



آموزش نرم افزار SPSS

گروه آمار
مدیریت آمار و فناوری اطلاعات



فهرست مطالب



- شروع هوشمندانه با SPSS
- مدیریت حرفه‌ای داده‌ها
- آمار توصیفی با چاشنی گرافیک
- آزمون‌های آماری با تفسیر کاربردی
- همبستگی و رگرسیون با نگاه تحلیلی
- جمع‌بندی



شروع ہوشمندانه با SPSS



معرفی SPSS

Statistical Package for the Social Sciences (SPSS) نرم‌افزاری قدرتمند برای تحلیل داده‌های آماری است که در علوم انسانی، اجتماعی، پزشکی و مدیریتی کاربرد فراوان دارد.

محیط SPSS از سه بخش اصلی تشکیل شده است:

بخش	توضیح	مثال
Data View	محل ورود داده‌ها: هر سطر = یک فرد هر ستون = یک متغیر	ستون «Age» برای سن، «Gender» برای جنسیت
Variable View	محل تعریف مشخصات متغیرها	نام متغیر (بدون فاصله و انگلیسی)، نوع داده، برچسب، داده‌های گم‌شده، تعداد اعشار....
Output Viewer	محل نمایش نتایج تحلیل‌ها	جدول‌ها و نمودارها

ورود و ویرایش داده‌ها

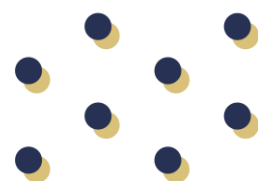


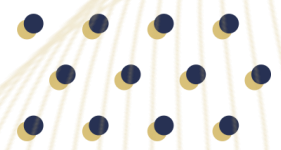
داده‌ها را می‌توانید به دو روش وارد کنید:

- دستی تایپ در Data View
- وارد کردن از فایل‌های دیگر Excel، CSV و... از مسیر:
File → Import Data

انتقال سریع داده‌های Excel به SPSS:

- داده‌ها را در Excel آماده کنید: ستون‌ها = متغیرها، سطرها = پاسخ‌دهندگان.
- کل محدوده داده را انتخاب و Copy کنید.
- در SPSS در بخش Data View کلیک کنید.
- گزینه Paste with Variable Names را انتخاب کنید.





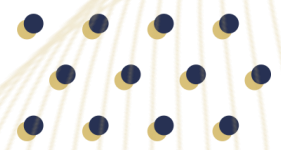
تنظیمات پیشرفته در Variable View

در این قسمت SPSS به شما اجازه می‌دهد رفتار هر متغیر را تنظیم کنید.

مثال	کاربرد	گزینه
نام متغیر Label → Sat: «رضایت از تدریس»	توضیح کامل متغیر برای نمایش در خروجی	Label
۱=زن، ۲=مرد	تعریف کدگذاری برای داده‌های کیفی	Values
مقدار ۹۹ به عنوان Missing	تعریف داده‌های گمشده	Missing
Nominal, Ordinal, Scale	نوع مقیاس اندازه‌گیری	Measure

💡 نکته: اگر داده‌ها چندمرحله‌ای یا گروه‌بندی شده هستند مثلاً کلاس‌ها یا شهرها، Measure را Ordinal یا Nominal درست انتخاب کنید تا در آزمون‌ها اشتباه نگیرید.





استفاده از Templates و متغیرهای هوشمند برای ورود و استانداردسازی داده های تکراری

در بسیاری از پژوهش ها مثلاً پرسشنامه ها، ساختار داده ها همیشه ثابت است.

به جای تعریف دوباره متغیرها، می توانید از **Templates** استفاده کنید:

یک فایل SPSS با ساختار متغیرهای دلخواه بسازید. یعنی متغیرهایی که از قبل با Label، Values و Missing تنظیم شده اند.

بدون وارد کردن داده ها، آن را با نامی مثل Template_Questionnaire.sav ذخیره کنید.

هر زمان داده ی جدید داشتید، فقط فایل را باز کنید و داده ها را Paste کنید.

✓ این روش باعث می شود در پروژه های پژوهشی مختلف، یکپارچگی در نام متغیرها و کدگذاری حفظ شود. سرعت کار چند برابر

می شود و خطاهای ورود داده تقریباً صفر می شود.





مدیریت داده‌ها مثل حرفه‌ای‌ها



آشنایی با مفاهیم پایه‌ای در مدیریت داده‌ها

قبل از شروع تحلیل، داده‌ها باید تمیز و آماده باشند. مشکلات رایج داده‌ها عبارت‌اند از:

- داده‌های گم‌شده (**Missing Values**): داده‌هایی که ثبت نشده یا ناقص هستند.
- داده‌های پرت (**Outliers**): مقادیری که بسیار متفاوت از بقیه داده‌ها هستند و می‌توانند تحلیل را منحرف کنند.

چرا پاک‌سازی داده‌ها مهم است؟

- آماده‌سازی داده برای تحلیل‌های پیشرفته
- جلوگیری از تحلیل نادرست
- نتایج دقیق‌تر و معتبرتر



آشنایی با مفاهیم پایه‌ای در مدیریت داده‌ها

🎯 هدف از Missing Value Analysis

در پژوهش‌های واقعی معمولاً داده‌ها ناقص‌اند (مثلاً پاسخگو به بعضی سؤالات جواب نداده).

تحلیل مقادیر گمشده کمک می‌کند بفهمیم:

چه مقدار از داده‌ها گم شده‌اند؟

آیا الگوی گمشده‌ها تصادفی است یا سیستماتیک؟

چه روشی برای برخورد با داده‌های گمشده مناسب‌تر است (حذف، میانگین‌گیری، برآورد، و غیره)؟



آشنایی با مفاهیم پایه‌ای در مدیریت داده‌ها

بررسی Missing Values و نحوه برخورد با آن‌ها

• شناسایی داده‌های گمشده:

- ۱) Analyze → Descriptive Statistics → Frequencies → Missing Values
- ۲) Analyze → Missing Value Analysis

• روش‌های برخورد:

۱. حذف رکوردهای گمشده Listwise/Pairwise
۲. جایگزینی با میانگین یا مد
۳. استفاده از تکنیک‌های پیشرفته مثلاً Multiple Imputation



آشنایی با مفاهیم پایه‌ای در مدیریت داده‌ها

داده های پرت

- **Outliers:** مقادیر خیلی بزرگ یا خیلی کوچک که می‌توانند تحلیل را خراب کنند.

- مسیر: Analyze → Descriptive Statistics → Explore

- فعال کردن Boxplot برای شناسایی مقادیر پرت

راهکارهای کنترل داده های پرت

- حذف مقادیر پرت

- اصلاح یا بازکدگذاری





دسته‌بندی هوشمند داده‌ها

- **Recode into Same Variables** به شما امکان می‌دهد مقادیر یک متغیر را بازکدگذاری کنید.

○ مثال: تبدیل سن به گروه‌های سنی

- **Recode into Different Variables** به شما امکان می‌دهد مقادیر یک متغیر را بازکدگذاری کنید و یک متغیر جدید بسازید.

○ مثال: تبدیل متغیر جنسیت از متن به عدد: مرد $\rightarrow$ ۱ زن $\rightarrow$ ۲

- **مزایا:**

○ حفظ متغیر اصلی

○ ایجاد گروه‌بندی‌های منطقی مثلاً گروه‌های سنی: جوان، میانسال، سالمند





ساخت شاخص‌های ترکیبی با Compute Variable

- **Compute Variable** برای ساخت متغیرهای جدید از داده‌های موجود استفاده می‌شود.

- مثال: شاخص سلامت = میانگین وزن، قد و سن

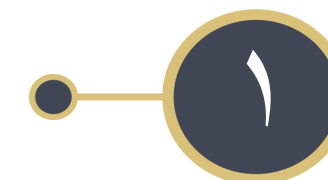
- مثال: شاخص رضایت از آموزش = میانگین ۵ سؤال پرسشنامه



نوشتن فرمول مورد نظر



وارد کردن نام متغیر جدید



انتخاب → Transform
Compute Variable

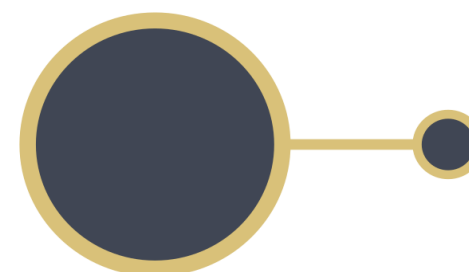


فیلتر کردن داده‌ها با Select Cases و Split File

Select Cases: فقط نمونه‌های مشخص را برای

تحلیل انتخاب می‌کند.

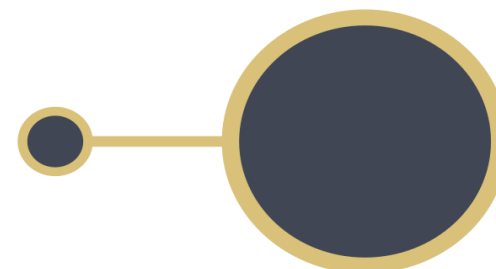
مثال: انتخاب افراد بالای ۱۸ سال



Split File: داده‌ها را به گروه‌های مختلف تقسیم و تحلیل

جداگانه انجام می‌دهد.

مثال: تحلیل جداگانه برای هر جنسیت یا هر گروه سنی



❗ نکته: همیشه بعد از اتمام تحلیل Split File را خاموش کنید، وگرنه تحلیل بعدی بدون توجه به گروه‌ها انجام می‌شود.





آمار توصیفی با چاشنی گرافیک



محاسبه آمارهای پایه: میانگین، میانه، انحراف معیار و درصدها

Descriptives

مسیر: Analyze → Descriptive Statistics → Descriptives

محاسبه: میانگین Mean، انحراف معیار Std. Deviation، حداقل و حداکثر، Median: میانه

برای متغیرهای نامتقارن،

Mode: نما برای داده‌های کیفی، Range: دامنه تغییرات



محاسبه آمارهای پایه: میانگین، میانه، انحراف معیار و درصدها

Frequencies

مسیر: Analyze → Descriptive Statistics → Frequencies

محاسبه: درصدها، تعداد فراوانی Count، نمودارهای سریع

گزینه: فعال کردن Display frequency tables برای مشاهده جدول دقیق



محاسبه آمارهای پایه: میانگین، میانه، انحراف معیار و درصدها

Explore

مسیر: Analyze → Descriptive Statistics → Explore

مناسب برای تحلیل عمیق تر و شناسایی Outlierها

ارائه جدولهای Summary، Percentiles و Boxplot

💡 نکته: استفاده همزمان از Explore و Descriptives کمک می کند توزیع، واریانس

و مقادیر پرت را بهتر بشناسیم.



مرور آمارهای توصیفی فراتر از میانگین و انحراف معیار



فراتر از مقادیر پایه، آمار توصیفی حرفه‌ای شامل:

Percentiles و Quartiles تحلیل پراکندگی و نقاط مرجع

Skewness و Kurtosis بررسی تقارن و کشیدگی توزیع

رسم نمودارهای سفارشی

SPSS ابزارهای متنوعی برای نمودارهای حرفه‌ای ارائه می‌دهد:

نوع نمودار	کاربرد
Bar Chart	مقایسه فراوانی گروه‌ها
Pie Chart	نمایش سهم درصدی در متغیرهای اسمی
Histogram	بررسی توزیع داده‌های عددی
Boxplot	شناسایی Outlier و پراکندگی و نمایش چارکها

❗ نکته: قبل از رسم نمودار، داده‌ها را دسته‌بندی کنید Split File یا Recode

از رنگ‌بندی، لیبل و افکت‌ها برای جذابیت و وضوح بیشتر استفاده کنید



استفاده از Chart Editor برای زیباسازی نمودارها

پس از ایجاد نمودار:

روی نمودار دوبار کلیک کنید → Chart Editor باز می شود

امکان تغییر:

رنگ ها و الگوها

فونت ها و اندازه ها

افزودن عنوان، برچسب و خطاهای استاندارد





مصورسازی حرفه‌ای با Chart Builder

Chart Builder ابزار قدرتمندی برای ساخت نمودارهای سفارشی است:

مسیر: Graphs → Chart Builder

امکانات حرفه‌ای:

انتخاب نوع نمودار بر اساس نوع متغیر اسمی، ترتیبی، عددی

انتخاب محور X و Y ، گروه‌بندی بر اساس متغیرها

ایجاد نمودارهای ترکیبی Combo: Bar + Line

استفاده از گزینه Clustered یا Stacked برای مقایسه گروه‌ها

استفاده از گزینه‌ی Element Properties برای تنظیمات دقیق

📌 نکته: نمودارهای ترکیبی برای مقایسه چند متغیر در یک نمودار بسیار کاربردی هستند.

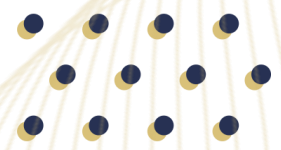
مثال: میانگین رضایت در گروه‌های جنسیتی همراه با درصد مشارکت نمایش داده شود.





آزمون‌های آماری با تفسیر کاربردی





قبل از هر آزمون: بررسی پیش فرض‌ها

تقریباً همه آزمون‌های پارامتریک در SPSS فرض می‌کنند که داده‌ها دارای شرایط زیر هستند:

پیش فرض	روش بررسی در SPSS	توضیح
نرمال بودن توزیع	Analyze → Descriptive Statistics → Explore → Plots → Normality plots with tests	خروجی آزمون Shapiro-Wilk یا Kolmogorov-Smirnov
برابری واریانس‌ها	ANOVA Levene's Test و t آزمون	اگر Sig. > 0.05 باشد، واریانس‌ها برابرند
استقلال داده‌ها	باید از طراحی پژوهش اطمینان حاصل شود	داده‌های هر گروه مستقل از گروه دیگر باشند

📌 نکته: پیش از اجرای آزمون، از گزینه‌ی Descriptives در منوی Analyze → Descriptive Statistics برای بررسی میانگین، انحراف معیار و شکل توزیع داده‌ها استفاده کنید تا تخمینی از رفتار متغیرها داشته باشید و برای تشخیص نقاط پرت و انتخاب آزمون مناسب.

💡 نکته: اگر داده‌ها نرمال نباشند، از آزمون‌های ناپارامتریک استفاده کنید Kruskal–Wallis، Wilcoxon، Mann–Whitney.





قبل از هر آزمون: بررسی پیش فرض‌ها

✈ نکته: پیش از اجرای آزمون، از گزینه‌ی Descriptives در منوی Analyze → Descriptive Statistics برای بررسی میانگین، انحراف معیار و شکل توزیع داده‌ها استفاده کنید تا تخمینی از رفتار متغیرها داشته باشید و برای تشخیص نقاط پرت و انتخاب آزمون مناسب.

💡 نکته: اگر داده‌ها نرمال نباشند، از آزمون‌های ناپارامتریک استفاده کنید Mann–Whitney، Wilcoxon، Kruskal–Wallis.



آزمون t مستقل Independent Samples T-Test

کاربرد: مقایسه میانگین دو گروه مستقل مثلاً مردان و زنان در میزان رضایت.

مراحل:

مسیر: Analyze → Compare Means → Independent-Samples T Test

متغیر وابسته: عددی مثلاً نمره رضایت در قسمت Test Variable

متغیر مستقل: گروهی مثلاً جنسیت در Grouping Variable

تعریف کد گروه‌ها مثلاً ۱=مرد، ۲=زن

تفسیر خروجی:

اگر در قسمت آزمون لون $\text{Sig.} > 0.05$ از سطر بالا "Equal variances assumed" استفاده کنید.

اگر $\text{Sig. 2-tailed} > 0.05$ باشد: تفاوت میانگین دو گروه معنادار است.

مثال: اگر میانگین رضایت مردان ۳,۲ و زنان ۴,۱ باشد و $\text{Sig.} = 0.03 \rightarrow$

نتیجه: «سطح رضایت زنان به طور معناداری بیشتر از مردان است.»



آزمون tزوجی Paired Samples T-Test

کاربرد: مقایسه میانگین دو وضعیت مرتبط پیش آزمون و پس آزمون، قبل و بعد از آموزش.

مراحل:

Analyze → Compare Means → Paired-Samples T Test

مسیر:

وارد کردن دو متغیر مرتبط مثلاً PreTest و PostTest

تفسیر خروجی:

Mean Difference: اختلاف میانگین ها

اگر Sig. ۲-tailed کمتر از ۰,۰۵ : تغییر معنادار است

مثال:

میانگین نمرات قبل = ۱۲,۴، بعد از آموزش = ۱۵,۱ → Sig.=۰/۰۰۲

نتیجه: «آموزش باعث افزایش معنادار نمره آزمودنی ها شده است.»





آزمون ANOVA تحلیل واریانس یک طرفه

کاربرد: مقایسه میانگین بیش از دو گروه مستقل مثلاً سه رشته تحصیلی مختلف.

مراحل:

Analyze → Compare Means → One-Way ANOVA

مسیر:

متغیر وابسته: مثلاً «رضایت از تدریس»

متغیر گروهی: مثلاً «رشته تحصیلی»

گزینه‌ی Post Hoc را برای مقایسه زوجی فعال کنید Bonferroni یا Tukey

گزینه‌های رایج: ➤

Tukey برای واریانس‌های برابر

Games-Howell برای واریانس‌های نابرابر





آزمون ANOVA تحلیل واریانس یک طرفه

تفسیر خروجی:

Sig. در جدول ANOVA: اگر $0.05 > 0.001 \rightarrow$ تفاوت معنادار بین گروه‌ها وجود دارد.

Post Hoc Test: مشخص می‌کند کدام گروه‌ها با هم تفاوت دارند.

مثال:

میانگین رضایت	رشته
۴,۳	روانشناسی
۳,۸	مدیریت
۳,۱	آمار

و $\rightarrow \text{Sig.} = 0.001$

نتیجه: «سطح رضایت بین سه رشته متفاوت است و دانشجویان روانشناسی رضایت بیشتری دارند.»





آزمون کای دو Chi-Square Test

کاربرد: بررسی رابطه بین دو متغیر طبقه‌ای کیفی. مثلاً بررسی رابطه بین جنسیت و نوع انتخاب رشته.

مراحل:

Analyze → Descriptive Statistics → Crosstabs

مسیر:

یکی از متغیرها در Rows و دیگری در Columns

گزینه‌ی Statistics → Chi-square را فعال کنید.

تفسیر خروجی:

Pearson Chi-Square: اگر $\rightarrow \text{Sig.} < 0.05$ رابطه بین متغیرها معنادار است.

از Crosstab Table برای تفسیر درصدها در هر گروه استفاده می‌شود.

مثال:

اگر بین جنسیت و نوع رشته تحصیلی $\rightarrow \text{Sig.} = 0.04$

نتیجه: «بین جنسیت و انتخاب رشته رابطه معناداری وجود دارد.»



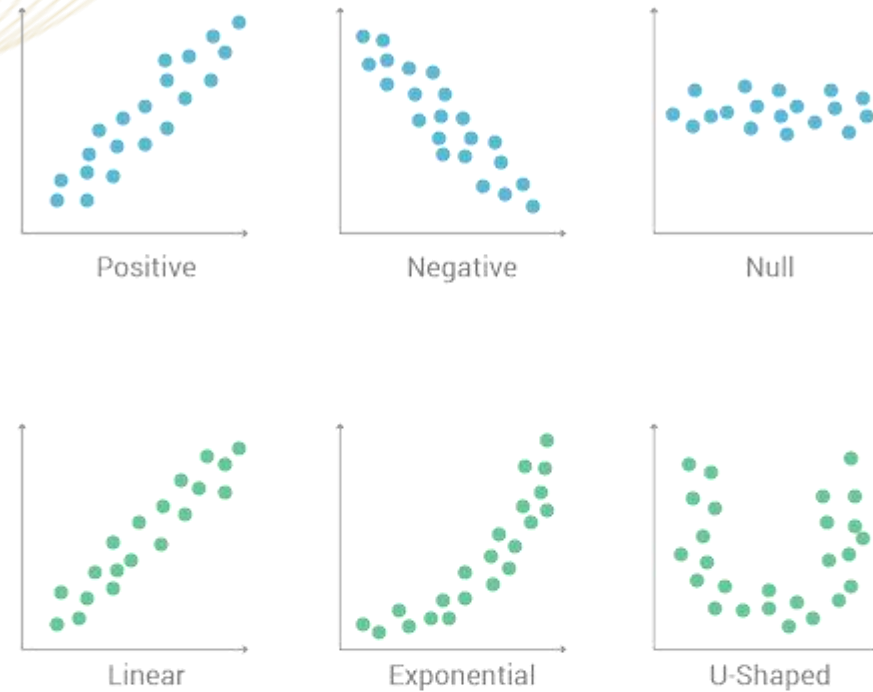


همبستگی و رگرسیون با نگاه تحلیلی



مفهوم همبستگی Correlation

Types of Correlations



هدف: بررسی میزان و جهت رابطه خطی بین دو متغیر عددی

نوع رابطه	توضیح	مثال
مثبت +	با افزایش یکی، دیگری هم افزایش می‌یابد	قد و وزن
منفی -	با افزایش یکی، دیگری کاهش می‌یابد	اضطراب و عملکرد تحصیلی
صفر	رابطه‌ای وجود ندارد	شماره کفش و معدل

Correlation Strength



نکته: تفاوت بین همبستگی و علیت (Correlation $\neq$ Causation)



آزمون همبستگی پیرسون Pearson Correlation

کاربرد: بررسی رابطه خطی بین دو متغیر فاصله‌ای یا نسبی که توزیع نرمال دارند.

مراحل:

Analyze → Correlate → Bivariate

مسیر:

دو متغیر عددی را انتخاب کنید.

گزینه‌ی Pearson را فعال و Two-tailed را انتخاب کنید.

تفسیر خروجی:

اگر $\text{Sig.} < 0.05 \rightarrow$ رابطه معنادار است.

علامت مثبت یا منفی جهت رابطه را مشخص می‌کند.

مثال:

همبستگی بین «ساعات مطالعه» و «نمره امتحان» $r = 0.68$, $\text{Sig.} = 0.002$

نتیجه: رابطه مثبت و قوی بین ساعات مطالعه و نمره وجود دارد.



مقدار ضریب r	شدت رابطه	تفسیر
0.00-0.19	بسیار ضعیف	تقریباً بدون رابطه
0.20-0.39	ضعیف	رابطه‌ی کم‌قدرت
0.40-0.59	متوسط	رابطه‌ی قابل توجه
0.60-0.79	قوی	رابطه‌ی بالا
0.80-1.00	بسیار قوی	تقریباً خطی کامل



آزمون همبستگی اسپیرمن Spearman Correlation

کاربرد: زمانی که داده‌ها نرمال نیستند یا مقیاس آن‌ها رتبه‌ای **Ordinal** است.

مراحل:

همان مسیر قبلی **Analyze → Correlate → Bivariate**

اما این بار گزینه‌ی **Spearman** را انتخاب کنید.

مثال:

$$r = ۰.۵۵, \text{Sig.} = ۰.۰۱$$

نتیجه:  بین رتبه رضایت مشتریان و میزان خرید رابطه مثبت و متوسط وجود دارد.





تحلیل رگرسیون خطی ساده Simple Linear Regression

هدف: پیش‌بینی یک متغیر وابسته عددی Y از روی یک متغیر مستقل X

بررسی اولیه:

رسم نمودارهای **scatter**

تحلیل همبستگی

فرمول:

$$Y = a + bX$$

که در آن:

a = Intercept عرض از مبدأ

b = ضریب رگرسیون شیب خط





تحلیل رگرسیون خطی ساده Simple Linear Regression

مراحل:

Analyze → Regression → Linear

مسیر:

متغیر وابسته در **Dependent**

متغیر پیش‌بین در **Independents**

تفسیر خروجی:

R: ضریب همبستگی بین پیش‌بین و وابسته

در **ANOVA table** اگر **Sig** > 0.05 : مدل رگرسیون معنادار است

B(ضریب رگرسیون): نشان می‌دهد با افزایش یک واحد در **X**، مقدار **Y** چقدر تغییر می‌کند





تحلیل رگرسیون خطی ساده Simple Linear Regression

مثال:

اگر

$$Y = 20 + 2/5X \text{ و } \text{Sig.} = 0/01$$

که در آن X میزان ساعت مطالعه و Y نمره ی کسب شده هست.

نتیجه: «با افزایش هر ساعت مطالعه، نمره به طور میانگین ۲.۵ واحد افزایش می یابد.»





تحلیل رگرسیون چندگانه Multiple Regression

هدف: پیش‌بینی یک متغیر وابسته از چند متغیر مستقل.

 فرمول:

$$Y = a + b_1X_1 + b_2X_2 + \dots + b_nX_n$$

مراحل:

Analyze → Regression → Linear

مسیر:

یک متغیر وابسته و چند متغیر مستقل انتخاب کنید.

در قسمت **Statistics**، گزینه‌های **Collinearity diagnostics** و **Durbin-Watson** را فعال کنید.





تحلیل رگرسیون چندگانه Multiple Regression

تفسیر خروجی:

Adjusted R²: دقت مدل پس از کنترل تعداد متغیرها و حجم نمونه نسخه‌ی تعدیل‌شده‌ی **R²**.

Beta Standardized Coefficient: اهمیت نسبی متغیرها در پیش‌بینی **Y**.

Multicollinearity: بررسی هم‌خطی چندگانه

اگر **Tolerance** > 0.1 باشد، مشکل هم‌خطی وجود دارد

اگر **VIF** < 10 باشد، هشدار هم‌خطی

مثال:

مدل پیش‌بینی رضایت شغلی بر اساس درآمد، ساعت کار، و سطح تحصیلات





نمودارهای تحلیلی در رگرسیون

مسیر: **Graphs → Chart Builder → Scatter/Dot → Simple Scatter**

متغیر وابسته روی محور **Y** و متغیر پیش‌بین روی **X**

✎ نکته: خط روند را از طریق **Chart Editor** اضافه کنید

Residuals: برای بررسی باقیمانده‌ها

از گزینه‌ی **Plots** در پنجره‌ی **Regression** استفاده کنید و **ZRESID vs ZPRED** را رسم کنید.

الگوی تصادفی: پیش‌فرض استقلال و نرمال بودن برقرار است.





✈ ترفند ویژه: ساخت مدل رگرسیون با انتخاب متغیرهای مؤثر Stepwise

در رگرسیون چندگانه، اگر تعداد متغیرهای پیش‌بین زیاد است، می‌توان از روش **Stepwise** برای انتخاب خودکار متغیرهای مؤثر استفاده کرد.

روش اجرا:

در پنجره‌ی **Linear Regression**، در قسمت **Method** گزینه‌ی **Stepwise** را انتخاب کنید.
SPSS به ترتیب، متغیرهایی را که بیشترین تأثیر را دارند وارد مدل می‌کند.

مزیت:

مدل نهایی ساده‌تر و دقیق‌تر است.
فقط متغیرهای معنی‌دار باقی می‌مانند.

📊 مثال:

در مدل پیش‌بینی رضایت شغلی از میان ۵ متغیر، **SPSS** فقط «درآمد» و «تعادل کار و زندگی» را وارد مدل می‌کند.
نتیجه: 📈 مدل نهایی مؤثر و بدون پیچیدگی اضافی است.





تर्फندهای طلایی و جمع‌بندی



آشنایی با Syntax در SPSS

Syntax زبان دستوری SPSS است که برای اجرای سریع، تکرارپذیر و مستندسازی تحلیل‌ها استفاده می‌شود.

مزیت‌ها:

- اجرای چند تحلیل به صورت هم‌زمان
- ثبت و بازتولید کامل مراحل تحلیل (مستندسازی و قابل ردیابی بودن مراحل تحلیل)
- اشتراک‌گذاری پروژه با دیگر پژوهشگران
- صرفه‌جویی در زمان در پروژه‌های بزرگ (تکرار آسان تحلیل‌ها برای داده‌های جدید)



اجرای سریع تحلیل‌ها با Syntax



- چطور Syntax را فعال کنیم:
از منو: File → New → Syntax
دستورات را بنویسید و سپس با دکمه ▶ یا Run → All اجرا کنید

مثال: اجرای هم‌زمان آمار توصیفی برای چند متغیر 

- هر تحلیلی که در محیط گرافیکی انجام می‌دهید، در واقع پشت‌صحنه‌اش Syntax دارد. ✓
با کلیک دکمه Paste در هر پنجره‌ی تحلیل، دستور آن در پنجره Syntax ذخیره می‌شود.

❗ نکته: هر بار که تحلیل جدیدی در SPSS انجام می‌دهید، روی Paste بزنید تا Syntax آن در فایل ذخیره شود. در پایان می‌توانید همه را در یک فایل sps ترکیب کنید و پروژه را خودکارسازی کنید.



مثال پروژه‌های از Syntax کامل:



* ورود داده‌ها:

* پاک‌سازی داده‌ها.

* آمار توصیفی.

* آزمون t مستقل.

* رگرسیون چندگانه.

* ذخیره خروجی.





ساخت گزارش حرفه‌ای از خروجی‌ها

SPSS خروجی‌ها را در قالب Output Viewer نمایش می‌دهد. برای ارائه‌ی پژوهشی بهتر، می‌توانید آن‌ها را با فرمت‌های زیر صادر یا ویرایش کنید:

فرمت خروجی	کاربرد	مسیر ذخیره‌سازی
spv.	خروجی اصلی SPSS	File → Save As
docx.	برای گزارش در Word	File → Export → Word
pptx.	برای ارائه در PowerPoint	File → Export → PowerPoint
xls.	انتقال داده‌ها به Excel	File → Export → Excel

📌 نکته: در Word، نمودارها و جداول SPSS کاملاً قابل ویرایش‌اند. برای زیباسازی سریع:

رنگ پس‌زمینه را حذف کنید

برچسب‌ها را خلاصه و شفاف بنویسید

عنوان و منبع هر جدول را درج کنید



جمع‌بندی نکات کلیدی دوره

ابزار SPSS	مهارت کلیدی	موضوع
Transform → Compute / Recode	پاک‌سازی، بازکدگذاری، استانداردسازی	مدیریت داده‌ها
Descriptives / Frequencies / Chart Builder	بررسی میانگین، درصد و نمودارها	آمار توصیفی
Compare Means / Crosstabs	Chi-Square, ANOVA, T	آزمون‌های آماری
Correlate / Regression	کشف روابط و پیش‌بینی	همبستگی و رگرسیون
Syntax Editor / Export	اجرای Syntax و گزارش‌نویسی	ترندهای حرفه‌ای

یادآوری: 

همیشه پیش‌فرض‌های آزمون‌ها را بررسی کنید.

داده‌های گم‌شده را با روش‌های منطقی اصلاح کنید نه میانگین ساده.

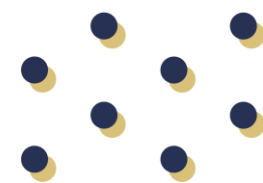
خروجی‌ها را با دقت تفسیر کنید، نه فقط خواندن اعداد Sig.

مستندسازی تحلیل با Syntax نشانه‌ی حرفه‌ای بودن است.

مرور نکات کلیدی و اشتباهات رایج

موضوع	نکته کلیدی	اشتباه رایج
ورود داده‌ها	همیشه مقیاس Scale/Ordinal/Nominal را درست انتخاب کن	همه متغیرها را پیش فرض "Scale" گذاشتن
کدگذاری	Labelها را دقیق تعریف کنید	استفاده از کد بدون برچسب
تحلیل‌ها	قبل از آزمون، پیش فرض‌ها را بررسی کنید	اجرای آزمون بدون بررسی نرمال بودن
تفسیر	فقط به مقدار Sig. نگاه نکنید	نادیده گرفتن جهت یا اندازه اثر
خروجی	خروجی را تمیز و خلاصه ارائه بده	ارسال خروجی خام SPSS به صورت کامل

✎ یادآوری: تحلیل داده یعنی درک داستان پشت اعداد، نه فقط گزارش نتایج.



سپاسی از توجه شما

