

Access

To read this article in full you may need to log in, make a payment or gain access through a site license (see right).

[nature.com](#) > [Journal home](#) > [Table of Contents](#)

Analysis

Nature Reviews Neuroscience **14**, 365–376 (May 2013) | doi:10.1038/nrn3475

There is an [Erratum \(May 1 2013\)](#) associated with this article.

Power failure: why small sample size undermines the reliability of neuroscience

Katherine S. Button, John P. A. Ioannidis, Claire Mokrysz, Brian A. Nosek, Jonathan Flint, Emma S. J. Robinson & Marcus R. Munafò

A study with low statistical power has a reduced chance of detecting a true effect, but it is less well appreciated that low power also reduces the likelihood that a statistically significant result reflects a true effect. Here, we show that the average statistical power of studies in the neurosciences is very low. The consequences of this include overestimates of effect size and low reproducibility of results. There are also ethical dimensions to this problem, as unreliable research is inefficient and wasteful. Improving reproducibility in neuroscience is a key priority and requires attention to well-established but often ignored methodological principles.

To read this article in full you may need to log in, make a payment or gain access through a site license (see right).

ARTICLE TOOLS

-  Send to a friend
-  Export citation
-  Export references
-  Rights and permissions
-  Order commercial reprints

SEARCH PUBMED FOR

- ▶ Katherine S. Button
- ▶ John P. A. Ioannidis
- ▶ Claire Mokrysz
- ▶ Brian A. Nosek
- ▶ Jonathan Flint
- ▶ Emma S. J. Robinson
- ▶ [more authors of this article](#)