

Understanding the Lower and Upper Limits of Sample Sizes in Clinical Research

Abdullah A. Hamad ^{1, 2}, Sirwan K. Ahmed ³

1. Faculty of Medicine, Menoufia University, Shibin El-Kom, EGY 2. Medical Research Group of Egypt, Negida Academy, Arlington, USA 3. College of Nursing, University of Raparin, Rania, IRQ

Corresponding author: Sirwan K. Ahmed, sirwan.k.ahmed@gmail.com

Review began 12/13/2024

Review ended 12/25/2024

Published 01/01/2025

© Copyright 2025

Hamad et al. This is an open access article distributed under the terms of the Creative Commons Attribution License CC-BY 4.0., which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

DOI: 10.7759/cureus.76724

Abstract

The determination of an appropriate sample size is pivotal in medical research, not only for achieving statistical adequacy but also for ensuring ethical integrity and resource efficiency. This editorial elucidates the complexities involved in defining the lower and upper limits of sample size across various research paradigms. The lower limit is essential for maintaining sufficient statistical power and precision, which depend on several factors, including the study's objectives, inherent population variability, and the desired accuracy of results. An insufficient sample size poses a risk of significant Type II errors and produces wide confidence intervals, thereby undermining the reliability and applicability of the research findings. Conversely, the upper limit encounters practical constraints related to resource allocation and ethical considerations, where the principle of diminishing returns becomes evident as sample sizes increase beyond a certain threshold. This scenario leads to minimal gains in precision at the cost of potential participant risk and resource overutilization. The editorial advocates for a methodical approach to sample size calculation, utilizing statistical tools such as sample size formulas, and G*Power and adopting innovative methodologies, including adaptive trial designs and Bayesian statistics. These strategies facilitate dynamic adjustments based on interim results and prior knowledge, respectively, promoting optimal resource utilization while preserving robust statistical power. Ultimately, the careful calibration of sample size enhances the validity and ethical integrity of medical research, thereby bolstering its contribution to scientific knowledge.

Categories: Public Health, Epidemiology/Public Health, Medical Education

Keywords: clinical research, education, reliability, research method, sample size, sample size and power computation, sample size estimation, sample size formula, sampling, validity

Editorial

In medical research, studying an entire population is often impractical. To address this, researchers rely on samples that aim to represent the population and allow for the generalization of findings [1]. Determining the appropriate sample size is a critical component of study design. A sample that is too small may fail to detect significant effects, while an excessively large sample can result in wasted resources and ethical challenges. Understanding the concepts of upper and lower limits of sample size is essential for conducting ethical, efficient, and scientifically robust research.

The lower limit of sample size ensures adequate statistical power and precision. However, there is no universal minimum sample size that applies to all studies. The appropriate lower limit depends on factors such as study objectives, population variability, and the desired level of precision [2]. A sample size that is too small can lead to significant limitations in research outcomes. Low statistical power reduces the ability to detect true effects, thereby increasing the risk of Type II errors (false negatives). Additionally, small sample sizes often produce wide confidence intervals, resulting in imprecise parameter estimates that limit the applicability and reliability of findings. Lastly, insufficient sample sizes increase the vulnerability of results to random error and potential bias, making conclusions less dependable. However, small sample sizes are sometimes necessary for specific studies, such as Phase 1 drug trials, where the primary goal is to assess initial safety rather than therapeutic effects, and formal sample size calculations may not be required.

Key factors to consider in determining an appropriate sample size include statistical power, Type I error, and Type II error [3]. Statistical power, typically set at 80% (0.8) in medical research, represents the probability of detecting a true effect when one exists. Type I error (α) is the risk of rejecting a true null hypothesis and is conventionally set at 5% (0.05), corresponding to a 95% confidence level. Type II error (β) is the risk of failing to reject a false null hypothesis and is inversely related to statistical power, calculated as $\beta = 1 - \text{power}$. Ensuring these parameters are balanced is critical for producing valid and reliable results. To assist in sample size calculation, researchers can use tools like the GPower software, which provides calculations based on study design and statistical tests, or apply sample size formulas that consider factors such as power, error rates, and expected effect size [2].

In contrast to the focus on lower sample size limits, discussions on the upper limit are less common. Larger

How to cite this article

Hamad A A, Ahmed S K (January 01, 2025) Understanding the Lower and Upper Limits of Sample Sizes in Clinical Research. Cureus 17(1): e76724. DOI 10.7759/cureus.76724

sample sizes generally enhance precision by minimizing sampling error, but practical, ethical, and economic considerations impose constraints. The principle of diminishing returns applies, where the marginal gain in precision decreases as the sample size grows. Beyond a certain point, additional participants contribute minimally to the accuracy of results [4]. Although this threshold is not a fixed point, it signifies a stage where the accuracy returns diminish progressively with oversampling (Figure 1). Resource constraints also play a significant role, as larger studies demand more time, funding, and personnel, potentially straining budgets and delaying results. Ethical considerations further underscore the need for restraint, as enrolling more participants than necessary exposes individuals to risks without adding meaningful scientific value. For example, in a hypertension drug trial, oversampling may expose participants to unnecessary side effects without improving the study's reliability, wasting resources and potentially harming vulnerable populations. Similarly, a small sample size that fails to detect a true effect is also considered unethical. This ethical dilemma emphasizes the importance of precise sample size calculation to ensure that research maximizes benefits while minimizing unnecessary harm and resource waste [5].

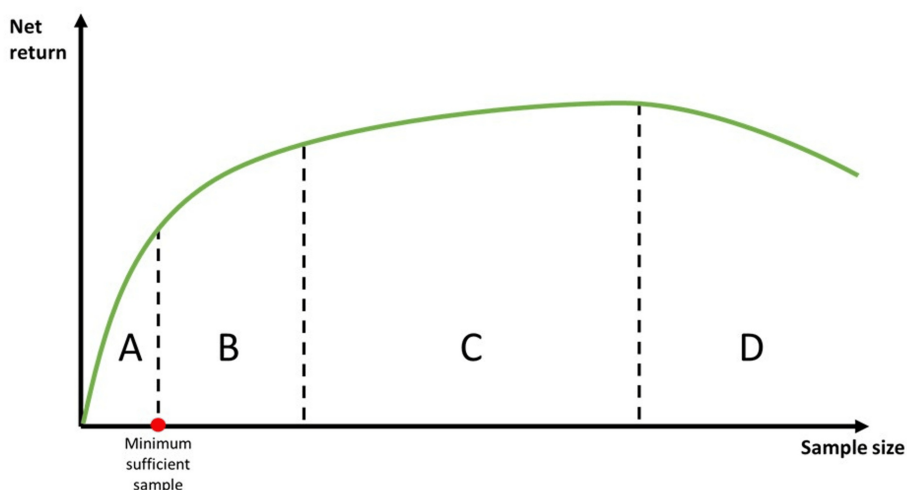


FIGURE 1: Demonstrates the relationship between sample size and return in research studies.

In area A, sample sizes are insufficient, lacking the statistical power required to reliably detect the effect estimate, thereby increasing the risk of Type II errors. In area B, sample sizes achieve adequate statistical power, with further increases significantly improving precision. This may be necessary to address weak sampling methods, account for variability, or adjust for potential confounders. Beyond this lies area C, marking the *diminishing returns* threshold, where additional increases in sample size yield minimal improvements in the precision or accuracy of results. Finally, in area D, increasing the sample size further leads to negative returns due to disproportionate costs, time, and resource demands, along with ethical concerns about unnecessarily exposing additional participants to study procedures for only marginal gains in accuracy.

Image credit: Abdullah A. Hamad.

Moreover, oversampling in intervention studies can lead to the detection of statistically significant but clinically irrelevant differences between groups. This may overstate the intervention's effectiveness and diminish its applicability to real-world contexts. For instance, while larger samples reduce sampling error, they can also overinflate trivial differences, undermining the practical relevance of findings.

Therefore, determining the optimal sample size requires balancing statistical rigor, practicality, and ethical principles. While the lower limit ensures sufficient power and precision, the upper limit reflects the constraints of diminishing returns, resource availability, and participant welfare. Researchers should prioritize obtaining a sufficient, representative sample over maximizing participant numbers. Employing robust sampling techniques, such as stratified or random sampling, can further enhance the quality of research outcomes without unnecessary participant recruitment. Additionally, newer approaches such as adaptive trial designs and Bayesian statistics offer innovative ways to optimize sample size determination. Adaptive trial designs adjust the sample size dynamically based on interim results, allowing researchers to optimize resources and maintain statistical power. Bayesian statistics use prior knowledge to iteratively update sample size decisions as new data is collected, allowing for more efficient resource use.

Additional Information

Author Contributions

All authors have reviewed the final version to be published and agreed to be accountable for all aspects of the

work.

Concept and design: Sirwan K. Ahmed, Abdullah A. Hamad

Acquisition, analysis, or interpretation of data: Sirwan K. Ahmed, Abdullah A. Hamad

Drafting of the manuscript: Sirwan K. Ahmed, Abdullah A. Hamad

Critical review of the manuscript for important intellectual content: Sirwan K. Ahmed, Abdullah A. Hamad

Supervision: Sirwan K. Ahmed

Disclosures

Conflicts of interest: In compliance with the ICMJE uniform disclosure form, all authors declare the following: **Payment/services info:** All authors have declared that no financial support was received from any organization for the submitted work. **Financial relationships:** All authors have declared that they have no financial relationships at present or within the previous three years with any organizations that might have an interest in the submitted work. **Other relationships:** All authors have declared that there are no other relationships or activities that could appear to have influenced the submitted work.

Acknowledgements

While preparing this work, the authors utilized ChatGPT (GPT-4o) to enhance the writing quality. After using this tool, the authors revised and edited the content as necessary and assumed full responsibility for the publication's content.

References

1. Elfil M, Negida A: Sampling methods in clinical research; an educational review . Arch Acad Emerg Med. 2017, 5:e52. [10.22037/aaem.v5i1.177](https://doi.org/10.22037/aaem.v5i1.177)
2. Ahmed SK: How to choose a sampling technique and determine sample size for research: a simplified guide for researchers. Oral Oncol Rep. 2024, 12:100662. [10.1016/j.oor.2024.100662](https://doi.org/10.1016/j.oor.2024.100662)
3. Wang X, Ji X: Sample size estimation in clinical research: from randomized controlled trials to observational studies. Chest. 2020, 158:S12-20. [10.1016/j.chest.2020.03.010](https://doi.org/10.1016/j.chest.2020.03.010)
4. Bacchetti P, Deeks SG, McCune JM: Breaking free of sample size dogma to perform innovative translational research. Sci Transl Med. 2011, 3:87ps24. [10.1126/scitranslmed.3001628](https://doi.org/10.1126/scitranslmed.3001628)
5. Andrade C: Sample size and its importance in research . Indian J Psychol Med. 2020, 42:102-3. [10.4103/IJPSYM.IJPSYM_504_19](https://doi.org/10.4103/IJPSYM.IJPSYM_504_19)