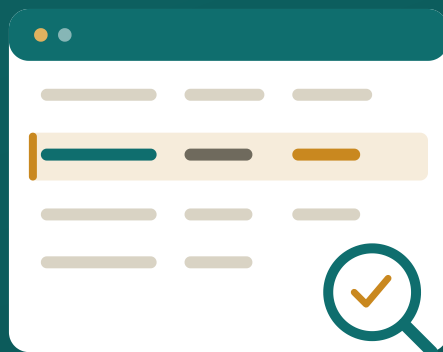




The SPSS Output Decoder

Which table to read, which row the assumption test sends you to, the effect sizes the output hides, and the diagnostics a committee checks before they trust your model. Reading statistical output the way an examiner does, from PhD statisticians.



Written by PhD statisticians

dissertationstatisticshelp.org

The software gives you everything, so the skill is knowing what to ignore and what it left out

Run a test and the output is a wall of tables. Two failures sink results chapters: reading the wrong row of the right table, and reporting only what the default output prints while missing the effect size or diagnostic a reviewer expects. This decoder covers both.

PRINCIPLE

For almost every analysis there is one table that holds the answer, one assumption test that decides *which row* of it you may read, one effect size (often not printed by default), and one or two diagnostics that decide whether the model is trustworthy at all. Find those, and you can both write the result and defend it.

Four things to locate, not three

1 The estimate and its interval

The difference, slope, or ratio with a confidence interval. Report the interval, not just the point estimate; it carries the precision the p-value does not.

2 The significance value, read from the licensed row

The Sig. column tells you compatibility with no effect. But for several tests an assumption check (Levene, Mauchly) determines which row's Sig. you are allowed to read. Reading the default row regardless is a classic error.

3 The effect size, computed if SPSS will not print it

Older SPSS prints no standardised effect for a t-test, and one-way ANOVA prints none unless you used the GLM dialog. You often have to compute or request it. A result without one is incomplete.

4 The diagnostics that gate the model

Collinearity, sphericity, goodness of fit, and influence decide whether the headline number means anything. Examiners check these before the p-value.

Which table, which licensed row, which effect size

For the common analyses: the table to read, the assumption test that picks the row, the significance number, and the effect size to report (with where to get it when SPSS hides it).

ANSWER TABLE, ROW-SELECTING ASSUMPTION, SIGNIFICANCE, AND EFFECT SIZE

Analysis	Table and row-selector	Significance	Effect size (and source)
Independent t-test	Independent Samples Test; Levene decides the row	Sig. (2-tailed)	Hedges' g; v27+ Effect Sizes table, else compute
Paired t-test	Paired Samples Test	Sig. (2-tailed)	Cohen's d on the difference; report direction
One-way ANOVA	ANOVA + Post Hoc; use GLM for effect size	Sig. for the factor	Omega squared (less biased) or partial eta squared from GLM
Repeated-measures ANOVA	Tests of Within-Subjects Effects; Mauchly decides the row	Sig. on the corrected row	Partial eta squared
Correlation	Correlations	Sig. (2-tailed)	r itself; r squared as variance shared
Linear regression	Coefficients + Model Summary	Sig. per predictor	Standardised beta; R squared and adjusted R squared
Chi-square	Chi-Square Tests	Asymptotic Sig. (or Fisher exact if cells sparse)	Cramer's V or phi
Logistic regression	Variables in the Equation	Sig. per predictor	Exp(B) odds ratio with its 95% CI

WATCH OUT FOR

The Levene trap: in the Independent Samples Test table, read Levene's Test first. If its Sig. is below 0.05 the variances differ, so you must read the "Equal variances not assumed" row, which applies the Welch-Satterthwaite correction and adjusts the degrees of freedom downward (they will be non-integer). Reading the equal-variances row when Levene is significant is the most common t-test error in submitted theses. The Mauchly trap is its repeated-measures twin: if Mauchly's test of sphericity is significant, do not read the "Sphericity Assumed" row; read the Greenhouse-Geisser corrected row (or Huynh-Feldt when epsilon exceeds about 0.75).

The numbers SPSS does not print by default, and how to get them

A reviewer asks for these even though the default output omits them. Knowing where each one lives separates a defensible chapter from an incomplete one.

MISSING-BY-DEFAULT, AND THE ROUTE TO IT

What is missing	Why it matters	How to get it in SPSS
Standardised effect for a t-test (pre-v27)	No effect size means an incomplete result	Compute Hedges' g from the means and pooled SD, or upgrade to the v27+ Effect Sizes table
Omega squared for ANOVA	Eta squared is positively biased	Compute omega squared from the ANOVA sums of squares; SPSS prints only eta squared family
Confidence interval on an effect size	Precision of the effect, not just its point	v27+ prints it for d; otherwise use a noncentral-t calculator or bootstrap
Robust intervals under non-normality	Default SEs assume normality	The Bootstrap module gives bias-corrected accelerated (BCa) intervals
Pseudo R-squared and fit for logistic	Tells whether the model fits at all	Nagelkerke R squared in Model Summary; Hosmer-Lemeshow you want non-significant

PRINCIPLE

Two reporting habits that read as careless: writing partial eta squared as if it were comparable across designs (it is not, and partial values can sum past one), and copying SPSS's "Sig. = .000". SPSS rounds; .000 means the value is smaller than 0.0005, so report it as $p < .001$, never $p = .000$.

The diagnostics an examiner reads before your p-value

For regression and logistic models, the headline coefficient is only as good as the assumptions behind it. These are the checks that get raised in a viva.

1 Collinearity, in the Coefficients table

Request Collinearity diagnostics. Read Tolerance and the variance inflation factor (VIF). A VIF above about 5 warrants concern and above 10 signals serious multicollinearity that destabilises the coefficients; near-zero tolerance says the same. Report them, do not just have them.

2 Independence and influence

The Durbin-Watson statistic in Model Summary flags autocorrelated residuals (values near 2 are clean). Save Cook's distance and standardised residuals: a Cook's distance above 1, or several cases beyond three standard residuals, points to influential cases you must address, not delete silently.

3 Logistic model fit and separation

Check the Omnibus Tests of Model Coefficients (the model should improve on the null), Nagelkerke R squared for explained variation, the Hosmer-Lemeshow test (you want it non-significant), and the classification table. A coefficient with an enormous standard error usually signals complete or quasi-complete separation, not a real mega-effect.

Turn the licensed numbers into a results sentence

With the estimate, the correctly read significance value, the effect size, and the interval, the sentence almost writes itself. Reporting follows American Psychological Association style, which most committees expect.

Results sentence templates, by test

Independent t-test (Levene non-significant, equal variances assumed):

An independent t-test showed [group 1] ($M = [mean]$, $SD = [sd]$) scored [higher/lower] than [group 2] (M, SD), $t([df]) = [t]$, $p = [exact]$, Hedges' $g = [g]$, 95% CI [lower, upper].

If Levene was significant, report the Welch row and its non-integer df.

Repeated-measures ANOVA (Mauchly significant, Greenhouse-Geisser):

There was a [significant] effect of [factor], $F([GG\ df1], [GG\ df2]) = [F]$, $p = [exact]$, partial eta squared = [value].

Logistic regression:

[Predictor] predicted [outcome], $B = [B]$, Wald = [W], $p = [exact]$, odds ratio = [Exp(B)], 95% CI [lower, upper].

Always write $p < .001$, never $p = .000$.

PRINCIPLE

Report the effect size and interval next to every significance value, and report the diagnostic you relied on (which Levene or Mauchly outcome, the VIF, the fit statistic). A results chapter that names the assumption check and the row it sent you to reads as the work of someone who understands the analysis, not someone copying tables.

The output reading checklist

Run each analysis through this list before you copy a single number into your chapter.

- ☒ You identified the one table that holds the answer for this analysis.
- ☒ You let the assumption test choose the row (Levene for the t-test, Mauchly for repeated measures).
- ☒ You reported an effect size, computing or requesting it when the default output omitted it.
- ☒ You used omega squared or named partial eta squared honestly, not as a cross-design constant.
- ☒ You wrote $p < .001$ rather than copying SPSS's .000.
- ☒ You checked and reported the gating diagnostics (VIF, Durbin-Watson, Cook's distance, model fit).
- ☒ You reported a confidence interval on the estimate and, where possible, on the effect size.
- ☒ Your results sentence names the test, the licensed numbers, the effect, and its direction.

WHEN THE OUTPUT HAS TO BE RIGHT

Have your analysis run and the results chapter written with you

If you would rather not decode the output alone, our PhD statisticians run the analysis, read every assumption-licensed row correctly, report the effect sizes and diagnostics SPSS hides, and hand you a committee-ready results section in your field's reporting style, with a plain-language explanation of each number. You stay the author and make every final call.

Get a quote for your statistics

dissertationstatisticshelp.org . Written by PhD statisticians . We work in writing, on your schedule.