

CANCER RESEARCH CONNECTIONS

From the Biostatistics Shared Resource (BSR) at Case Comprehensive Cancer Center

Volume 1 | Spring 2026 | Published Quarterly | Editors: Mireya Díaz and Sujata Patil

FEATURE

Shared Resource Spotlight

Free Biostatistics Consultations for Researchers

Looking for personalized support in planning your research study? The Biostatistics Shared Resource at Case Comprehensive Cancer Center offers complimentary consultations to help you design a strong analytic plan tailored to your grant proposal. Whether you're working on a traditional clinical trial or exploring innovative designs like adaptive trials, master protocols, emulated trials, drug repurposing, or high-dimensional data analysis, our team is here to collaborate with you.

What We Offer

Complimentary, personalized consultations to help you design a rigorous analytic plan seamlessly integrated into your grant proposal. We also provide hands-on data analysis support, including for retrospective and observational studies across most data types and designs.

AREAS OF EXPERTISE

Our team supports a wide range of study designs and analytical approaches:

✓ Traditional or Adaptive Clinical Trials	✓ Retrospective & Observational Data Analysis
✓ Master Protocols	✓ Emulated / Pragmatic Trials
✓ Drug Repurposing Studies	✓ High-Dimensional Data Analysis

We'll listen to your ideas and help integrate your analytic plan seamlessly into your proposal. Our experienced team of PhD- and MS-level statisticians has partnered with hundreds of researchers over the years, contributing to successful studies funded by government agencies, foundations and nonprofits, and industry sponsors.

Don't wait to turn your idea into the next big discovery!

Schedule your appointment today and let us help you become our next success story!

Case Comprehensive Cancer Center

Book a consultation online:

[Schedule via REDCap →](#)

Cleveland Clinic Foundation

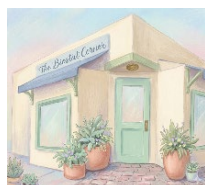
Email to book your appointment:

[patils2@ccf.org →](mailto:patils2@ccf.org)

CANCER RESEARCH CONNECTIONS

From the Biostatistics Shared Resource (BSR) at Case Comprehensive Cancer Center

Volume 1 | Spring 2026 | Published Quarterly | Editors: Mireya Díaz and Sujata Patil



THE BIOSTAT CORNER

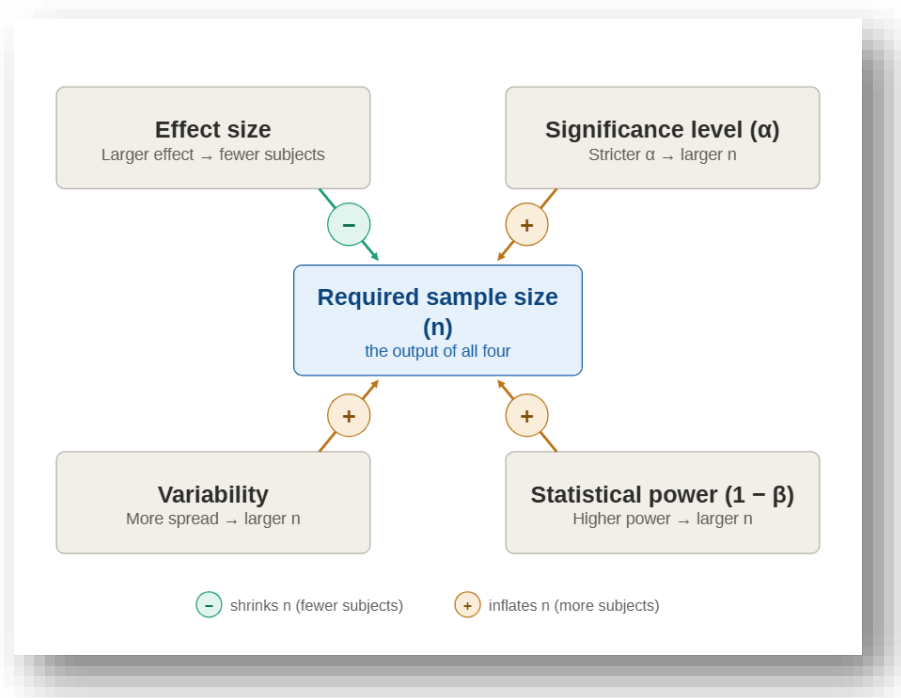
Expert Perspectives on Statistical Methods in Cancer Research

The Art and Science of Sample Size Estimation

Roughly 20% of trials in high-impact journals report negative results, and most were simply underpowered — unable to detect even large relative differences of 25–50%. A 2025 review found a comparable rate of trials with a disconnect between clinical and statistical significance. The usual culprits: choosing a convenient sample size instead of estimating one, making unrealistic assumptions, never testing those assumptions across a range of scenarios, and ignoring losses during follow-up.

In oncology phase III trials, 66% recruited at least 90% of their planned sample, yet only 58% reported the expected effect size for both arms — despite this being a CONSORT 2010 recommendation. Sound design, including a proper sample size or power calculation, makes a successful trial more likely, and reporting the parameters behind that calculation supports transparency and the FAIR principles. Many of these details can be found in trial protocols on ClinicalTrials.gov.

ELEMENTS OF SAMPLE SIZE CALCULATION



Sample size estimation is part science, part art. The science is probability theory; the art lies in adapting standard formulas to complex designs that have no ready-made solution.

Every calculation rests on four elements: the effect size, the variability of the measure used to capture it, the type I error (significance level), and the type II error (power). Outcome type (continuous,

CANCER RESEARCH CONNECTIONS

From the Biostatistics Shared Resource (BSR) at Case Comprehensive Cancer Center

Volume 1 | Spring 2026 | Published Quarterly | Editors: Mireya Díaz and Sujata Patil

categorical, time-to-event, longitudinal) and design (single-arm, parallel, cross-over) change the specific values, but the structure stays the same. These are the inputs a (bio)statistician will ask you for. One more design choice — whether your hypothesis is one-sided (directional) or two-sided — feeds into the type I error term.

Effect size and variability are often combined into the standardized effect size (effect size ÷ variability), or Cohen's d. By convention, 0.2 is small, 0.5 medium, and 0.8 large. Smaller effects demand larger samples — so always ask whether an effect that small is clinically, biologically, or practically meaningful. Outcome type matters too: continuous outcomes need the smallest samples, proportions more, and time-to-event outcomes the most.

Single-Arm Studies

Single-arm studies — common in observational work and later-phase trials — may be comparative or non-comparative. Comparative studies test whether the outcome matches a predefined value or historical control. Non-comparative (descriptive) studies instead seek a target precision.

Two-arm studies of superiority

Comparing two arms head-to-head is a superiority test: the alternative hypothesis is that the estimates (means, proportions, or hazard rates) differ. The mirror image — equivalence and non-inferiority — is a topic for a future issue.

Adjustments

Two adjustments come up constantly:

- Attrition. Participants lost before their outcome is observed shrink your effective sample. Inflate the estimate by dividing by $(1 - \text{attrition proportion})$; for example, with 10% attrition, divide by 0.9.
- Correlation. When outcomes are measured repeatedly in the same individuals, observations are correlated, which changes the variance. Multiply the variance by a variance inflation factor that captures the correlation structure: $VIF = 1 + (m - 1) \times ICC$. Other VIF may apply depending on the association structure.
- Two more, briefly: for unequal allocation, adjust the allocation-ratio term in the denominator; for multiple endpoints, apply a Bonferroni (or similar) correction to α .

ONLINE CALCULATORS

Online calculators are useful but not always correct. Verify any result by running it through two independent tools — or better, bring your design to the Biostatistics Shared Resources. That's what we're here for.

KEY REFERENCES

- Moher D, Dulberg CS, Wells GA. Statistical power, sample size, and their reporting in randomized controlled trials. *JAMA*. 1994;272(2):122–4.
- Esterhuizen TM, et al. Discordance between clinical and statistical significance. *BMJ Open*. 2025;15(8):e100411.
- Schulz KF, Altman DG, Moher D; CONSORT Group. CONSORT 2010 statement. *BMJ*. 2010;340:c332.
- Cohen J. *Statistical Power Analysis for the Behavioral Sciences*. 2nd ed. Hillsdale, NJ: Erlbaum; 1988.
- Lachin JM. Introduction to sample size determination in clinical trials. *Control Clin Trials*. 1981;2(2):93–113.
- Campbell MJ, Donner A, Klar N. Developments in cluster randomized trials. *Stat Med*. 2007;26:2–19.

Interested in a future topic? Reach out to the BSR — we'd love to hear from you!