۴۹ سال تعلق دارد. ($\text{Mean Rank} = ۶۸۴/۳۵$)

Test Statistics ^{a,b}	
	موسیقی کلاسیک
Chi-Square	8.197
df	3
Asymp. Sig.	.042
a. Kruskal Wallis Test	
b. Grouping Variable: Age Categories	

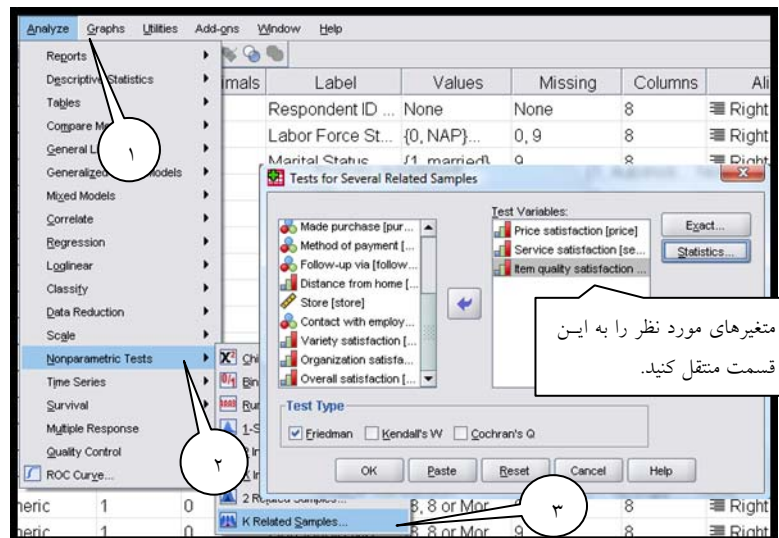
همانگونه که ملاحظه می‌شود با توجه به میزان خطای بدست آمده می‌توان با ۹۵ درصد اطمینان ادعا کرد که همانند نمونه مورد بررسی، اعضای جامعه آماری نیز در سنین مختلف ارزیابی‌های متفاوتی از موسیقی کلاسیک دارند. ($\text{chi-square} = 8.197, \text{sig} = 0.042$)

۵- مقایسه چند متغیر کیفی رتبه‌ای یا کمی غیرنرمال در یک گروه

• آزمون فریدمن

اگر قصد داشته باشید چندین متغیر کیفی رتبه‌ای یا کمی غیرنرمال را همزمان در یک گروه مقایسه کنید، باید از آزمون فریدمن استفاده کنید. در این آزمون متغیرها برای هر فرد از حیث نمره، رتبه‌بندی می‌شوند. سپس میانگین رتبه‌های هر متغیر محاسبه شده و با یکدیگر مقایسه می‌شود. برای انجام این آزمون باید به شرح تصویر ذیل عمل شود.

شکل شماره (۷۴) شیوه اجرای آزمون فریدمن



اگر قصد پاسخ به این سؤال را داشته باشید که "آیا رضایت مشتریان از قیمت کالا و تنوع محصولات و خدمات موجود در فروشگاه مواد غذایی متفاوت است؟" با انجام این آزمون جداول خروجی زیر مشاهده خواهد شد.

جدول شماره (۶۵) نتایج آزمون فریدمن در زمینه مقایسه رضایت مشتریان از قیمت و تنوع کالا و خدمات یک فروشگاه

Descriptive Statistics					
	N	Mean	Std. Deviation	Minimum	Maximum
رضایت از قیمت کالا	582	3.08	1.261	1	5
رضایت از تنوع محصولات	582	3.08	1.274	1	5
رضایت از خدمات	582	3.09	1.341	1	5

این جدول نشان می دهد که نمونه مورد بررسی ۵۸۲ نفر بوده است. با توجه به میانگین بدست آمده، می توان استنباط کرد که رضایت مشتریان از قیمت و تنوع محصولات و همچنین خدمات فروشگاه تقریباً مشابه است و از این حیث تفاوت های بسیار اندکی در نمونه مورد بررسی مشاهده می شود.

درس چهاردهم - روشهای مناسب برای مقایسه متغیرهای رتبه‌ای □ ۱۹۳

Ranks	
	Mean Rank
رضایت از قیمت کالا	2.01
رضایت از تنوع محصولات	2.00
رضایت از خدمات	1.99

این جدول میانگین رتبه‌های رضایت مشتریان را از قیمت کالا، تنوع محصولات و خدمات فروشگاه نشان می‌دهد. همانگونه که ملاحظه می‌شود رضایت مشتریان مورد بررسی از قیمت کالا در مقایسه با خدمات و تنوع محصولات بیشتر است.

Test Statistics ^a	
N	582
Chi-Square	.063
df	2
Asymp. Sig.	.969

a. Friedman Test

با توجه به میزان خطای بدست آمده ($\text{sig} = 0/969$ ، $\text{chi-square} = 0/063$) می‌توان ادعا کرد تفاوت مشاهده شده در رضایت مشتریان از قیمت کالا، تنوع محصولات و همچنین خدمات ارائه شده در فروشگاه ناچیز بوده و قابلیت تعمیم به جامعه آماری را ندارد.

تمرین درس چهاردهم

۱- در یک تحقیق پیمایشی نیازهای ورزشی شهروندان در پنج حیطه به شرح جدول زیر سنجیده شده است. آیا بین ابعاد پنج گانه نیاز براساس جداول زیر تفاوت معناداری وجود دارد؟ نتایج مربوطه را تفسیر کنید.

Descriptive Statistics

	N	Mean	Std. Deviation	Minimum	Maximum
بعد فضاهای ورزشی	1197	2.3017	.68629	1.00	4.00
بعد تجهیزات ورزشی	1197	2.2470	.60181	1.00	4.00
بعد فرهنگی اجتماعی	1197	2.1891	.67184	1.00	4.00
بعد بهداشت و سلامت	1197	2.2238	.72897	1.00	4.00
بعد اقتصاد ورزش	1197	2.6124	1.06354	1.00	4.00

Friedman Test

Ranks

	Mean Rank
بعد فضاهای ورزشی	2.96
بعد تجهیزات ورزشی	2.91
بعد فرهنگی اجتماعی	2.88
بعد بهداشت و سلامت	2.82
بعد اقتصاد ورزش	3.42

Test Statistics^a

N	1197
Chi-Square	115.493
df	4
Asymp. Sig.	.000

a. Friedman Test

۲- آیا نتایج مندرج در جداول زیر تفاوت معناداری را بین میزان تماشای تلویزیون در بین شهروندان زن و مرد نشان می دهد؟ نتایج حاصله را تفسیر کنید.

درس چهاردهم – روشهای مناسب برای مقایسه متغیرهای رتبه‌ای □ ۱۹۵

Ranks			
جنسیت	N	Mean Rank	Sum of Ranks
زن میزان تماشای تلویزیون	612	641.90	392845.50
مرد	582	550.81	320569.50
Total	1194		

Test Statistics ^a	
	میزان تماشای
Mann-Whitney U	150916.500
Wilcoxon W	320569.500
Z	-4.664
Asymp. Sig. (2-tailed)	.000

a. Grouping Variable: جنسیت

۳- آیا نتایج مندرج در جداول زیر تفاوت معناداری را بین میزان تحرک بدنی شاغلین مناطق مختلف شهر نشان می‌دهد؟ نتایج حاصله را تفسیر کنید.

Ranks		
مناطق شهری	N	Mean Rank
شمال میزان تحرک بدنی در شغل	78	222.15
شرق	60	221.52
غرب	48	141.12
مرکز	102	161.49
جنوب	96	208.91
Total	384	

Test Statistics^{a,b}

	میزان تحرک بدنی در شغل
Chi-Square	31.485
df	4
Asymp. Sig.	.000

a. Kruskal Wallis Test

b. Grouping Variable: مناطق شهری

۴- براساس داده‌های فایل GSS subset.sav،

الف- آیا بین تحصیلات پدر (pdeg) برحسب نژاد (Race) تفاوت معناداری وجود دارد؟

ب- آیا بین تحصیلات (Educ) زنان و مردان، تفاوت معناداری وجود دارد؟

ج- آیا بین تحصیلات پدران (pdeg) و مادران (madeg) اعضای نمونه تفاوت معناداری وجود دارد؟

درس پانزدهم

آزمون مناسب برای

تعیین جهت و شدت تأثیر متغیر مستقل بر وابسته

در پایان این درس انتظار می‌رود دانشجو

- ۱- موارد کاربرد رگرسیون چندگانه را شناسایی و مفروضات آن را بداند.
- ۲- بتواند نتایج مربوط به معناداری رگرسیون چندگانه را تفسیر کند.
- ۳- بتواند جهت و شدت ضرایب تأثیر متغیرهای مستقل را تفسیر کند.
- ۴- بتواند شدت تأثیر متغیرهای مستقل را با یکدیگر مقایسه کند.
- ۵- توانایی بررسی مفروضات رگرسیون چندگانه را داشته باشد.

۱- مقدمه

اگر در جریان تحقیقی بخواهید جهت و شدت تأثیر یک یا چند متغیر مستقل را بر یک یا چند متغیر متغیر وابسته بررسی کنید و نهایتاً به معادله‌ای برای پیش‌بینی متغیر وابسته برسید، باید از آزمون رگرسیون استفاده کنید. با توجه به تعداد متغیر وابسته، رگرسیون انواع مختلفی پیدا می‌کند که در اینجا به اختصار معرفی می‌شود.

- اگر تعداد متغیرهای وابسته در فرضیه بیش از یک مورد باشد، باید از رگرسیون چند متغیری^۱ استفاده کرد.

- اگر تعداد متغیرهای وابسته در فرضیه تنها یک مورد باشد، باید از رگرسیون چندگانه^۲ استفاده کرد.

رگرسیون چندگانه نیز براساس سطح سنجش متغیر وابسته، به انواع متعددی تقسیم می‌شود که عبارتند از:

الف- رگرسیون خطی^۳ و غیرخطی^۴ (ویژه متغیرهای وابسته‌ای که در سطح سنجش فاصله‌ای یا نسبی هستند)^۵

ب- رگرسیون لجستیک^۶ (ویژه متغیرهای وابسته‌ای که اسمی دو مقوله‌ای هستند).

ج- رگرسیون اسمی چند مقوله‌ای^۷ (ویژه متغیرهای وابسته‌ای که اسمی چند مقوله‌ای هستند).

د- رگرسیون رتبه‌ای^۸ (ویژه متغیرهای وابسته‌ای که سطح سنجش رتبه‌ای دارند).

در این بخش با توجه به اهداف تعیین شده کتاب، صرفاً رگرسیون چندگانه خطی توضیح داده می‌شود.^۹

1-Multivariate Regression

2-Multiple Regression

3-Linear

4-Nonlinear

۵- اگر چه این آزمون برای متغیرهای کمی (فاصله‌ای و نسبی) است، اما می‌توان با دو وجهی نمودن متغیرهای کیفی (اسمی و رتبه‌ای) و یا متغیرهایی که خود کیفی و دو وجهی هستند (مثل جنسیت)، از این آزمون استفاده کرد.

6-Logistic

7-Multinomial

8-Ordinal

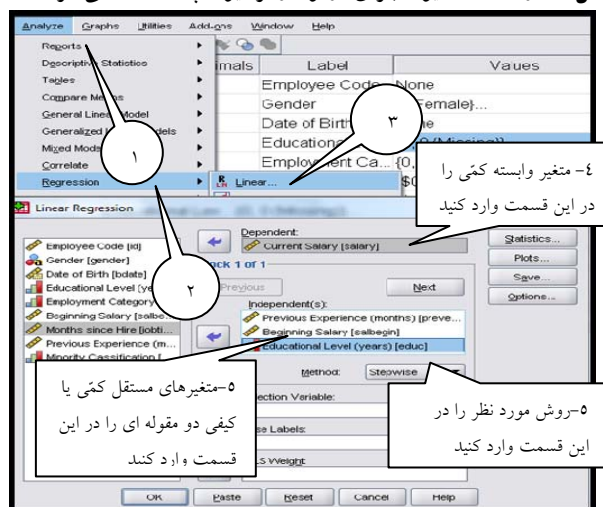
۱- مفروضات رگرسیون چندگانه خطی عبارتست از:

درس پانزدهم – آزمون مناسب برای تعیین جهت و شدت تأثیر متغیر مستقل بر وابسته □ ۱۹۹

۲- رگرسیون چندگانه خطی

اگر بخواهید به این سؤال پاسخ دهید که "حقوق بدو استخدام، تجربیات قبلی و سطح تحصیلات تا چه اندازه بر میزان حقوق فعلی کارمندان تأثیر دارد؟"، به دلایلی که پیش از این گفته شد، باید از رگرسیون چندگانه خطی به شیوه زیر استفاده کنید.

شکل شماره (۷۵) شیوه اجرای آزمون رگرسیون چندگانه خطی (مرحله ۱)



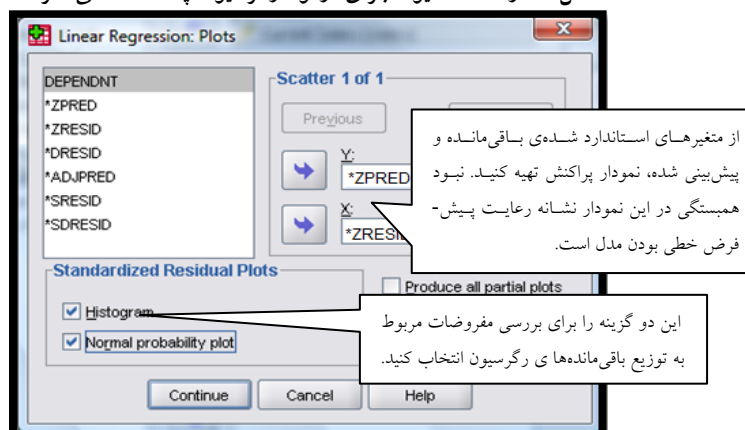
برای انجام رگرسیون چندگانه خطی، روش‌های مختلفی برای ورود متغیرهای مستقل به معادله وجود دارد که از آن جمله می‌توان به روش‌های جبری^۱، گام‌به‌گام^۲، پیش‌رونده^۳،

- متغیر یا متغیرهای مستقل با متغیر وابسته رابطه خطی داشته باشند.
- واریانس عبارت خطا در تمام موارد ثابت باشد.
- توزیع فراوانی مقادیر باقی‌مانده باید از شکل نرمال برخوردار باشد.
- باقی‌مانده‌ها نباید وابسته به هم باشند
- بین متغیرهای مستقل رابطه هم‌خطی وجود نداشته باشد.

- 1 - Enter
- 2 - Stepwise
- 3 - Forward

پس‌رونده^۱ اشاره کرد. در روش جبری تمامی متغیرهای مستقل وارد معادله می‌شود. در روش گام‌به‌گام در هر مرحله یکی از متغیرهای مستقل (بر حسب شدت تأثیر خالص و نیز معناداری) وارد معادله می‌شود. این کار تا ورود آخرین متغیری که تأثیر معنادار بر ضریب تعیین داشته باشد ادامه می‌یابد. در روش پیش‌رونده نیز همانند روش گام‌به‌گام ملاک ورود متغیر مستقل به معادله، داشتن بیشتر تأثیر معنادار است. با این تفاوت که در روش گام‌به‌گام اثرات خالص متغیرهای مستقل پس از ورود به معادله مجدداً مورد بررسی قرار می‌گیرد اما در روش پیش‌رونده این کار صورت نمی‌گیرد. اما مبنای ورود متغیرهای مستقل در روش حذف پس‌رونده چنان است که ابتدا تمامی متغیرهای مستقل وارد معادله می‌شوند، سپس اثرات خالص یا توانان تمامی آنها مورد بررسی قرار می‌گیرد. آنگاه متغیری که حذف آن کمترین کاهش غیرمعنادار را در ضریب تعیین داشته باشد، از معادله حذف می‌شود. این کار تا حذف آخرین متغیر ضعیف ادامه می‌یابد. برای بررسی مفروضات رگرسیون در زمینه توزیع باقی‌مانده‌ها، باید دو گزینه زیر را همانند تصویر زیر فعال کنید.

شکل شماره (۷۶) شیوه اجرای آزمون رگرسیون چندگانه خطی (مرحله ۲)



بعد از آن آماره دوربین و اتسون را برای بررسی پیش‌فرض استقلال باقی‌مانده‌ها مطابق تصویر زیر انتخاب کنید.

درس پانزدهم - آزمون مناسب برای تعیین جهت و شدت تأثیر متغیر مستقل بر وابسته □ ۲۰۱

شکل شماره (۷۷) شیوه اجرای آزمون رگرسیون چندگانه خطی (مرحله ۳)



جدول شماره (۶۶) معرفی متغیرهای مستقل و وابسته و روش بکارگرفته شده در انجام رگرسیون

Variables Entered/Removed ^a			
Model	Variables Entered	Variables Removed	Method
1	Beginning Salary	.	Stepwise (Criteria: Probability-of-F-to-enter <= .050, Probability-of-F-to-remove >= .100).
2	Previous Experience (months)	.	Stepwise (Criteria: Probability-of-F-to-enter <= .050, Probability-of-F-to-remove >= .100).
3	Educational Level (years)	.	Stepwise (Criteria: Probability-of-F-to-enter <= .050, Probability-of-F-to-remove >= .100).

a. Dependent Variable: Current Salary

همانگونه که ملاحظه می‌شود این جدول متغیرهای مستقل وارد شده به مدل رگرسیون را در مراحل مختلف نشان می‌دهد. متغیر وابسته، "حقوق فعلی" است و متغیرهای مستقل به ترتیب ورود به معادله عبارتند از: "حقوق بدو استخدام"، "تجربیات قبلی"، "سطح تحصیلات". ستون چهارم روش انجام رگرسیون را به شیوه گام‌به‌گام (Stepwise) معرفی می‌کند.^۱

۱- در این روش کلیه متغیرهای مستقل بیرون از معادله قرار دارند و هر یک که تأثیر بیشتر و معنادارتری داشته باشند، وارد معادله می‌شود.

جدول شماره (۶۷) ضرایب همبستگی و تعیین در گام های مختلف

Model Summary ^a					
Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.880 ^a	.775	.774	\$8,115.356	
2	.891 ^b	.793	.793	\$7,776.652	
3	.895 ^c	.802	.800	\$7,631.917	1.832

a. Predictors: (Constant), Beginning Salary
b. Predictors: (Constant), Beginning Salary, Previous Experience (months)
c. Predictors: (Constant), Beginning Salary, Previous Experience (months), Educational Level (years)

مقایسه گام به گام نتایج جدول فوق نشان می دهد عملیات رگرسیون تا سه مرحله انجام شده است و میزان ضریب همبستگی چندگانه و ضریب تعیین در هر مرحله نسبت به مرحله قبل افزایش معناداری داشته و خطای برآورد در مراحل پایانی در مقایسه با مرحله اول کاهش معنادار داشته است. بررسی متغیرهای مستقل نشان می دهد که هر سه متغیر مورد بررسی (حقوق بدو استخدام، تجربیات قبلی و سطح تحصیلات) نقش تعیین کننده ای در تبیین متغیر وابسته (حقوق فعلی) دارند. همانگونه که ملاحظه می شود این سه متغیر به میزان (۰/۸۹ = r^2) با متغیر وابسته همبستگی دارند. ضریب تعیین (R^2 square) حاصله، نیز بیانگر آن است که ۸۰ درصد از تغییرات موجود در حقوق فعلی کارمندان، تابع متغیرهای مستقل مزبور است. از آنجا که مقدار آماره دوربین- واتسون به ۲ نزدیک است می توان اطمینان داشت که پیش فرض استقلال باقی مانده ها نیز در تحلیل رگرسیونی حاضر رعایت شده است.

درس پانزدهم - آزمون مناسب برای تعیین جهت و شدت تأثیر متغیر مستقل بر وابسته □ ۲۰۳

جدول شماره (۶۸) نتایج آزمون معناداری رگرسیون در گام‌های مختلف

ANOVA ^a					
Model	Sum of Squares	df	Mean Square	F	Sig.
1 Regression	1.068E11	1	1.068E11	1622.118	.000 ^a
Residual	3.109E10	472	6.586E7		
Total	1.379E11	473			
2 Regression	1.094E11	2	5.472E10	904.752	.000 ^b
Residual	2.848E10	471	6.048E7		
Total	1.379E11	473			
3 Regression	1.105E11	3	3.685E10	632.607	.000 ^c
Residual	2.738E10	470	5.825E7		
Total	1.379E11	473			

a. Predictors: (Constant), Beginning Salary
b. Predictors: (Constant), Beginning Salary, Previous Experience (months)
c. Predictors: (Constant), Beginning Salary, Previous Experience (months), Educational Level (years)
d. Dependent Variable: Current Salary

جدول فوق به بررسی معناداری مدل رگرسیونی به تفکیک مراحل سه‌گانه پرداخته است. همانطور که ملاحظه می‌شود نسبت میانگین مجذورات رگرسیون بر باقی‌مانده، در هر سه مرحله به طرز قابل توجهی بالاست ($F=632.607$, $\text{sig}=0.000$). نگاهی به سطح معناداری مدل‌ها نیز حاکی از این واقعیت است که مدل‌های رگرسیونی به دست آمده در هر سه مرحله به طرز معناداری قابلیت پیش‌بینی متغیر وابسته را دارد.

جدول شماره (۶۹) متغیرهای وارد شده در معادله رگرسیون و ضرایب تأثیر آنها در گام‌های

مختلف

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	1928.206	888.680		2.170	.031
Beginning Salary	1.909	.047	.880	40.276	.000
2 (Constant)	3850.718	900.633		4.276	.000
Beginning Salary	1.923	.045	.886	42.283	.000
Previous Experience (months)	-22.445	3.422	-.137	-6.558	.000
3 (Constant)	-3661.517	1935.490		-1.892	.059
Beginning Salary	1.749	.060	.806	29.198	.000
Previous Experience (months)	-16.730	3.605	-.102	-4.641	.000
Educational Level (years)	735.956	168.689	.124	4.363	.000

Dependent Variable: Current Salary

نتایج مندرج در جدول بالا شدت و جهت تأثیر متغیرهای مستقل را بر وابسته نشان می‌دهد. همانگونه که ملاحظه می‌شود کلیه متغیرهای مستقل (حقوق بدو استخدام، تجربیات قبلی و سطح تحصیلات) در مرحله نهایی (گام سوم)، وارد معادله رگرسیون شده‌اند. معادله رگرسیون برای پیش‌بینی متغیر وابسته، براساس ضرایب استاندارد نشده به شرح زیر است:

$$Y' = -3661.517 + 1.749 (\text{Beginning Salary}) - 16.73 (\text{Previous Experience (months)}) + 735.956 (\text{Educational Level (years)})$$

به عنوان مثال می‌توان پیش‌بینی کرد که اگر فردی حقوق بدو استخدامش ۱۲۰۰۰ دلار باشد و سابقه شغلی ۳۶ ماهه داشته باشد و دارای تحصیلاتی در حد لیسانس (۱۶ سال تحصیل موفقیت‌آمیز) باشد، با توجه به معادله فوق درآمد فعلی او بالغ بر ۲۹۷۰۴.۰۵۹ دلار خواهد شد:

$$Y' = -3661.517 + 1.749 (12000) - 16.73 (36) + 735.956 (16) = 29704.059$$

درس پانزدهم - آزمون مناسب برای تعیین جهت و شدت تأثیر متغیر مستقل بر وابسته □ ۲۰۵

معادله رگرسیون را بر حسب ضرایب استاندارد شده β نیز می‌توان نوشت، با این تفاوت که برای پیش‌بینی نمره استاندارد متغیر وابسته، باید نمره استاندارد متغیرهای مستقل را در این معادله قرار داد.

$$Z_Y = 0.806 (Z_{\text{Beginning Salary}}) - 0.102 (Z_{\text{Previous Experience (months)}}) + 0.124 (Z_{\text{Educational Level (years)}})$$

رگرسیون استاندارد صفر است. علاوه بر این ضرایب استاندارد شده β امکان دیگری نیز در اختیار محقق قرار می‌دهد و آن قابلیت مقایسه شدت تأثیر متغیرهای مستقل است. همانطور که ملاحظه می‌شود حقوق بدو استخدام بیشترین تأثیر و سابقه خدمت به ماه کمترین تأثیر را بر حقوق فعلی افراد می‌گذارد. بررسی جهت تأثیر متغیرهای مستقل نشان می‌دهد که با حقوق فعلی افراد، در بین کسانی که حقوق بدو استخدام $(\beta = 0.806)$ شان بیشتر بوده و از سطح تحصیلات $(\beta = 0.124)$ بالاتری برخوردارند بیشتر است و در مقابل سابقه خدمت $(\beta = 0.102)$ تأثیر معکوس بر درآمد فعلی افراد دارد.

جدول شماره (۷۰) متغیرهای خارج شده از معادله رگرسیون

Excluded Variables ^c					
Model	Beta In	t	Sig.	Partial Correlation	Collinearity Statistics Tolerance
1 Previous Experience (months)	-.137 ^a	-6.558	.000	-.289	.998
Educational Level (years)	.172 ^a	6.356	.000	.281	.599
2 Educational Level (years)	.124 ^b	4.363	.000	.197	.520
a. Predictors in the Model: (Constant), Beginning Salary					
b. Predictors in the Model: (Constant), Beginning Salary, Previous Experience (months)					
c. Dependent Variable: Current Salary					

جدول فوق متغیرهای مستقل خارج از معادله رگرسیون را به تفکیک گام‌های مختلف نشان می‌دهد. همانگونه که ملاحظه می‌شود در مرحله اول دو متغیر (تجربیات قبلی و سطح تحصیلات) و در مرحله دوم یک متغیر (سطح تحصیلات) در خارج از معادله قرار دارند (متغیر تجربیات قبلی در گام دوم وارد معادله شده است) اما در مرحله سوم هیچ متغیری نیست که خارج از معادله باقی‌مانده باشد.

جدول شماره (۷۱) آماره‌های مربوط به باقیمانده‌ها

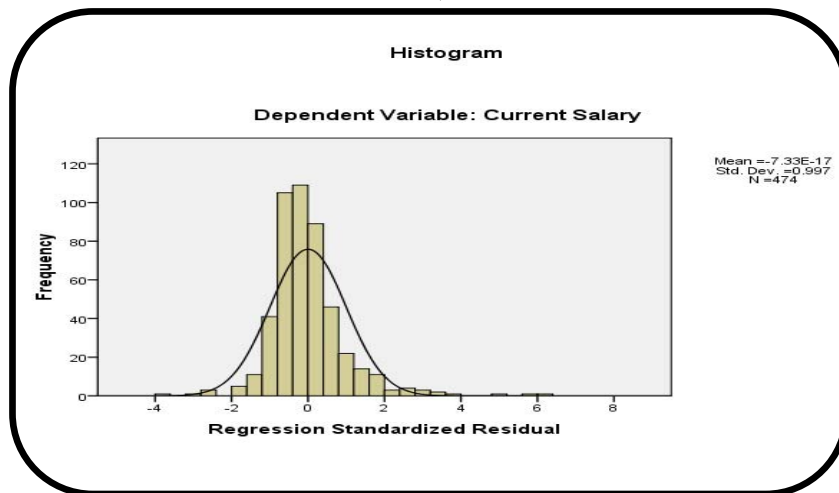
Residuals Statistics ^a					
	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	\$12,382.90	\$146,851.62	\$34,419.57	\$15,287.298	474
Residual	\$-28,852.994	\$48,701.199	\$-2.151E-13	\$7,607.676	474
Std. Predicted Value	-1.442	7.355	.000	1.000	474
Std. Residual	-3.781	6.381	.000	.997	474

همانطور که می‌دانیم یک مدل رگرسیونی قوی، مدلی است که بتواند متغیر وابسته را بطور دقیق پیش‌بینی کند. مثلاً در داده‌های مورد بررسی، اگر بتوانیم درآمد فعلی افراد را همانطور که واقعاً هست، پیش‌بینی کنیم، مدل دقیقی فراهم شده است. در مثال فعلی چنانچه مدل رگرسیونی حاصله را بر روی داده‌های مورد بررسی اعمال و مقادیر متغیرهای مستقل را در معادله رگرسیون قرار دهیم، انتظار می‌رود مقادیر پیش‌بینی شده به مقادیر واقعی بسیار نزدیک باشد و تفاضل آنها که اصطلاحاً باقی‌مانده نامیده می‌شود، خیلی کوچک باشند. در یک مدل رگرسیونی دقیق باقی‌مانده‌ها بسیار کوچک خواهند بود. همانطور که مشاهده می‌شود کمترین مقدار پیش‌بینی شده درآمد ۱۲۳۸۲.۹ دلار است که ۲۸۸۵۲.۹۹ از مقدار واقعی بیشتر است و بزرگترین مقدار پیش‌بینی شده ۱۵۲۸۷.۲۹ است که ۴۸۷۰۱.۱۹ کمتر از مقدار واقعی است. در جدول شاخص‌های آماری مربوط به مقادیر پیش‌بینی شده و باقی‌مانده استاندارد نیز آمده است.

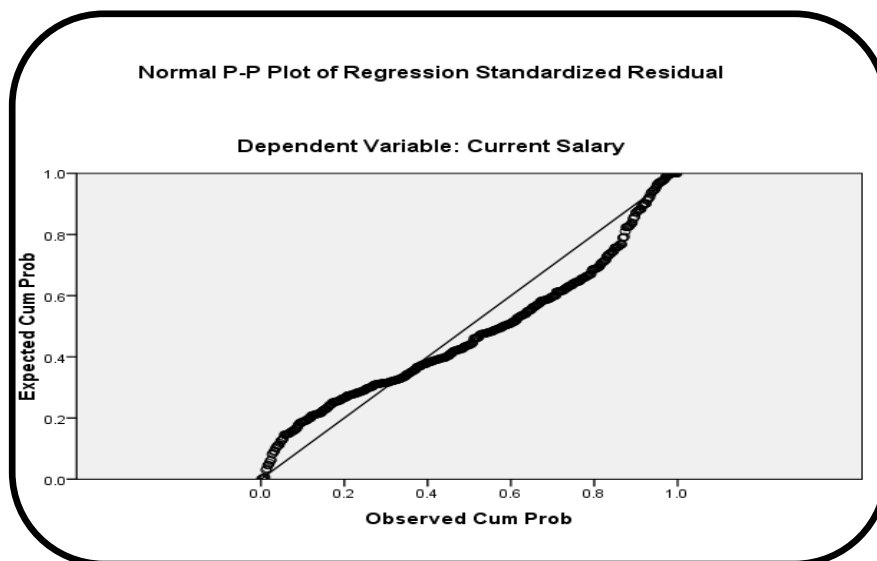
همانطور که در ابتدای مبحث رگرسیون آمد یکی از مفروضات استفاده از رگرسیون، نرمال بودن توزیع باقی‌مانده‌هاست. نمودارهای زیر نشان می‌دهند شکل توزیع باقیمانده‌های استاندارد شده رگرسیون دارای حالت نرمال بوده و بین آنها و باقی‌مانده‌های استاندارد مورد انتظار همبستگی بالایی وجود دارد (توزیع نقاط بر روی قطر نمودار بعدی، مؤید همین واقعیت است).

درس پانزدهم - آزمون مناسب برای تعیین جهت و شدت تأثیر متغیر مستقل بر وابسته □ ۲۰۷

شکل شماره (۷۸) نمودار هیستوگرام باقی مانده‌های استاندارد شده رگرسیون



شکل شماره (۷۹) نمودار احتمال نرمال بودن باقی مانده‌های استاندارد شده رگرسیون



تمرین درس پانزدهم

۱- جداول زیر نتایج رگرسیون مرگ و میر کودکان را بر متغیرهای مستقل درصد شهرنشینی، درصد باسوادان و میزان کالری مصرفی روزانه نشان می دهد. با دقت در نتایج به سؤالات مربوطه پاسخ دهید.

Model Summary^d

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.920 ^a	.846	.843	15.3293
2	.941 ^b	.886	.883	13.2606
3	.946 ^c	.894	.890	12.8525

a. Predictors: (Constant), People who read (%)

b. Predictors: (Constant), People who read (%), Daily calorie intake

c. Predictors: (Constant), People who read (%), Daily calorie intake, People living in cities (%)

d. Dependent Variable: Infant mortality (deaths per 1000 live births)

الف- ضرایب همبستگی و تعیین از گام اول تا سوم چه تغییری کرده اند و چرا؟

ب- خطای برآورد از گام اول تا سوم چه تغییری کرده است و چرا؟

ANOVA^d

Model	Sum of Squares	Df	Mean Square	F	Sig.
1 Regression	92648.098	1	92648.098	394.267	.000 ^a
Residual	16919.152	72	234.988		
Total	109567.250	73			
2 Regression	97082.367	2	48541.184	276.048	.000 ^b
Residual	12484.883	71	175.843		
Total	109567.250	73			
3 Regression	98004.162	3	32668.054	197.764	.000 ^c
Residual	11563.088	70	165.187		
Total	109567.250	73			

a. Predictors: (Constant), People who read (%)

b. Predictors: (Constant), People who read (%), Daily calorie intake

c. Predictors: (Constant), People who read (%), Daily calorie intake, People living in cities (%)

d. Dependent Variable: Infant mortality (deaths per 1000 live births)

ج- آیا نتایج رگرسیون در گام های مختلف قابلیت تعمیم به جامعه آماری را دارد؟

د- میانگین مجذورات باقی مانده در گام های مختلف چه تغییری کرده است؟

درس پانزدهم - آزمون مناسب برای تعیین جهت و شدت تأثیر متغیر مستقل بر وابسته □ ۲۰۹

Coefficients ^a					
Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	163.791	6.120		26.762	.000
People who read (%)	-1.540	.078	-.920	-19.856	.000
2 (Constant)	192.042	7.725		24.859	.000
People who read (%)	-1.227	.092	-.732	-13.373	.000
Daily calorie intake	-.019	.004	-.275	-5.022	.000
3 (Constant)	185.781	7.943		23.390	.000
People who read (%)	-1.140	.096	-.681	-11.861	.000
Daily calorie intake	-.015	.004	-.215	-3.653	.000
People living in cities (%)	-.212	.090	-.137	-2.362	.021

a. Dependent Variable: Infant mortality (deaths per 1000 live births)

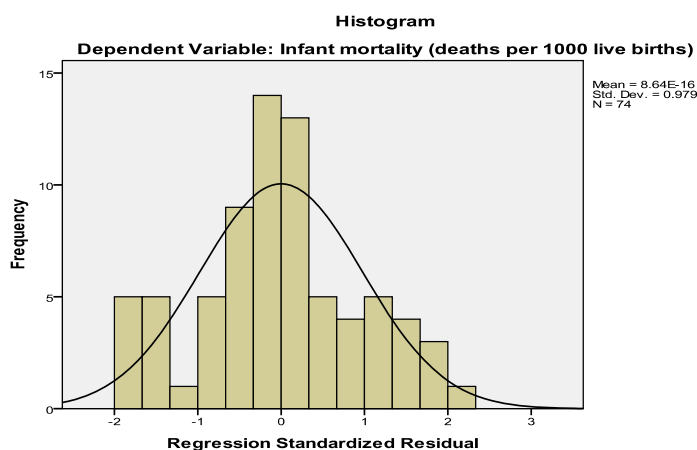
ه- معادله رگرسیونی برای پیش‌بینی نرخ مرگ‌ومیر کودکان را براساس نتایج جدول فوق بنویسید.

و- مقدار عرض از مبدأ (مقدار ثابت) چند شده است؟

ز- کدام متغیر بیشترین تأثیر را بر نرخ مرگ‌ومیر کودکان گذاشته است کدام است؟ جهت تأثیر این متغیر چگونه است؟

ح- به ازای یک واحد افزایش مصرف کالری روزانه در کشورها، نرخ مرگ‌ومیر کودکان چه تغییری می‌کند و چقدر؟

ط- توزیع باقی‌مانده‌های رگرسیونی چه حالتی دارد؟



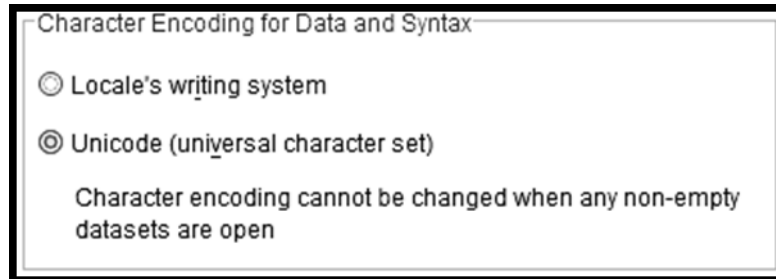
۲- براساس فایل Word95، رگرسیون رشد جمعیت را بر متغیرهای مستقل درصد

شهرنشینی، درصد باسوادان و میزان کالری مصرفی روزانه محاسبه و تفسیر نمایید.

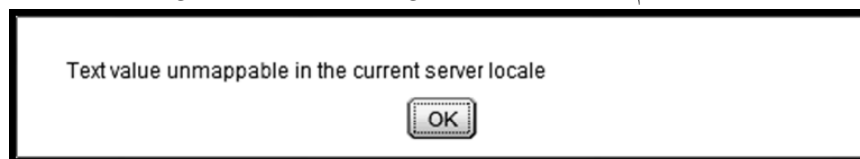
ضمیمه شماره ۱

برخی نکات لازم در هنگام نصب برنامه SPSS

- ۱- همیشه قبل از نصب برنامه، راهنمای نصب نرم افزار را که معمولاً در قالب یک فایل متنی در لوح فشرده نرم افزار وجود دارد به دقت بخوانید. در این فایل پیش نیازهای سخت افزاری یا نرم افزاری برنامه و همچنین مراحل نصب آن نوشته شده است.
- ۲- در هنگام نصب همیشه به شماره سریال این نرم افزار که معمولاً در قالب یک فایل متنی در لوح فشرده آن وجود دارد، توجه داشته باشید و آن را بطور دقیق وارد کنید.
- ۳- در برخی نسخه های این نرم افزار، بعد از عملیات نصب باید فایل های Crack را در لوح فشرده این برنامه تعبیه شده است، کپی و در پوشه ای این برنامه در سیستم خودتان الصاق (Paste) کنید. اگر در هنگام الصاق با سؤالی با مضمون بازنویسی (Overwrite) مواجه شدید، بر روی بلی کلیک کنید.
- ۴- در برخی نسخه ها، فایلی که در پوشه Crack قرار دارد یک فایل اجرایی به نام Patch.exe است که بعد از کپی و الصاق در پوشه SPSS، باید آن را اجرا کنید. این اجرا در صورتی درست است که پیام دال بر موفقیت اجرا، نمایش داده شود.
- ۵- به خاطر داشته باشید که چنانچه عملیات Crack یا Patch را انجام ندهید، برنامه spss اجازه ضبط داده ها و انجام عملیات آماری را به شما نخواهد داد. به عبارتی دیگر اگر چه برنامه باز می شود و شما حتی می توانید داده های آن را در آن وارد کنید، اما نه تنها نمی توانید هیچگونه عملیات آماری بر روی داده های آن انجام دهید، بلکه حتی نمی توانید داده های وارد شده را در جایی ضبط کنید.
- ۶- گاهی اوقات این نرم افزار پیام خطاری به شما می دهد با این مضمون که برای جلوگیری از منقضی شدن تاریخ استفاده از این نرم افزار، فایل Expoff.bat را اجرا کنید. توصیه می شود در صورت مواجهه این کار را انجام دهید.
- ۷- یکی از مشکلات کار با این نرم افزار، امکان تایپ فارسی برچسب متغیرها و مقادیر و نمایش آنها در بخش خروجی هاست. برای حل این مشکل توصیه می شود نسخه های 16 به بعد این نرم افزار را نصب کنید. برای فارسی کردن این نسخه ها کافی است که برنامه را اجرا کنید (بی آنکه هیچ فایل داده ای باز باشد)، سپس قسمت Unicode را از لایه General از منوی Edit/Option فعال کنید. با اینکار مشکل نوشتن و نمایش برچسب های فارسی در فایل های داده و خروجی برطرف می شود.



در غیر اینصورت هنگام نوشتن برچسب فارسی با خطای زیر مواجه می شوید.



- ۸- گاهی اوقات نسخه های جدید این نرم افزار برچسب های فارسی داده هایی را که در نسخ قبلی وارد شده است به شکل خوانا نمایش نمی دهند. برای این منظور باید از منوی کنترل پنل دستور **Region and Language** را اجرا و سپس بر روی **Administrative** کلیک کرده و **Language for non-Unicode programs** را به فارسی تغییر دهید.
- ۹- تأکید می شود نام فایل ها یا پوشه های مربوط به داده هایتان را به زبان فارسی ننویسید. چون در برخی نسخه ها این دسته از فایل ها اجرا نمی شوند.

ضمیمه شماره ۲: برخی منابع برای مطالعه بیشتر

ردیف	مبحث	منابع مناسب برای مطالعه بیشتر
۱	آمار توصیفی	۱. نایی، هوشنگ، ۱۳۸۹، آمار توصیفی برای علوم اجتماعی، تهران، انتشارات سمت. ۲. ساعی، علی، ۱۳۸۱، تحلیل آماری در علوم اجتماعی با نرم-افزار Spss for windows، تهران، انتشارات کیان مهر
۲	روایی و پایایی	۳. حیدری چروده، مجید. ۱۳۸۹، راهنمای سنجش روایی و پایایی در پژوهش‌های فرهنگی و اجتماعی، مشهد، انتشارات جهاد دانشگاهی مشهد. ۴. کارمینز، ادوارد. جی و ریچارد آزلر، ۱۳۸۶، سنجش روایی و پایایی، ترجمه سیمین فروغ‌زاده و مریم اسکافی، مشهد، انتشارات جهاد دانشگاهی مشهد.
۳	آزمون کی دو و ضرایب مربوط به جداول توافقی	۵. نگهبان، علیرضا، ۱۳۸۲، راهنمای روش تحقیق به کمک پرسشنامه، تهران نشر جهاد دانشگاهی. ۶. حسینی، سید یعقوب، ۱۳۸۲، آمار ناپارامتریک - روش تحقیق و نرم‌افزار آماری SPSS 10.0، تهران، انتشارات دانشگاه علامه طباطبائی.
۴	آزمون‌های برابری میانگین‌ها و تحلیل واریانس	۷. نصیری، رسول، آموزش گام‌به‌گام SPSS13، تهران، نشر گستر. ۸. زرگر، محمود، ۱۳۸۴، راهنمای جامع SPSS13، تهران، نشر بهینه. ۹. حبیب‌پور، کرم و رضا صفری، ۱۳۸۸، صفری شالی، راهنمای جامع کاربرد SPSS در تحقیقات پیمایشی. تهران، نشر لویه. ۱۰. کرامر، دانکن، ۱۳۸۶، درآمدی به کاربرد آمار در تحقیقات اجتماعی، ترجمه یحیی علی بابایی و امیر ملکی، تهران، انتشارات دانشگاه تهران.
۵	رگرسیون	۱۱. بابایی، علی، ۱۳۷۷، رگرسیون چند متغیره در نرم‌افزار

آماري SPSS، تهران، نشر جهاد دانشگاهي. ۱۲. افشانی، سید علیرضا، مرتضی نوریان و زینب حسینی رامشه، ۱۳۸۴، فرازی بر SPSS14، تهران، نشر بیشه. ۱۳. فتوحی، اکبر و علیرضا اصغری، ۱۳۸۲، آنالیز آماری داده‌ها در SPSS11، تهران، نشر کانون. ۱۴. حبیب‌پور، کرم و رضا صفری، ۱۳۸۸، صفری شالی، راهنمای جامع کاربرد SPSS در تحقیقات پیمایشی. تهران، نشر لوپه.	خطی چندگانه	
---	--------------------	--

منابع و مآخذ

- افشانی، سیدعلیرضا، مرتضی نوریان و زینب حسینی رامشه، ۱۳۸۴، فرازی بر SPSS14، تهران، نشر بیشه.
- بابایی، علی، ۱۳۷۷، رگرسیون چند متغیره در نرم افزار آماری SPSS، تهران، نشر جهاد دانشگاهی.
- حبیب پور، کرم و رضا صفری، ۱۳۸۸، صفری شالی، راهنمای جامع کاربرد SPSS در تحقیقات پیمایشی. تهران، نشر لویه.
- حسینی، سید یعقوب، ۱۳۸۲، آمار ناپارامتریک - روش تحقیق و نرم افزار آماری SPSS 10.0، تهران، انتشارات دانشگاه علامه طباطبائی.
- حیدری چروده، مجید. ۱۳۸۳، بررسی عملکرد خوابگاه های مدارس شبانه روزی، مشهد، سازمان مدیریت و برنامه ریزی استان خراسان رضوی.
- حیدری چروده، مجید. ۱۳۸۴، بررسی تطبیقی آگاهی ها، نگرش ها و رفتارهای دانش آموزان مدارس دارا و فاقد مربی بهداشت در زمینه بهداشت، مشهد، جهاد دانشگاهی مشهد.
- حیدری چروده، مجید. ۱۳۸۸ نیازسنجی شهروندان تهرانی در زمینه ورزش های همگانی، تهران، دفتر مطالعات اجتماعی و فرهنگی شهرداری تهران.
- حیدری چروده، مجید. ۱۳۸۹، راهنمای سنجش روایی و پایایی در پژوهش های فرهنگی و اجتماعی، مشهد، انتشارات جهاد دانشگاهی مشهد.
- زرگر، محمود، ۱۳۸۴، راهنمای جامع SPSS13، تهران، نشر بهینه.
- ساعی، علی، ۱۳۸۱، تحلیل آماری در علوم اجتماعی با نرم افزار Spss for windows، تهران، انتشارات کیان مهر
- فتوحی، اکبر و علیرضا اصغری، ۱۳۸۲، آنالیز آماری داده ها در SPSS11، تهران، نشر کانون.
- فروغ زاده، سیمین، ۱۳۸۵، بررسی میزان استفاده دانش آموزان از مواد خوراکی موجود در فروشگاه های مدارس و راهکارهایی برای استفاده از مواد خوراکی مفید و سالم، سازمان آموزش و پرورش استان خراسان شمالی.
- کارمینز، ادوارد. جی و ریچارد آ. زلر، ۱۳۸۶، سنجش روایی و پایایی، ترجمه سیمین فروغ زاده و مریم اسکافی، مشهد، انتشارات جهاد دانشگاهی مشهد

کرامر، دانکن، ۱۳۸۶، درآمدی به کاربرد آمار در تحقیقات اجتماعی، ترجمه یحیی علی بابایی و امیر ملکی، تهران، انتشارات دانشگاه تهران.

نایی، هوشنگ، ۱۳۸۹، آمار توصیفی برای علوم اجتماعی، تهران، انتشارات سمت.

نصیری، رسول، آموزش گام به گام SPSS13، تهران، نشر گستر.

نگهبان، علیرضا. (۱۳۸۲)، راهنمای روش تحقیق به کمک پرسشنامه -spss11، تهران، نشر جهاد دانشگاهی تهران

تازه‌های انتشارات جامعه‌شناسان

- ۱- نظریه‌سازی در علوم اجتماعی، شومیکر و همکاران، محمد عبداللہی
- ۲- نظریات جدید جامعه‌شناسی و ورزش، محمدرضا مهرآئین
- ۳- ماهیت و انواع نظریه‌های جامعه‌شناسی (جلد ۱)، دان مارتیندال، حمید عبداللہیان
- ۴- زنان در عرصه عمومی (پژوهش برتر سال)، محمد عبداللہی
- ۵- کاربرد نظریه‌های ارتباطات، سون ویندال و همکاران، علیرضا دھقان
- ۶- پژوهش، پژوهشگری و پژوهشنامه‌نویسی (جلد ۲)، خلیل میرزایی
- ۷- مقدمه‌ای بر مدل‌سازی معادله ساختاری با کاربرد نرم‌افزارهای Lisrel، spss، EQS (با Cd) شوماخر و لومکس، وحید قاسمی
- ۸- جامعه‌شناسی صلح، گاستون بوتول، هوشنگ فرخجسته
- ۹- جامعه‌شناسی جنگ، کاپلو و ونسن، هوشنگ فرخجسته
- ۱۰- جامعه‌شناسی روزنامه‌نگاری، برایان مک‌نیر، علی اصغر کیا - محمد رسولی
- ۱۱- اندیشه اجتماعی و سیاسی سبز، اندرو دابسون، محسن ثلاثی
- ۱۲- جامعه‌شناسی نظم اجتماعی، پیتر ورسلی، سعید معیدفر
- ۱۳- فرهنگ و اژگان علوم اجتماعی، شیوا شکاری/ وجهه معصوم
- ۱۴- جامعه‌شناسی ورزش (وندالیسم و اوباشگری در ورزش فوتبال)، وحید قاسمی، وحید ذوالاکتاف، علی نورعلی‌وند
- ۱۵- جامعه‌شناسی جنسیت، خدیجه سفیری، سارا ایمانیان
- ۱۶- جامعه‌شناسی مصرف (توریسم و خرید)، علی اصغر سعیدی - محمدحسین آبادی
- ۱۷- مبانی توسعه پایدار شهری، مهرداد نوابخش - اسحق ارجمند سیاهپوش
- ۱۸- نفوذ اجتماعی، کاظم سام دلیری
- ۱۹- کاربرد آمار در علوم اجتماعی با استفاده از نرم‌افزار SPSS و نحوه تفسیر خروجی‌ها، سعید گودرزی
- ۲۰- سیاست اجتماعی و توسعه، آنتونی حال - جیمز میجلی، مهدی ابراهیمی - علیرضا صادقی
- ۲۱- درک ایرانیان از عدالت، داریوش یعقوبی
- ۲۲- درآمدی بر فرهنگ‌های سایبر، دیوید بل، مسعود کوثری - حسین حسینی
- ۲۳- راهنمای بین‌المللی برآورد پیامدهای اجتماعی، فرنک ونکلی - هنک بکر، هادی جلیلی
- ۲۴- جامعه‌شناسی کشورهای اسلامی، غلامعباس توسلی و بقیه نویسندگان
- ۲۵- جامعه‌شناسی علم، محمد توکل
- ۲۶- مدل‌سازی معادله ساختاری با کاربرد نرم‌افزار Amos Graphics (با Cd)، وحید قاسمی
- ۲۷- جامعه‌شناسی کیفیت زندگی، مرضیه مختاری - جواد نظری
- ۲۸- جامعه‌شناسی سیاسی کلان، محمدرضا طالبان
- ۲۹- طرح و پایان‌نامه نویسی، خلیل میرزایی
- ۳۰- فرا روش، (بنیان‌های فلسفی و عملی تحقیق ترکیبی در علوم اجتماعی و رفتاری) احمد محمدپور
- ۳۱- روش در روش (درباره ساخت معرفت در علوم انسانی)، احمد محمدپور
- ۳۲- ضد روش جلد ۱ (منطق و طرح در روش‌شناسی کیفی)، احمد محمدپور
- ۳۳- الگوها و فرایندها، رابرت ال‌بی، رحیم فرخ‌نیا
- ۳۴- افکار عمومی در ایران، مرکز افکارسنجی دانشجویان ایران (نشر مشترک ایسپا و جامعه‌شناسان)
- ۳۵- درآمدی بر روش‌های کمی و کیفی تحقیق در علوم اجتماعی، احمد سعیدی و... (نشر مشترک ایسپا و جامعه‌شناسان)
- ۳۶- نظریه‌های جامعه‌شناسی، عباس محمدی اصل
- ۳۷- فرهنگ و رفتار اجتماعی، هری س. تری یاندیس، نصرت فتی

- ۳۸- Spectrum سیستم مدل‌سازی سیاست‌گذاری، جان استور و...، حاتم حسینی
- ۳۹- زندگی و اندیشه‌ی بزرگان انسان‌شناسی، جری. دی. مور، هاشم آقا بیگ پوری - جعفر احمدی
- ۴۰- نظریه اجتماعی و سیاست هویت، کریچ کالهن، قلی پور - محمدزاده
- ۴۱- زنانی که مرد می‌شوند، آنتونیا یونگ، منیژه مقصودی - محمد فتوره‌چی
- ۴۲- مقدمه‌ای بر سیاست‌گذاری اجتماعی، کن بلیک مور، علی اصغر سعیدی - سعید صادقی جقه
- ۴۳- سیاست اجتماعی (نظریه‌های کلاسیک، مدرن، پست‌مدرن و مطالعات مقایسه‌ای) فرید همتی
- ۴۴- سنجش مفاهیم در پیمایش‌های اجتماعی، فردین علیخواه
- ۴۵- واژه‌نامه ارتباطات و رسانه‌ها، حسین حسینی - جعفر احمدی
- ۴۶- کمیت و کیفیت در تحقیقات اجتماعی، آلن براین، هاشم آقا بیگ پوری
- ۴۷- نخستین‌های تاریخ روزنامه‌نگاری ایران، رامین رسول‌اف
- ۴۸- مجموعه مقالات همایش ملی خودکشی (علل، پیامدها و راهکارها)، دفتر تحقیقات کاربردی فرماندهی انتظامی استان ایلام.
- ۴۹- ارزیابی تأثیرات اجتماعی، محمد فاضل
- ۵۰- انسان‌شناسی فلسفی، رحیم فرخ‌نیا - علیرضا غفاری
- ۵۱- فراتحلیل در پژوهش‌های اجتماعی و رفتاری، محمود قاضی طباطبایی - ابوعلی ودادبیر
- ۵۲- شیوه‌ی عملی مقاله‌نویسی، خلیل میرزایی
- ۵۳- جامعه‌شناسی روستایی، حسین بهروان
- ۵۴- ارتباطات کامپیوتری (اینترنت و جامعه)، آلیس تومیک و همکاران، سروناز تربتی
- ۵۵- گنج رایگان (درآمدی بر جامعه‌شناسی ایران امروز)، عباس محمدی اصل
- ۵۶- روش ارزیابی مشارکتی روستایی، چمبرز و همکاران، منیژه مقصودی
- ۵۷- جامعه‌شناسی جوانان (سبک‌های زندگی جوانان در دنیای متغیر) استیون مایلز، مینا قریب - نعیمه جوان
- ۵۸- استراتژی‌های پژوهش اجتماعی، نورمن بلیکی، جعفر احمدی و آقا بیگ پوری
- ۵۹- اصول آمار مقدماتی در علوم رفتاری، با تأکید بر SPSS فاروق امین مظفری - خورشید پاداش اصل
- ۶۰- اصول آمار پیشرفته در علوم رفتاری با تأکید بر SPSS فاروق امین مظفری - خورشید پاداش اصل
- ۶۱- جامعه‌شناسی حقوق، اسحق ارجمند سیاهپوش
- ۶۲- جامعه‌شناسی تکنولوژی، محمد توکل
- ۶۳- روابط عمومی الکترونیکی، رحمان سعیدی، علی اصغر کیا
- ۶۴- جامعه‌شناسی روابط بین نسلی، تقی آزاد ارمکی
- ۶۵- آموزش SPSS / ۱۶ و نحوه تفسیر خروجی‌ها، سعید گودرزی
- ۶۶- شاهد تن (خشونت بنیادین علیه زنان جنوب آسیایی ساکن ایالات متحده) شمیتا، داس داسگویتا، بنفشه جوادی الیادرائی
- ۶۷- تصمیمات اخلاقی در رویه مددکاری اجتماعی، رالفا دولگوف، فرانک م. لونبرگ، دونا هارینگتون، فضا نوپور
- ۶۸- جامعه‌شناسی مسائل اجتماعی در ایران، جواد افشار کهن
- ۶۹- تجارت خارجی ایران (در عهد قاجار و پهلوی)، میثم موسایی
- ۷۰- دین، توسعه و فرهنگ، میثم موسایی
- ۷۱- تکنیک‌های خاص تحقیق، ذبیح‌الله صدقی - سکیته بابایی
- ۷۲- چندرسانه‌ای‌ها، آن کرنی فرانسیس، علی اصغر کیا
- ۷۳- سیستم‌های استنباط فازی در پژوهش‌های اجتماعی، وحید قاسمی
- ۷۴- مجله مطالعات اجتماعی ایران (دوره سوم، شماره ۲، تابستان ۱۳۸۸ و ویژه‌نامه شهر و شهروندی)، انجمن جامعه‌شناسان
- ۷۵- مجله مطالعات اجتماعی ایران (دوره سوم، شماره ۳، پاییز ۱۳۸۸ و ویژه‌نامه خراسان رضوی)، انجمن جامعه‌شناسان
- ۷۶- مجله مطالعات اجتماعی ایران (دوره سوم، شماره ۴، زمستان ۱۳۸۸ و ویژه‌نامه روش‌شناسی)، انجمن جامعه‌شناسان

فهرست کتاب‌های در دست انتشار

- ۱- جامعه‌شناسی و جامعه ایرانی، تقی آزاد ارمکی
- ۲- نظام اجتماعی، تالکت پارسنز، نادر سالارزاده امیری
- ۳- مدرن گرایی و سرمایه اجتماعی خانواده، عالیه شکریگی
- ۴- جامعه‌شناسی کار (دنیای شگفت‌انگیز کار در مدرنیته بازاندیشانه)، اولریش بک، سعید صادقی - آیت نباتی حصارلو
- ۵- مشاوره با زنان معتاد، مونیک کوهن، فریده همتی
- ۶- دموکراسی و سنت‌های مدنی، روبرت پاتنام، محمدتقی دلفروز
- ۷- خشونت خانگی، آزاده مومنی
- ۸- طرح چگونگی وحدت ملل، اسرافیل رحیمی موفر
- ۹- تأثیرات زندان بر زندانی، عباس عبدی
- ۱۰- کاربرد تحلیل گفتمان در تحقیقات اجتماعی، حسین علی قجری - جواد نظری
- ۱۱- پارسنز و شکل‌گیری جامعه امروز، عباس محسنی اصل
- ۱۲- علم و تکنیک به مثابه ایدئولوژی در باب منطق علوم اجتماعی، غلامرضا خدیوی
- ۱۳- هجرت طبقه خلاق یا فرار مغزها از آمریکا، ریچارد فلوریدا، ابراهیم انصاری - محمداسماعیل انصاری
- ۱۴- شهرها و طبقه خلاق، ریچارد فلوریدا، ابراهیم انصاری - محمداسماعیل انصاری
- ۱۵- دموکراسی، چارلز تیلی، یعقوب احمدی
- ۱۶- تحلیل نسلی، افسانه کمالی
- ۱۷- تجاوز جنسی در آمریکا، ساسان ودیعه
- ۱۸- در جستجوی نظریه‌های عالمان مسلمان، محمد ثقفی
- ۱۹- جامعه‌شناسی سازمان‌ها و فرهنگ سیاسی، مهرداد نوابخش - حسین شریعت
- ۲۰- فرهنگ جامع جامعه‌شناسی و علوم انسانی، کرامت الله راسخ
- ۲۱- سیاست، طرح‌ها و برنامه‌های بهداشت روان، فردین علیپور
- ۲۲- مفاخر جامعه‌شناسی ایران - ۱ (غلامعباس توسلی)، به کوشش محمدباقر تاج‌الدین
- ۲۳- جامعه‌شناسی روابط عمومی در ایران، عباس محمدی اصل
- ۲۴- برداشت‌هایی از نظریه جامعه‌شناسی معاصر، آنتونی الیوت، فرهنگ ارشاد
- ۲۵- صنعت فرهنگی، تئودور آدرنو، سمیه رنجبر
- ۲۶- هابرماس به روایت فمینیست‌ها، سعید اسلامی، شیوا شکاری
- ۲۷- جنسیت و تضاد نقش‌های زناشویی، زهرا زارع
- ۲۸- تحلیل شبکه‌های اجتماعی به همراه آموزش نرم افزار USINET، ابوالفضل رضائی، علی میرزامحمدی
- ۲۹- راهنمای گام به گام استفاده از نرم‌افزار MAXqda2 در تحقیق کیفی، ابوالفضل رضائی
- ۳۰- موازین بین‌المللی حقوق زنان (کنوانسیون‌ها، میثاق‌ها، قطعنامه‌ها، پروتکل‌ها، اعلامیه‌ها، بیانیه‌ها، توصیه‌نامه‌ها، مقاله‌نامه‌ها...)، (جلد اول) شهیندخت مولاوردی
- ۳۱- موازین بین‌المللی حقوق زنان (اسناد مصوب کنفرانس‌های جهانی و اجلاس‌های ویژه) جلد دوم، شهیندخت مولاوردی
- ۳۲- درآمدی بر شبکه‌های اجتماعی، ییروئن بروگمن، خلیل میرزایی

- ۳۳- نظریه پردازی در جامعه‌شناسی (فقر نظریه پردازی در ایران) مریم رحیمی سجاسی
- ۳۴- هویت انسانی توسعه، مرتضی نظری
- ۳۵- مساله روش در مطالعات فرهنگی، حمید عبداللهیان، علی حاجی محمدی، علی صباغی
- ۳۶- نظریه جامعه‌شناختی معاصر و ریشه‌های کلاسیک آن، جرج ریتزر، خلیل میرزایی، علی بقایی
- ۳۷- جامعه‌شناسی کاپوچینو، مهرداد ناظری، ایمان هنریور
- ۳۸- روش تحقیق کیفی در جامعه‌شناسی، دیوید سیلورمن، محسن ثلاثی
- ۳۹- چگونه مسائل اجتماعی را حل کنیم، جیمز کرون، مهرداد نوابخش، فاطمه کرمی
- ۴۰- جامعه‌شناسی جنگ، علیرضا شایان مهر
- ۴۱- دایرة المعارف خانواده نظریه‌ها و پژوهش‌های خانواده، لینگستون، عالیہ شکریگی، شهلا قهرمانی
- ۴۲- روش تحقیق در انسان‌شناسی، دیویدنی هانتر، ماریان بی‌فولی، علیرضا قبادی، فاطمه دیهیمی، مریم زینی
- ۴۳- تعاون در جامعه جدید، مارک ون دو، تام. ار. تایلر، آندره بیلز، سعید وصالی
- ۴۴- نظریه جامعه‌شناسی مدرن، جورج ریتزر، داگلاس جی. گودمن، عباس لطفی‌زاده،
- ۴۵- جامعه‌شناسی معاصر شهری، ابراهیم انصاری
- ۴۶- جامعه‌شناسی سلطه، مهران صولتی
- ۴۷- ابن خلدون نخستین جامعه‌شناس مسلمان، محمد ثقفی
- ۴۸- تحلیل عاملی تاییدی با استفاده از نرم‌افزارهای LISREL, AMOS, MPLUS جرمی. جی. آلبریگت، حسین صفری
- ۴۹- مقدمه‌ای بر جامعه‌شناسی انواع دینداری، سید محمد میرسندسی