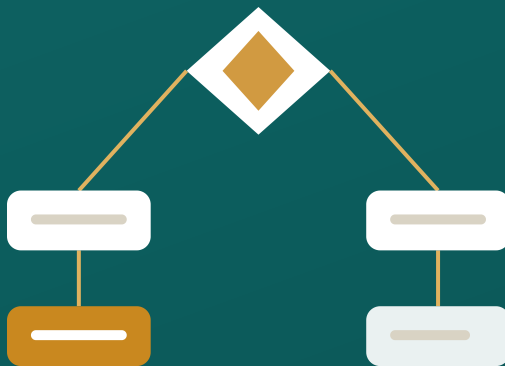


Which Statistical Test Should You Use?

A plain-language decision guide that takes you from your research question and your data to the correct test, and shows you what to do when the assumptions do not hold. Written by PhD statisticians.



Written by PhD statisticians dissertationstatisticshelp.org

Choosing a test is a logic chain, not a memory test

You do not pick a statistical test by remembering a name. You answer a short series of questions about your data and your comparison, and the test falls out at the end. Once you can run that chain, the choice stops feeling arbitrary and starts feeling obvious.

PRINCIPLE

Almost every test choice is decided by four things: the type of your outcome variable, how many groups or measurements you are comparing, whether those groups are independent or related, and whether the assumptions of the test are reasonably met. Answer those four and you have your test.

The four questions that decide your test

1 What type is your outcome variable?

Continuous (a measured quantity such as score, time, or concentration), ordinal (ranked categories), or categorical (named groups, including yes or no). This single answer rules out most of the options immediately.

2 How many groups or conditions are you comparing?

One sample against a known value, two groups, or three or more groups. The number of things being compared separates a t test from analysis of variance, and a chi-square from a simple proportion test.

3 Are the groups independent or related?

Independent means different people in each group. Related (paired) means the same people measured twice, or matched pairs. The same comparison uses a different test depending on this answer.

4 Are the assumptions reasonably met?

Most classic tests assume roughly normal data and similar variances. When those hold, you use the standard test. When they do not, you move to the matched alternative in the table that follows.

From your comparison to the right test

Find the row that matches what you are comparing and the type of your outcome. This covers the analyses that appear in the large majority of dissertations and theses.

COMPARISON TO TEST, FOR A CONTINUOUS OUTCOME

What you are comparing	Groups are independent	Groups are related (paired)
One group against a known value	One sample t test	Not applicable
Two groups	Independent samples t test	Paired samples t test
Three or more groups, one factor	One way analysis of variance	Repeated measures analysis of variance
Three or more groups, two factors	Two way analysis of variance	Mixed analysis of variance
Relationship between two continuous variables	Pearson correlation	Not applicable
Predicting one continuous outcome	Linear regression	Not applicable

COMPARISON TO TEST, FOR A CATEGORICAL OR ORDINAL OUTCOME

What you are comparing	Typical test
Association between two categorical variables	Chi-square test of independence
One proportion against a known value	Binomial or one sample proportion test
Two groups on a ranked or non-normal outcome	Mann-Whitney U test
Three or more groups on a ranked or non-normal outcome	Kruskal-Wallis test
Two related measurements on a ranked outcome	Wilcoxon signed-rank test
Predicting a yes or no outcome	Logistic regression

The parametric to non-parametric swap

If your continuous outcome is badly skewed, has strong outliers, or the sample is small, the classic test may not be trustworthy. Each one has a non-parametric counterpart that makes fewer assumptions. The comparison you are testing does not change, only the test does.

IF THE NORMALITY ASSUMPTION DOES NOT HOLD, SWAP ACROSS

Classic test	Non-parametric alternative	Same question, fewer assumptions
Independent samples t test	Mann-Whitney U test	Are two independent groups different
Paired samples t test	Wilcoxon signed-rank test	Did the paired measurements change
One way analysis of variance	Kruskal-Wallis test	Do three or more groups differ
Pearson correlation	Spearman correlation	Do two variables move together

WATCH OUT FOR

Do not run both the classic and the non-parametric test and then report whichever gives the smaller probability value. Decide the rule in advance: check the assumption, and if it fails, use the alternative. Choosing the test after seeing the result is the kind of thing a committee is trained to spot.

PRINCIPLE

Two refinements that mark a more sophisticated choice. First, the non-parametric swap quietly changes the hypothesis: a Mann-Whitney U test does not compare means, it tests whether one group tends to produce larger values (stochastic dominance), so a significant result is not a statement about average difference unless you can assume identically shaped distributions. Report the median and a rank-biserial effect size accordingly, not the mean. Second, there is a third path the table does not show: robust methods. Welch's t-test, which does not assume equal variances, is now recommended as the *default* over Student's t even when variances look equal, because it costs almost nothing when they are and protects you when they are not. For skew with outliers, a trimmed-means t-test or a bootstrap confidence interval keeps the question about means while resisting the violation, which non-parametric tests cannot.

How to check the assumptions without overdoing it

Assumption checking is where many students either skip a step or go too far and over-interpret a single number. A calm, proportionate check is what you want.

- ☒ Look at the shape of the data with a histogram or a normal quantile plot before trusting any single significance test of normality.
- ☒ On larger samples, a normality test will flag tiny, harmless departures, so judge the plot and the sample size together, not the probability value alone.
- ☒ Check that the variability is similar across groups, because a large imbalance changes which version of the test you should run.
- ☒ Confirm the observations are independent, which is an assumption of design, not something a plot can rescue after the fact.
- ☒ If an assumption is clearly violated, move to the matched alternative rather than reporting a result the data do not support.

Run the chain on a real question

Here is the four-question chain applied to a common dissertation comparison, so you can see how quickly it resolves.

Worked example: comparing exam scores across three teaching methods

Question: Do three teaching methods produce different exam scores?

1. Outcome type: Exam score is continuous.
2. Groups compared: Three methods, so three groups.
3. Independent/related: Different students in each method, so independent.
4. Assumptions: Scores roughly normal, variances similar.

Chain result: One way analysis of variance.

If assumptions failed: Kruskal-Wallis test instead.

Follow up if significant: Post-hoc comparisons to find which methods differ.

PRINCIPLE

Notice that the test was never guessed. Each answer narrowed the field until one option remained. Write your own question into the same four lines and the test will reveal itself the same way.

The test selection checklist

Run your analysis plan through this list. If every box is ticked, you can defend the choice of test to anyone who asks.

- ☒ You have classified your outcome variable as continuous, ordinal, or categorical.
- ☒ You have counted the groups or conditions you are comparing.
- ☒ You have decided whether those groups are independent or related.
- ☒ You have checked the key assumptions and chosen the alternative if one failed.
- ☒ Where you switched to a non-parametric test, you reported the matching effect size and median, not the mean.
- ☒ You considered a robust option (Welch by default, trimmed means, or bootstrap) rather than assuming non-parametric is the only fallback.
- ☒ You can state, in one sentence, why this test answers your research question.
- ☒ You know which follow-up comparison you will run if the overall test is significant.
- ☒ You decided the test before looking at which option gives the smaller probability value.

WHEN YOU WANT THE ANALYSIS DONE RIGHT

Have the right test chosen, run, and explained for your data

If you would rather hand the analysis to someone who does this every day, our PhD statisticians choose the correct test for your design, run it, check the assumptions, and give you committee-ready tables with a plain-language interpretation so you can explain every result. 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.