

Commentary

Sample size in clinical research: why it matters and how it is determined

Nitin Chandurkar, Kumar Naidu*

Department of Clinical Research and Development, Ipca Laboratories Ltd. Mumbai, Maharashtra, India

Received: 10 April 2026

Revised: 11 July 2026

Accepted: 13 July 2026

*Correspondence:

Kumar Naidu,

E-mail: kumarnaidu2@gmail.com

Copyright: © the author(s), publisher and licensee Medip Academy. This is an open-access article distributed under the terms of the Creative Commons Attribution Non-Commercial License, which permits unrestricted non-commercial use, distribution, and reproduction in any medium, provided the original work is properly cited.

ABSTRACT

Sample size is a critical component of clinical research that directly influences the validity, reliability, and interpretability of study findings. Determining an appropriate sample size is not merely a statistical exercise but a multidisciplinary decision involving clinical relevance, study design, variability, and ethical considerations. An inadequately powered study may fail to detect meaningful treatment effects, whereas an excessively large study may expose unnecessary numbers of participants to interventions and increase study costs. This article provides a practical overview of why sample size matters in clinical research and explains the major factors that influence its determination. Examples from different study settings, including superiority, non-inferiority, and bioequivalence studies, are used to bridge statistical concepts with practical application. Common misconceptions surrounding sample size are also discussed to improve understanding among clinical research professionals.

Keywords: Sample size, Clinical trials, Biostatistics, Power, Variability, Study design

INTRODUCTION

Sample size is one of the most fundamental elements in the design of clinical study. It plays a central role in determining whether a study can answer its research question in a scientifically meaningful and reliable manner. Despite this, sample size is often perceived as a routine statistical output rather than a carefully reasoned component of study design. In practice, however, it is a key decision that has implications not only for statistical analysis but also for ethics, feasibility, timelines, costs, and regulatory acceptability.^{1,2}

In clinical research, every study is designed to answer a specific question, such as demonstrating superiority, establishing non-inferiority, or showing bioequivalence between treatments.

The required sample size varies depending on the study objective, endpoint, and expected treatment effect.³

A study with too few participants may fail to identify a true treatment effect even when one exists, while an excessively large study may unnecessarily expose participants to intervention and increase costs. Therefore, sample size determination is both a scientific and ethical necessity.⁴

WHY SAMPLE SIZE MATTERS

The importance of sample size in clinical research extends beyond statistical calculations and directly affects the credibility and usefulness of study findings.

A key consideration is the ability to detect a true treatment effect. Studies with insufficient sample size are prone to type II error, meaning that a real effect may go undetected.⁵ This can result in potentially effective treatments being incorrectly dismissed.

On the other hand, very large sample sizes may detect statistically significant differences that are not clinically

meaningful.⁶ This distinction between statistical significance and clinical relevance is critical in interpreting study results.

Ethical considerations also play an important role. Clinical trials involve human participants, and enrolling subjects in a study that is underpowered or unnecessarily large raises ethical concerns. Regulatory and ethical guidelines emphasize that studies should be designed to answer research questions efficiently without exposing participants to avoidable risk.⁷

From an operational perspective, sample size directly impacts timelines, number of sites, monitoring effort, and overall study cost. Thus, it is not only a statistical parameter but also a key driver of study feasibility.

KEY FACTORS THAT INFLUENCE SAMPLE SIZE

Sample size determination is influenced by several interrelated factors that reflect both statistical principles and clinical assumptions.

Effect size represents the magnitude of the difference between treatment groups and is one of the most influential determinants of sample size.¹ In clinical trials, effect size may be expressed as the difference in clinical outcomes, such as blood pressure reduction or survival rates. In bioequivalence studies, it is often expressed as the ratio of pharmacokinetic parameters like AUC/Cmax.

For example, if the expected response rate is 70% in one treatment group and 50% in the other, the effect size is a 20% absolute difference. Larger effect sizes are easier to detect statistically and therefore require smaller sample sizes, whereas smaller differences demand larger sample sizes, as illustrated in Figure 1.⁹

Another important factor is variability in data. Increased variability makes it harder to detect true differences between groups and therefore requires larger sample sizes, shown in Figure 2.⁸ This is particularly relevant in clinical endpoints with high inter/intra-subject variability.

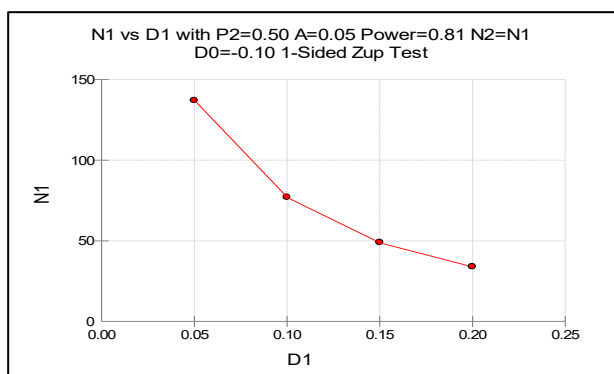


Figure 1: Changes in sample size with proportion differences.

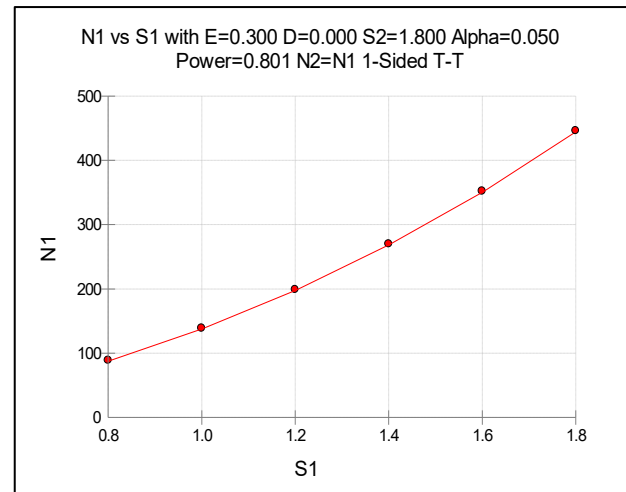


Figure 2: Change in sample size with variability.

The significance level (α) represents the probability of committing a type I error, which occurs when a study incorrectly concludes that a treatment effect exists when it actually does not.¹ In most biomedical research, α is set at 0.05, corresponding to a 95% confidence level. This means that there is a 5% chance of falsely declaring that treatment is effective when not effective. For example, in a clinical trial evaluating a new antihypertensive drug, $\alpha=0.05$ implies that 5 out of 100 similar studies might incorrectly show a beneficial effect purely by chance.

Reducing α to a more stringent level such as 0.01 increases confidence in the findings but also demands a larger sample size because stronger evidence is required to reject the null hypothesis.¹ Conversely, increasing α reduces the required sample size but compromises reliability. Therefore, selecting α involves balancing statistical rigor with feasibility.

Statistical power reflects the probability that a study will correctly detect a true treatment effect.¹ It is defined as $1-\beta$, where β is the probability of a type II error (failing to detect a real effect). Typically, studies are designed with 80% or 90% power. For instance, in a clinical trial investigating a new cancer therapy, a power of 80% means that the study has an 80% chance of identifying a true improvement in survival outcomes.

Higher power increases the likelihood of detecting meaningful differences but requires a larger sample size.² On the other hand, studies with lower power may be easier to conduct but risk producing inconclusive or misleading results. Thus, power is directly related to the sensitivity and credibility of the study, and its impact on sample size is demonstrated in Figure 3.³

The type of endpoint—continuous, binary, or time-to-event—also influences sample size calculations, as each requires different statistical approaches.⁴

Finally, practical considerations such as dropout rate must be incorporated to ensure that the final evaluable sample remains adequate. Failure to account for attrition can lead to underpowered analyses.

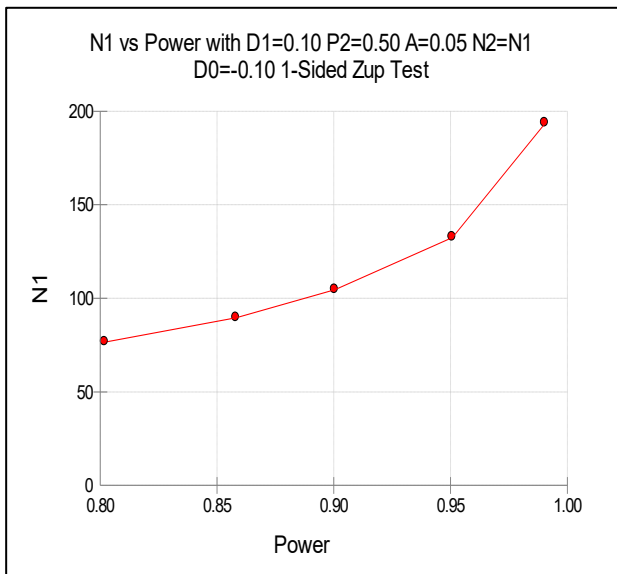


Figure 3: Changes in sample size with power.

HOW SAMPLE SIZE DIFFERS ACROSS STUDY DESIGNS

Sample size requirements can vary considerably depending on the study objective, endpoint, and design framework. This is because different study designs ask different scientific questions and therefore require different levels of statistical precision.

In superiority trials, the objective is to demonstrate that one treatment is better than another, such as a new drug showing greater reduction in HbA1c or blood pressure compared with placebo or standard therapy. In these studies, sample size is mainly driven by the expected treatment difference and the variability of the endpoint. If the expected difference is large, fewer subjects may be sufficient. However, if the expected improvement is modest, a larger sample size is required to reliably detect that difference.²

In non-inferiority trials, the objective is not necessarily to show that the new treatment is better, but rather that it is not unacceptably worse than an active comparator by more than a predefined margin. This margin represents the maximum difference that would still be considered clinically acceptable. Such designs are commonly used when the new treatment may offer other advantages, such as better tolerability, convenience, safety, or cost. In these studies, the choice of the non-inferiority margin becomes a critical driver of sample size. A narrower margin requires greater precision and therefore leads to a larger sample size Figure 4.⁹ In practical terms, the smaller the allowed loss in efficacy, the larger the study must be.

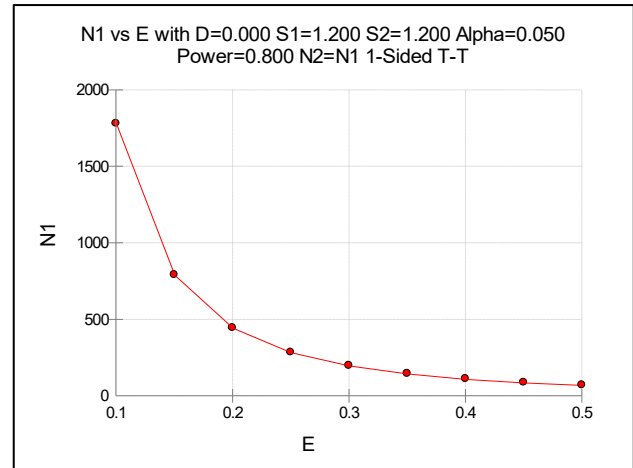


Figure 4: Change in sample size with margin.

In bioequivalence studies, the objective is to demonstrate that the pharmacokinetic performance of a test formulation is comparable to that of a reference product. Unlike efficacy trials, sample size in bioequivalence studies is often driven less by treatment effect and more by intra-subject variability, particularly for pharmacokinetic parameters such as C_{max} and AUC. Drugs with low variability may require only a modest number of subjects, whereas highly variable drugs can require substantially larger sample sizes. In such situations, regulatory guidance may allow or encourage the use of replicate study designs or reference-scaled approaches to address high variability more efficiently.¹⁰

For studies with binary endpoints, such as responder/non-responder, healed/not healed, or event/no event outcomes, sample size depends strongly on the assumed response rates in the treatment groups. Even relatively small changes in these assumed proportions can have a substantial impact on the required sample size.⁴ For example, a study comparing response rates of 70% versus 50% may require a manageable sample size, whereas a comparison of 70% versus 65% may require many more subjects.

Similarly, studies with time-to-event endpoints, such as survival, recurrence, or progression-free survival, are often driven not only by the number of patients enrolled but also by the number of events observed. In such studies, sample size planning depends on expected event rates, follow-up duration, and dropout patterns. As a result, two studies enrolling the same number of patients may still differ in statistical power depending on how many events occur during the study period.

Thus, sample size is highly context-dependent and cannot be generalized across all studies. The same number of subjects may be adequate for one design but completely insufficient for another. For this reason, sample size should always be interpreted within the framework of the study objective, endpoint, variability, and design characteristics.

PRACTICAL EXAMPLES FROM CLINICAL RESEARCH

The practical implications of sample size determination are best understood through examples.

In a diabetes study evaluating HbA1c reduction, a difference of 0.5% with a standard deviation of 1.0% may require a moderate sample size. However, if variability increases, the required sample size increases substantially, highlighting the sensitivity of sample size to variability.²

In bioequivalence studies, drugs with low intra-subject variability may require relatively small sample sizes, whereas highly variable drugs may require significantly larger studies.¹⁰

In non-inferiority trials, reducing the non-inferiority margin increases the required sample size, reflecting the need for greater precision.⁹

These examples illustrate that sample size is driven not only by statistical formulas but also by clinical assumptions and study objectives.

COMMON MISCONCEPTIONS ABOUT SAMPLE SIZE

Several misconceptions about sample size persist in clinical research.

One common belief is that larger sample sizes are always better, which is not necessarily true. Excessively large studies may detect trivial differences that lack clinical significance.⁶

Another misconception is that sample size is purely a statistical decision. In reality, it requires clinical judgment to define meaningful effect sizes and appropriate study assumptions.³

Pilot studies are often misunderstood in sample size planning. While pilot studies play an important role—particularly in bioequivalence and early-phase research by providing preliminary information on treatment performance and variability, their results must be interpreted with caution. Due to their typically small sample size, pilot studies may produce unstable estimates of variability or treatment effect. Over-reliance on such estimates can lead to underpowered or overestimated sample size calculations.

Ignoring dropout rates is another common mistake that can compromise study power if not appropriately accounted for during planning.

In practice, it is also observed that some studies are conducted with inadequate sample sizes due to optimistic assumptions, such as overestimating treatment

differences or using overly lenient margins. Such assumptions may lead to insufficient power and inconclusive results.

Finally, statistical significance is often mistaken for clinical importance, which may lead to misinterpretation of study findings.

CONCLUSION

Sample size determination is a critical component of clinical research that brings together statistical principles, clinical relevance, ethical responsibility, and operational feasibility. It plays a central role in ensuring that a study can answer its research questions with sufficient confidence and precision.

A well-justified sample size helps to strike the right balance between scientific rigor and practical feasibility. If the sample size is too small, the study may fail to detect a true treatment effect and produce inconclusive results. If it is unnecessarily large, it may lead to avoidable cost, complexity, and participant burden. Therefore, determining the right sample size is not simply a technical requirement, but a key decision that influences the overall quality and value of a clinical study.

As discussed throughout this article, sample size is influenced by several important factors, including the study objective, expected treatment difference, variability, statistical power, significance level, study design, and anticipated dropout rate. These factors must be considered carefully and interpreted in the context of the specific research question.

A sound understanding of sample size principles is essential not only for statisticians, but also for clinicians, clinical operations teams, medical writers, regulatory professionals, and all others involved in study planning and interpretation. Greater awareness of these concepts can contribute to better-designed studies, more efficient use of resources, and more meaningful interpretation of study findings.

Ultimately, an appropriate sample size helps ensure that clinical research is not only statistically robust, but also clinically relevant, ethically justified, and decision-enabling.

Funding: No funding sources

Conflict of interest: None declared

Ethical approval: Not required

REFERENCES

1. Chow SC, Shao J, Wang H, Lokhnygina Y. *Sample Size Calculations in Clinical Research*. 3rd ed. Boca Raton, FL: CRC Press; 2017: 1–510.

2. Julious SA. Sample Sizes for Clinical Trials. Boca Raton, FL: CRC Press; 2009: 1–358.
3. Friedman LM, Furberg CD, DeMets DL, Reboussin DM, Granger CB. Fundamentals of Clinical Trials. 5th ed. Cham: Springer International Publishing; 2015: 1–560.
4. Piantadosi S. Clinical Trials: A Methodologic Perspective. 3rd ed. Hoboken, NJ: John Wiley & Sons; 2017: 1–816.
5. Cohen J. Statistical Power Analysis for the Behavioral Sciences. 2nd ed. Hillsdale, NJ: Lawrence Erlbaum Associates; 1988: 1–567.
6. Pocock SJ. Clinical Trials: A Practical Approach. Chichester: John Wiley & Sons; 1983: 1–266.
7. ICH E9 Expert Working Group. Statistical Principles for Clinical Trials. 1998. Available at: <https://www.ema.europa.eu/en/ich-e9-statistical-principles-clinical-trials-scientific-guideline>. Accessed on 3 June 2026.
8. Julious SA. Sample size of 12 per group rule of thumb for a pilot study. *Pharm Stat*. 2005;4(4):287-91.
9. Machin D, Campbell MJ, Tan SB, Tan SH. Sample Size Tables for Clinical Studies. 3rd ed. Wiley-Blackwell. 2009.
10. European Medicines Agency. Guideline on the Investigation of Bioequivalence. 2010. Available at: <https://www.ema.europa.eu/en/investigation-bioequivalence-scientific-guideline>. Accessed on 3 June 2026.

Cite this article as: Chandurkar N, Naidu K. Sample size in clinical research: why it matters and how it is determined. *Int J Clin Trials* 2026;13(3):363-7.