

Statistical Power, Sample Size, and Minimum Detectable Effect

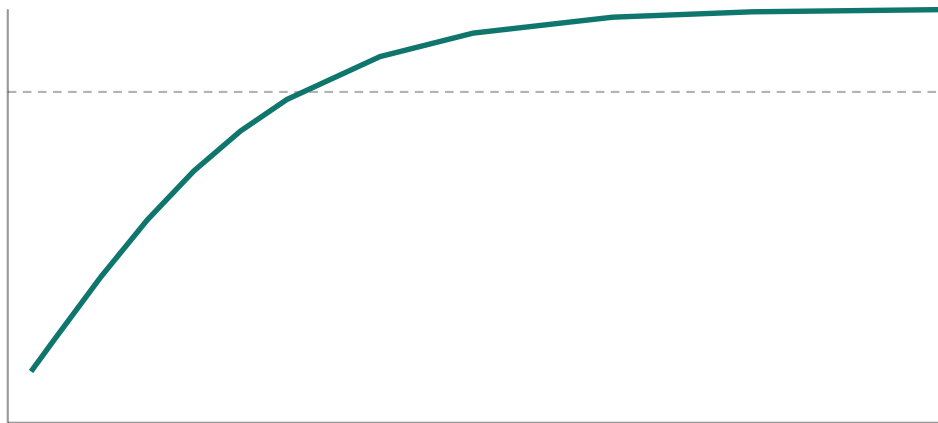
Power planning makes a study design transparent before data are collected. Adjust the standardized effect size below to see how the power curve, and the sample size needed to reach 80% power, both shift.



Statistical Power, Sample Size, and Minimum Detectable Effect

Exported from statreason.com. Deterministic, browser-local content; not professional statistical or research advice.

Standardized effect size (Cohen's d): 0.50



n per group for 80% power (alpha = .05) 63

At n=63 per group, this design reaches 80% power for an effect size of 0.50. Dashed line marks 80% power.

What this answers

Power planning asks how likely a chosen analysis is to distinguish a stated effect from its null model when the effect, variability, sample size, and assumptions are as planned. The curve above shows achieved power across a range of sample sizes for an independent two-group t-test design, at your chosen effect size.

The four pieces that trade off

Alpha is the selected false positive threshold, fixed at .05 here. Power is the planned probability of detecting an effect of the chosen size under the model. Effect size is the difference you want the study to be able to detect, adjustable above. Sample size is typically the output when the other pieces are fixed, but any one of these can be solved for when the rest are known, which is exactly what the table above does: solving for the n per group that reaches 80% power at your chosen effect size.

Worked example

For a two-group study wanting 80 percent power at alpha .05 to detect a standardized mean difference of .5 with equal group sizes, the required sample size is in the low 60s per group. If you expect 15 percent attrition, the recruitment target should be higher; that adjustment is a design choice, not a correction any calculator can infer for you.

Minimum detectable effect reverses the question

Minimum detectable effect asks: given the sample you can realistically obtain, what standardized or raw difference would the planned study have power to detect? This is often more useful than asking for a universal recommended sample size. A small study may still detect a very large effect, but it may be unable to distinguish smaller effects that matter in practice.

Power is conditional on your assumptions

An optimistic effect estimate, underestimated variability, or unrealistic event rate can make the required sample look smaller than it should be. Use a sensitivity view across several plausible effect sizes, the same reason this page lets you adjust the slider rather than showing one fixed number. Record the selected alternative direction, allocation ratio, confidence level, and attrition rule in your own planning notes.

Do not use post hoc power on a completed study

Once a study is complete, the observed estimate and its confidence interval are more direct descriptions of what the sample found and how uncertain it is. Power calculated backward from an already-observed result tends to track the p-value mechanically and adds little genuine information beyond it; use a power calculator to plan future data collection, not to explain a finished one.

Source

This page follows the power-analysis conventions documented by statsmodels and the shared statistical reasoning contract every StatReason engine is built against.

Limitations

This visual covers only the independent two-group mean-difference design at alpha .05; the engines linked below cover paired designs, proportions, correlations, and ANOVA designs with their own dedicated sample-size and power calculations.

Next action: use the [Sample Size Calculator](#) router for your specific design, or the [Minimum Detectable Effect Calculator](#) when your sample size is already fixed.