

## RESEARCH ARTICLE SUMMARY

## PSYCHOLOGY

# Estimating the reproducibility of psychological science

Open Science Collaboration\*

**INTRODUCTION:** Reproducibility is a defining feature of science, but the extent to which it characterizes current research is unknown. Scientific claims should not gain credence because of the status or authority of their originator but by the replicability of their supporting evidence. Even research of exemplary quality may have irreproducible empirical findings because of random or systematic error.

**RATIONALE:** There is concern about the rate and predictors of reproducibility, but limited evidence. Potentially problematic practices include selective reporting, selective analysis, and insufficient specification of the conditions necessary or sufficient to obtain the results. Direct replication is the attempt to recreate the conditions believed sufficient for obtaining a pre-

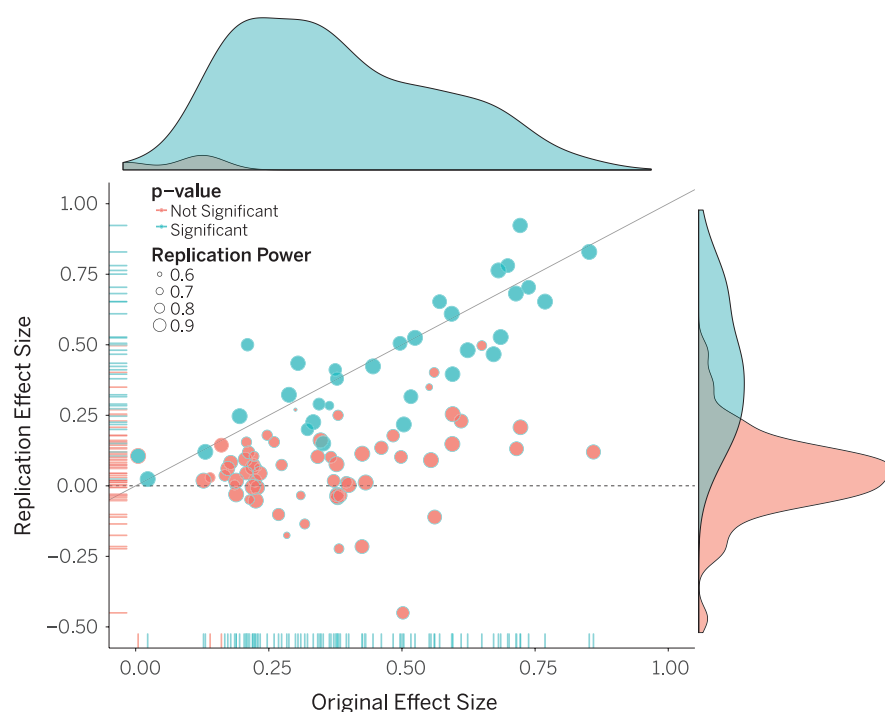
viously observed finding and is the means of establishing reproducibility of a finding with new data. We conducted a large-scale, collaborative effort to obtain an initial estimate of the reproducibility of psychological science.

**RESULTS:** We conducted replications of 100 experimental and correlational studies published in three psychology journals using high-powered designs and original materials when available. There is no single standard for evaluating replication success. Here, we evaluated reproducibility using significance and *P* values, effect sizes, subjective assessments of replication teams, and meta-analysis of effect sizes. The mean effect size (*r*) of the replication effects ( $M_r = 0.197$ ,  $SD = 0.257$ ) was half the magnitude of the mean effect size of the original effects ( $M_o = 0.403$ ,  $SD = 0.188$ ), representing a

substantial decline. Ninety-seven percent of original studies had significant results ( $P < .05$ ). Thirty-six percent of replications had significant results; 47% of original effect sizes were in the 95% confidence interval of the replication effect size; 39% of effects were subjectively rated to have replicated the original result; and if no bias in original results is assumed, combining original and replication results left 68% with statistically significant effects. Correlational tests suggest that replication success was better predicted by the strength of original evidence than by characteristics of the original and replication teams.

**CONCLUSION:** No single indicator sufficiently describes replication success, and the five indicators examined here are not the only ways to evaluate reproducibility. Nonetheless, collectively these results offer a clear conclusion: A large portion of replications produced weaker evidence for the original findings despite using materials provided by the original authors, review in advance for methodological fidelity, and high statistical power to detect the original effect sizes. Moreover, correlational evidence is consistent with the conclusion that variation in the strength of initial evidence (such as original *P* value) was more predictive of replication success than variation in the characteristics of the teams conducting the research (such as experience and expertise). The latter factors certainly can influence replication success, but they did not appear to do so here.

Reproducibility is not well understood because the incentives for individual scientists prioritize novelty over replication. Innovation is the engine of discovery and is vital for a productive, effective scientific enterprise. However, innovative ideas become old news fast. Journal reviewers and editors may dismiss a new test of a published idea as unoriginal. The claim that “we already know this” belies the uncertainty of scientific evidence. Innovation points out paths that are possible; replication points out paths that are likely; progress relies on both. Replication can increase certainty when findings are reproduced and promote innovation when they are not. This project provides accumulating evidence for many findings in psychological research and suggests that there is still more work to do to verify whether we know what we think we know. ■



**Original study effect size versus replication effect size (correlation coefficients).** Diagonal line represents replication effect size equal to original effect size. Dotted line represents replication effect size of 0. Points below the dotted line were effects in the opposite direction of the original. Density plots are separated by significant (blue) and nonsignificant (red) effects.

The list of author affiliations is available in the full article online.  
\*Corresponding author. E-mail: nosek@virginia.edu  
Cite this article as Open Science Collaboration, *Science* 349, aac4716 (2015). DOI: 10.1126/science.aac4716

## RESEARCH ARTICLE

## PSYCHOLOGY

# Estimating the reproducibility of psychological science

Open Science Collaboration\*†

Reproducibility is a defining feature of science, but the extent to which it characterizes current research is unknown. We conducted replications of 100 experimental and correlational studies published in three psychology journals using high-powered designs and original materials when available. Replication effects were half the magnitude of original effects, representing a substantial decline. Ninety-seven percent of original studies had statistically significant results. Thirty-six percent of replications had statistically significant results; 47% of original effect sizes were in the 95% confidence interval of the replication effect size; 39% of effects were subjectively rated to have replicated the original result; and if no bias in original results is assumed, combining original and replication results left 68% with statistically significant effects. Correlational tests suggest that replication success was better predicted by the strength of original evidence than by characteristics of the original and replication teams.

Reproducibility is a core principle of scientific progress (1–6). Scientific claims should not gain credence because of the status or authority of their originator but by the replicability of their supporting evidence. Scientists attempt to transparently describe the methodology and resulting evidence used to support their claims. Other scientists agree or disagree whether the evidence supports the claims, citing theoretical or methodological reasons or by collecting new evidence. Such debates are meaningless, however, if the evidence being debated is not reproducible.

Even research of exemplary quality may have irreproducible empirical findings because of random or systematic error. Direct replication is the attempt to recreate the conditions believed sufficient for obtaining a previously observed finding (7, 8) and is the means of establishing reproducibility of a finding with new data. A direct replication may not obtain the original result for a variety of reasons: Known or unknown differences between the replication and original study may moderate the size of an observed effect, the original result could have been a false positive, or the replication could produce a false negative. False positives and false negatives provide misleading information about effects, and failure to identify the necessary and sufficient conditions to reproduce a finding indicates an incomplete theoretical understanding. Direct replication provides the opportunity to assess and improve reproducibility.

There is plenty of concern (9–13) about the rate and predictors of reproducibility but limited evidence. In a theoretical analysis, Ioannidis estimated that publishing and analytic practices make it likely that more than half of research

results are false and therefore irreproducible (9). Some empirical evidence supports this analysis. In cell biology, two industrial laboratories reported success replicating the original results of landmark findings in only 11 and 25% of the attempted cases, respectively (10, 11). These numbers are stunning but also difficult to interpret because no details are available about the studies, methodology, or results. With no transparency, the reasons for low reproducibility cannot be evaluated.

Other investigations point to practices and incentives that may inflate the likelihood of obtaining false-positive results in particular or irreproducible results more generally. Potentially problematic practices include selective reporting, selective analysis, and insufficient specification of the conditions necessary or sufficient to obtain the results (12–23). We were inspired to address the gap in direct empirical evidence about reproducibility. In this Research Article, we report a large-scale, collaborative effort to obtain an initial estimate of the reproducibility of psychological science.

## Method

Starting in November 2011, we constructed a protocol for selecting and conducting high-quality replications (24). Collaborators joined the project, selected a study for replication from the available studies in the sampling frame, and were guided through the replication protocol. The replication protocol articulated the process of selecting the study and key effect from the available articles, contacting the original authors for study materials, preparing a study protocol and analysis plan, obtaining review of the protocol by the original authors and other members within the present project, registering the protocol publicly, conducting the replication, writing the final report, and auditing the process and analysis for quality control. Project coordinators

facilitated each step of the process and maintained the protocol and project resources. Replication materials and data were required to be archived publicly in order to maximize transparency, accountability, and reproducibility of the project (<https://osf.io/ezcuq>).

In total, 100 replications were completed by 270 contributing authors. There were many different research designs and analysis strategies in the original research. Through consultation with original authors, obtaining original materials, and internal review, replications maintained high fidelity to the original designs. Analyses converted results to a common effect size metric [correlation coefficient ( $r$ )] with confidence intervals (CIs). The units of analysis for inferences about reproducibility were the original and replication study effect sizes. The resulting open data set provides an initial estimate of the reproducibility of psychology and correlational data to support development of hypotheses about the causes of reproducibility.

## Sampling frame and study selection

We constructed a sampling frame and selection process to minimize selection biases and maximize generalizability of the accumulated evidence. Simultaneously, to maintain high quality, within this sampling frame we matched individual replication projects with teams that had relevant interests and expertise. We pursued a quasi-random sample by defining the sampling frame as 2008 articles of three important psychology journals: *Psychological Science* (PSC), *Journal of Personality and Social Psychology* (JPSP), and *Journal of Experimental Psychology: Learning, Memory, and Cognition* (JEP: LMC). The first is a premier outlet for all psychological research; the second and third are leading disciplinary-specific journals for social psychology and cognitive psychology, respectively [more information is available in (24)]. These were selected a priori in order to (i) provide a tractable sampling frame that would not plausibly bias reproducibility estimates, (ii) enable comparisons across journal types and sub-disciplines, (iii) fit with the range of expertise available in the initial collaborative team, (iv) be recent enough to obtain original materials, (v) be old enough to obtain meaningful indicators of citation impact, and (vi) represent psychology sub-disciplines that have a high frequency of studies that are feasible to conduct at relatively low cost.

The first replication teams could select from a pool of the first 20 articles from each journal, starting with the first article published in the first 2008 issue. Project coordinators facilitated matching articles with replication teams by interests and expertise until the remaining articles were difficult to match. If there were still interested teams, then another 10 articles from one or more of the three journals were made available from the sampling frame. Further, project coordinators actively recruited teams from the community with relevant experience for particular articles. This approach balanced competing goals: minimizing selection bias by

\*All authors with their affiliations appear at the end of this paper.

†Corresponding author. E-mail: [nosek@virginia.edu](mailto:nosek@virginia.edu)

having only a small set of articles available at a time and matching studies with replication teams' interests, resources, and expertise.

By default, the last experiment reported in each article was the subject of replication. This decision established an objective standard for study selection within an article and was based on the intuition that the first study in a multiple-study article (the obvious alternative selection strategy) was more frequently a preliminary demonstration. Deviations from selecting the last experiment were made occasionally on the basis of feasibility or recommendations of the original authors. Justifications for deviations were reported in the replication reports, which were made available on the Open Science Framework (OSF) (<http://osf.io/ezcu1>). In total, 84 of the 100 completed replications (84%) were of the last reported study in the article. On average, the to-be-replicated articles contained 2.99 studies ( $SD = 1.78$ ) with the following distribution: 24 single study, 24 two studies, 18 three studies, 13 four studies, 12 five studies, 9 six or more studies. All following summary statistics refer to the 100 completed replications.

For the purposes of aggregating results across studies to estimate reproducibility, a key result from the selected experiment was identified as the focus of replication. The key result had to be represented as a single statistical inference test or an effect size. In most cases, that test was a  $t$  test,  $F$  test, or correlation coefficient. This effect was identified before data collection or analysis and was presented to the original authors as part of the design protocol for critique. Original authors occasionally suggested that a different effect be used, and by default, replication teams deferred to original authors' judgments. Nonetheless, because the single effect came from a single study, it is not necessarily the case that the identified effect was central to the overall aims of the article. In the individual replication reports and subjective assessments of replication outcomes, more than a single result could be examined, but only the result of the single effect was considered in the aggregate analyses [additional details of the general protocol and individual study methods are provided in the supplementary materials and (25)].

In total, there were 488 articles in the 2008 issues of the three journals. One hundred fifty-eight of these (32%) became eligible for selection for replication during the project period, between November 2011 and December 2014. From those, 111 articles (70%) were selected by a replication team, producing 113 replications. Two articles had two replications each (supplementary materials). And 100 of those (88%) replications were completed by the project deadline for inclusion in this aggregate report. After being claimed, some studies were not completed because the replication teams ran out of time or could not devote sufficient resources to completing the study. By journal, replications were completed for 39 of 64 (61%) articles from *PSCI*, 31 of 55 (56%) articles from *JPSP*, and 28 of 39 (72%) articles from *JEP:LMC*.

The most common reasons for failure to match an article with a team were feasibility constraints for conducting the research. Of the 47 articles from the eligible pool that were not claimed, six (13%) had been deemed infeasible to replicate because of time, resources, instrumentation, dependence on historical events, or hard-to-access samples. The remaining 41 (87%) were eligible but not claimed. These often required specialized samples (such as macaques or people with autism), resources (such as eye tracking machines or functional magnetic resonance imaging), or knowledge making them difficult to match with teams.

### Aggregate data preparation

Each replication team conducted the study, analyzed their data, wrote their summary report, and completed a checklist of requirements for sharing the materials and data. Then, independent reviewers and analysts conducted a project-wide audit of all individual projects, materials, data, and reports. A description of this review is available on the OSF (<https://osf.io/xtine>). Moreover, to maximize reproducibility and accuracy, the analyses for every replication study were reproduced by another analyst independent of the replication team using the R statistical programming language and a standardized analytic format. A controller R script was created to regenerate the entire analysis of every study and recreate the master data file. This R script, available at <https://osf.io/fkmwg>, can be executed to reproduce the results of the individual studies. A comprehensive description of this reanalysis process is available publicly (<https://osf.io/a2eyg>).

### Measures and moderators

We assessed features of the original study and replication as possible correlates of reproducibility and conducted exploratory analyses to inspire further investigation. These included characteristics of the original study such as the publishing journal; original effect size,  $P$  value, and sample size; experience and expertise of the original research team; importance of the effect, with indicators such as the citation impact of the article; and rated surprisingness of the effect. We also assessed characteristics of the replication such as statistical power and sample size, experience and expertise of the replication team, independently assessed challenge of conducting an effective replication, and self-assessed quality of the replication effort. Variables such as the  $P$  value indicate the statistical strength of evidence given the null hypothesis, and variables such as "effect surprisingness" and "expertise of the team" indicate qualities of the topic of study and the teams studying it, respectively. The master data file, containing these and other variables, is available for exploratory analysis (<https://osf.io/5wup8>).

It is possible to derive a variety of hypotheses about predictors of reproducibility. To reduce the likelihood of false positives due to many tests, we aggregated some variables into summary indicators: experience and expertise of original team, experience and expertise of replication team, challenge of replication, self-assessed quality of repli-

cation, and importance of the effect. We had no a priori justification to give some indicators stronger weighting over others, so aggregates were created by standardizing [mean ( $M$ ) = 0,  $SD = 1$ ] the individual variables and then averaging to create a single index. In addition to the publishing journal and subdiscipline, potential moderators included six characteristics of the original study and five characteristics of the replication (supplementary materials).

### Publishing journal and subdiscipline

Journals' different publishing practices may result in a selection bias that covaries with reproducibility. Articles from three journals were made available for selection: *JPSP* ( $n = 59$  articles), *JEP:LMC* ( $n = 40$  articles), and *PSCI* ( $n = 68$  articles). From this pool of available studies, replications were selected and completed from *JPSP* ( $n = 32$  studies), *JEP:LMC* ( $n = 28$  studies), and *PSCI* ( $n = 40$  studies) and were coded as representing cognitive ( $n = 43$  studies) or social-personality ( $n = 57$  studies) subdisciplines. Four studies that would ordinarily be understood as "developmental psychology" because of studying children or infants were coded as having a cognitive or social emphasis. Reproducibility may vary by subdiscipline in psychology because of differing practices. For example, within-subjects designs are more common in cognitive than social psychology, and these designs often have greater power to detect effects with the same number of participants.

### Statistical analyses

There is no single standard for evaluating replication success (25). We evaluated reproducibility using significance and  $P$  values, effect sizes, subjective assessments of replication teams, and meta-analyses of effect sizes. All five of these indicators contribute information about the relations between the replication and original finding and the cumulative evidence about the effect and were positively correlated with one another ( $r$  ranged from 0.22 to 0.96, median  $r = 0.57$ ). Results are summarized in Table 1, and full details of analyses are in the supplementary materials.

### Significance and $P$ values

Assuming a two-tailed test and significance or  $\alpha$  level of 0.05, all test results of original and replication studies were classified as statistically significant ( $P \leq 0.05$ ) and nonsignificant ( $P > 0.05$ ). However, original studies that interpreted nonsignificant  $P$  values as significant were coded as significant (four cases, all with  $P$  values  $< 0.06$ ). Using only the nonsignificant  $P$  values of the replication studies and applying Fisher's method (26), we tested the hypothesis that these studies had "no evidential value" (the null hypothesis of zero-effect holds for all these studies). We tested the hypothesis that the proportions of statistically significant results in the original and replication studies are equal using the McNemar test for paired nominal data and calculated a CI of the reproducibility parameter. Second, we compared the central tendency of the distribution of  $P$  values of original and

**Table 1. Summary of reproducibility rates and effect sizes for original and replication studies overall and by journal/discipline.** *df/N* refers to the information on which the test of the effect was based (for example, *df* of *t* test, denominator *df* of *F* test, sample size –3 of correlation, and sample size for *z* and  $\chi^2$ ). Four original results had *P* values slightly higher than 0.05 but were considered positive results in the original article and are treated that way here. Exclusions (explanation provided in supplementary materials, A3) are “replications  $P < 0.05$ ” (3 original nulls excluded; *n* = 97 studies); “mean original and replication effect sizes” (3 excluded; *n* = 97 studies); “meta-analytic mean estimates” (27 excluded; *n* = 73 studies); “percent meta-analytic ( $P < 0.05$ )” (25 excluded; *n* = 75 studies); and, “percent original effect size within replication 95% CI” (5 excluded, *n* = 95 studies).

|                            | Effect size comparison                                      |         |                                            |                                   | Original and replication combined       |                                      |                                 |                                            |                                                    |                                                                       |                                                             |
|----------------------------|-------------------------------------------------------------|---------|--------------------------------------------|-----------------------------------|-----------------------------------------|--------------------------------------|---------------------------------|--------------------------------------------|----------------------------------------------------|-----------------------------------------------------------------------|-------------------------------------------------------------|
|                            | Replications<br><i>P</i> < 0.05<br>in original<br>direction | Percent | Mean<br>(SD)<br>original<br>effect<br>size | Median<br>original<br><i>df/N</i> | Mean (SD)<br>replication<br>effect size | Median<br>replication<br><i>df/N</i> | Average<br>replication<br>power | Meta-<br>analytic<br>mean (SD)<br>estimate | Percent<br>meta-<br>analytic<br>( <i>P</i> < 0.05) | Percent<br>original<br>effect size<br>within<br>replication<br>95% CI | Percent<br>subjective<br>“yes” to<br>“Did it<br>replicate?” |
| Overall                    | 35/97                                                       | 36      | 0.403 (0.188)                              | 54                                | 0.197 (0.257)                           | 68                                   | 0.92                            | 0.309 (0.223)                              | 68                                                 | 47                                                                    | 39                                                          |
| <i>JPSP</i> , social       | 7/31                                                        | 23      | 0.29 (0.10)                                | 73                                | 0.07 (0.11)                             | 120                                  | 0.91                            | 0.138 (0.087)                              | 43                                                 | 34                                                                    | 25                                                          |
| <i>JEP:LMC</i> , cognitive | 13/27                                                       | 48      | 0.47 (0.18)                                | 36.5                              | 0.27 (0.24)                             | 43                                   | 0.93                            | 0.393 (0.209)                              | 86                                                 | 62                                                                    | 54                                                          |
| <i>PSCI</i> , social       | 7/24                                                        | 29      | 0.39 (0.20)                                | 76                                | 0.21 (0.30)                             | 122                                  | 0.92                            | 0.286 (0.228)                              | 58                                                 | 40                                                                    | 32                                                          |
| <i>PSCI</i> , cognitive    | 8/15                                                        | 53      | 0.53 (0.2)                                 | 23                                | 0.29 (0.35)                             | 21                                   | 0.94                            | 0.464 (0.221)                              | 92                                                 | 60                                                                    | 53                                                          |

**Table 2. Spearman’s rank-order correlations of reproducibility indicators with summary original and replication study characteristics.** Effect size difference computed after converting *r* to Fisher’s *z*. *df/N* refers to the information on which the test of the effect was based (for example, *df* of *t* test, denominator *df* of *F* test, sample size –3 of correlation, and sample size for *z* and  $\chi^2$ ). Four original results had *P* values slightly higher than 0.05 but were considered positive results in the original article and are treated that way here. Exclusions (explanation provided in supplementary materials, A3) are “replications  $P < .05$ ” (3 original nulls excluded; *n* = 97 studies), “effect size difference” (3 excluded; *n* = 97 studies); “meta-analytic mean estimates” (27 excluded; *n* = 73 studies); and, “percent original effect size within replication 95% CI” (5 excluded, *n* = 95 studies).

|                                              | Replications<br><i>P</i> < 0.05 in<br>original direction | Effect size<br>difference | Meta-analytic<br>estimate | Original effect<br>size within<br>replication 95% CI | Subjective “yes”<br>to “Did it replicate?” |
|----------------------------------------------|----------------------------------------------------------|---------------------------|---------------------------|------------------------------------------------------|--------------------------------------------|
| Original study characteristics               |                                                          |                           |                           |                                                      |                                            |
| Original <i>P</i> value                      | –0.327                                                   | –0.057                    | –0.468                    | 0.032                                                | –0.260                                     |
| Original effect size                         | 0.304                                                    | 0.279                     | 0.793                     | 0.121                                                | 0.277                                      |
| Original <i>df/N</i>                         | –0.150                                                   | –0.194                    | –0.502                    | –0.221                                               | –0.185                                     |
| Importance of original result                | –0.105                                                   | 0.038                     | –0.205                    | –0.133                                               | –0.074                                     |
| Surprising original result                   | –0.244                                                   | 0.102                     | –0.181                    | –0.113                                               | –0.241                                     |
| Experience and expertise of original team    | –0.072                                                   | –0.033                    | –0.059                    | –0.103                                               | –0.044                                     |
| Replication characteristics                  |                                                          |                           |                           |                                                      |                                            |
| Replication <i>P</i> value                   | –0.828                                                   | 0.621                     | –0.614                    | –0.562                                               | –0.738                                     |
| Replication effect size                      | 0.731                                                    | –0.586                    | 0.850                     | 0.611                                                | 0.710                                      |
| Replication power                            | 0.368                                                    | –0.053                    | 0.142                     | –0.056                                               | 0.285                                      |
| Replication <i>df/N</i>                      | –0.085                                                   | –0.224                    | –0.692                    | –0.257                                               | –0.164                                     |
| Challenge of conducting replication          | –0.219                                                   | 0.085                     | –0.301                    | –0.109                                               | –0.151                                     |
| Experience and expertise of replication team | –0.096                                                   | 0.133                     | 0.017                     | –0.053                                               | –0.068                                     |
| Self-assessed quality of replication         | –0.069                                                   | 0.017                     | 0.054                     | –0.088                                               | –0.055                                     |

replication studies using the Wilcoxon signed-rank test and the *t* test for dependent samples. For both tests, we only used study-pairs for which both *P* values were available.

### Effect sizes

We transformed effect sizes into correlation coefficients whenever possible. Correlation coefficients have several advantages over other effect size measures, such as Cohen’s *d*. Correlation coefficients are bounded, well known, and therefore more readily interpretable. Most important for our purposes, analysis of correlation coefficients is straightforward because, after ap-

plying the Fisher transformation, their standard error is only a function of sample size. Formulas and code for converting test statistics *z*, *F*, *t*, and  $\chi^2$  into correlation coefficients are provided in the appendices at <http://osf.io/ezum7>. To be able to compare and analyze correlations across study-pairs, the original study’s effect size was coded as positive; the replication study’s effect size was coded as negative if the replication study’s effect was opposite to that of the original study.

We compared effect sizes using four tests. We compared the central tendency of the effect size distributions of original and replication studies using both a paired two-sample *t* test and the

Wilcoxon signed-rank test. Third, we computed the proportion of study-pairs in which the effect of the original study was stronger than in the replication study and tested the hypothesis that this proportion is 0.5. For this test, we included findings for which effect size measures were available but no correlation coefficient could be computed (for example, if a regression coefficient was reported but not its test statistic). Fourth, we calculated “coverage,” or the proportion of study-pairs in which the effect of the original study was in the CI of the effect of the replication study, and compared this with the expected proportion using a goodness-of-fit  $\chi^2$  test. We carried

out this test on the subset of study pairs in which both the correlation coefficient and its standard error could be computed [we refer to this data set as the meta-analytic (MA) subset]. Standard errors could only be computed if test statistics were  $r$ ,  $t$ , or  $F(1, df_2)$ . The expected proportion is the sum over expected probabilities across study-pairs. The test assumes the same population effect size for original and replication study in the same study-pair. For those studies that tested the effect with  $F(df_1 > 1, df_2)$  or  $\chi^2$ , we verified coverage using other statistical procedures (computational details are provided in the supplementary materials).

#### Meta-analysis combining original and replication effects

We conducted fixed-effect meta-analyses using the R package metafor (27) on Fisher-transformed correlations for all study-pairs in subset MA and on study-pairs with the odds ratio as the dependent variable. The number of times the CI of all these meta-analyses contained 0 was calculated. For studies in the MA subset, estimated effect sizes were averaged and analyzed by discipline.

#### Subjective assessment of “Did it replicate?”

In addition to the quantitative assessments of replication and effect estimation, we collected subjective assessments of whether the replication provided evidence of replicating the original result. In some cases, the quantitative data anticipate a straightforward subjective assessment of replication. For more complex designs, such as multivariate interaction effects, the quantitative analysis may not provide a simple interpretation. For subjective assessment, replication teams answered “yes” or “no” to the question, “Did your results replicate the original effect?” Additional subjective variables are available for analysis in the full data set.

#### Analysis of moderators

We correlated the five indicators evaluating reproducibility with six indicators of the origi-

nal study (original  $P$  value, original effect size, original sample size, importance of the effect, surprising effect, and experience and expertise of original team) and seven indicators of the replication study (replication  $P$  value, replication effect size, replication power based on original effect size, replication sample size, challenge of conducting replication, experience and expertise of replication team, and self-assessed quality of replication) (Table 2). As follow-up, we did the same with the individual indicators comprising the moderator variables (tables S3 and S4).

### Results

#### Evaluating replication effect against null hypothesis of no effect

A straightforward method for evaluating replication is to test whether the replication shows a statistically significant effect ( $P < 0.05$ ) with the same direction as the original study. This dichotomous vote-counting method is intuitively appealing and consistent with common heuristics used to decide whether original studies “worked.” Ninety-seven of 100 (97%) effects from original studies were positive results (four had  $P$  values falling a bit short of the 0.05 criterion— $P = 0.0508, 0.0514, 0.0516$ , and  $0.0567$ —but all of these were interpreted as positive effects). On the basis of only the average replication power of the 97 original, significant effects [ $M = 0.92$ , median ( $Mdn$ ) =  $0.95$ ], we would expect approximately 89 positive results in the replications if all original effects were true and accurately estimated; however, there were just 35 [36.1%; 95% CI = (26.6%, 46.2%)], a significant reduction [McNemar test,  $\chi^2(1) = 59.1, P < 0.001$ ].

A key weakness of this method is that it treats the 0.05 threshold as a bright-line criterion between replication success and failure (28). It could be that many of the replications fell just short of the 0.05 criterion. The density plots of  $P$  values for original studies (mean  $P$  value = 0.028) and replications (mean  $P$  value = 0.302) are shown in Fig. 1, left. The 64 nonsignificant

$P$  values for replications were distributed widely. When there is no effect to detect, the null distribution of  $P$  values is uniform. This distribution deviated slightly from uniform with positive skew, however, suggesting that at least one replication could be a false negative,  $\chi^2(128) = 155.83, P = 0.048$ . Nonetheless, the wide distribution of  $P$  values suggests against insufficient power as the only explanation for failures to replicate. A scatterplot of original compared with replication study  $P$  values is shown in Fig. 2.

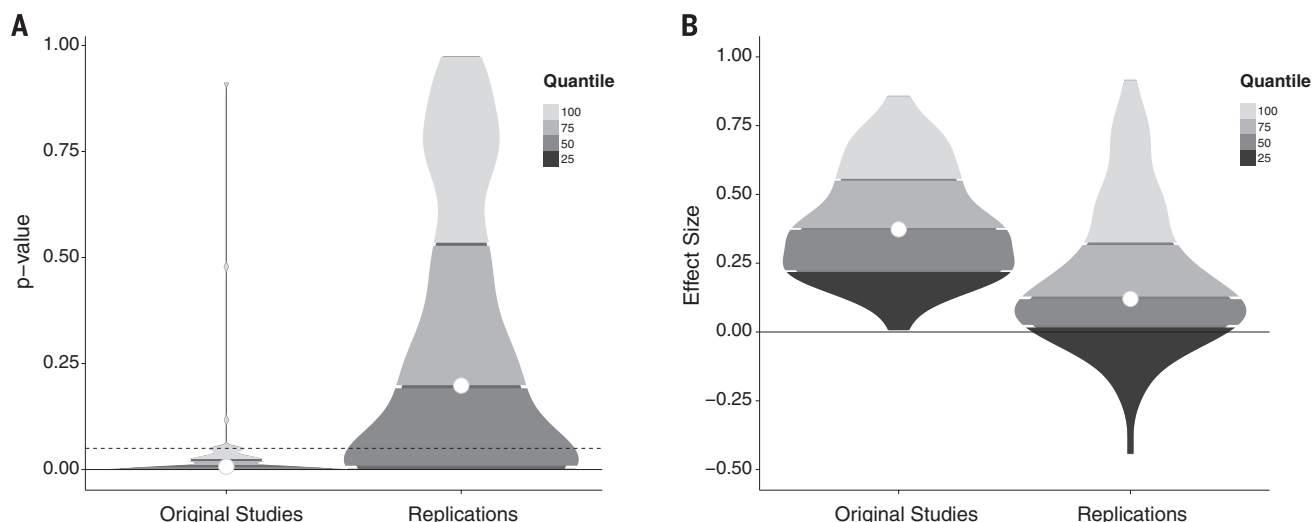
#### Evaluating replication effect against original effect size

A complementary method for evaluating replication is to test whether the original effect size is within the 95% CI of the effect size estimate from the replication. For the subset of 73 studies in which the standard error of the correlation could be computed, 30 (41.1%) of the replication CIs contained the original effect size (significantly lower than the expected value of 78.5%,  $P < 0.001$ ) (supplementary materials). For 22 studies using other test statistics [ $F(df_1 > 1, df_2)$  and  $\chi^2$ ], 68.2% of CIs contained the effect size of the original study. Overall, this analysis suggests a 47.4% replication success rate.

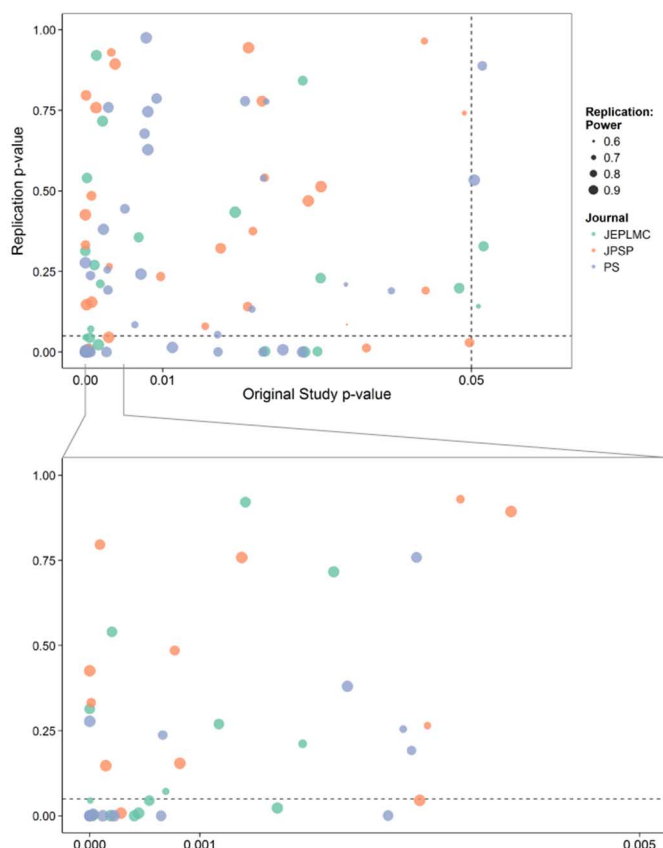
This method addresses the weakness of the first test that a replication in the same direction and a  $P$  value of 0.06 may not be significantly different from the original result. However, the method will also indicate that a replication “fails” when the direction of the effect is the same but the replication effect size is significantly smaller than the original effect size (29). Also, the replication “succeeds” when the result is near zero but not estimated with sufficiently high precision to be distinguished from the original effect size.

#### Comparing original and replication effect sizes

Comparing the magnitude of the original and replication effect sizes avoids special emphasis on  $P$  values. Overall, original study effect sizes



**Fig. 1. Density plots of original and replication  $P$  values and effect sizes.** (A)  $P$  values. (B) Effect sizes (correlation coefficients). Lowest quantiles of  $P$  values are not visible because they are clustered near zero.



**Fig. 2. Scatterplots of original study and replication  $P$  values for three psychology journals.**

Data points scaled by power of the replication based on original study effect size. Dotted red lines indicate  $P = 0.05$  criterion. Subplot below shows  $P$  values from the range between the gray lines ( $P = 0$  to  $0.005$ ) in the main plot above.

( $M = 0.403$ ,  $SD = 0.188$ ) were reliably larger than replication effect sizes ( $M = 0.197$ ,  $SD = 0.257$ ), Wilcoxon's  $W = 7137$ ,  $P < 0.001$ . Of the 99 studies for which an effect size in both the original and replication study could be calculated (30), 82 showed a stronger effect size in the original study (82.8%;  $P < 0.001$ , binomial test) (Fig. 1, right). Original and replication effect sizes were positively correlated (Spearman's  $r = 0.51$ ,  $P < 0.001$ ). A scatterplot of the original and replication effect sizes is presented in Fig. 3.

### Combining original and replication effect sizes for cumulative evidence

The disadvantage of the descriptive comparison of effect sizes is that it does not provide information about the precision of either estimate or resolution of the cumulative evidence for the effect. This is often addressed by computing a meta-analytic estimate of the effect sizes by combining the original and replication studies (28). This approach weights each study by the inverse of its variance and uses these weighted estimates of effect size to estimate cumulative evidence and precision of the effect. Using a fixed-effect model, 51 of the 75 (68%) effects for which a meta-analytic estimate could be computed had 95% CIs that did not include 0.

One qualification about this result is the possibility that the original studies have inflated effect sizes due to publication, selection, reporting, or other biases (9, 12–23). In a discipline with low-powered research designs and an emphasis on positive results for publication, effect sizes will be systematically overestimated in the published literature. There is no publication bias in the replication studies because all results are reported. Also, there are no selection or reporting biases because all were confirmatory tests based on pre-analysis plans. This maximizes the interpretability of the replication  $P$  values and effect estimates. If publication, selection, and reporting biases completely explain the effect differences, then the replication estimates would be a better estimate of the effect size than would the meta-analytic and original results. However, to the extent that there are other influences, such as moderation by sample, setting, or quality of replication, the relative bias influencing original and replication effect size estimation is unknown.

### Subjective assessment of “Did it replicate?”

In addition to the quantitative assessments of replication and effect estimation, replication teams

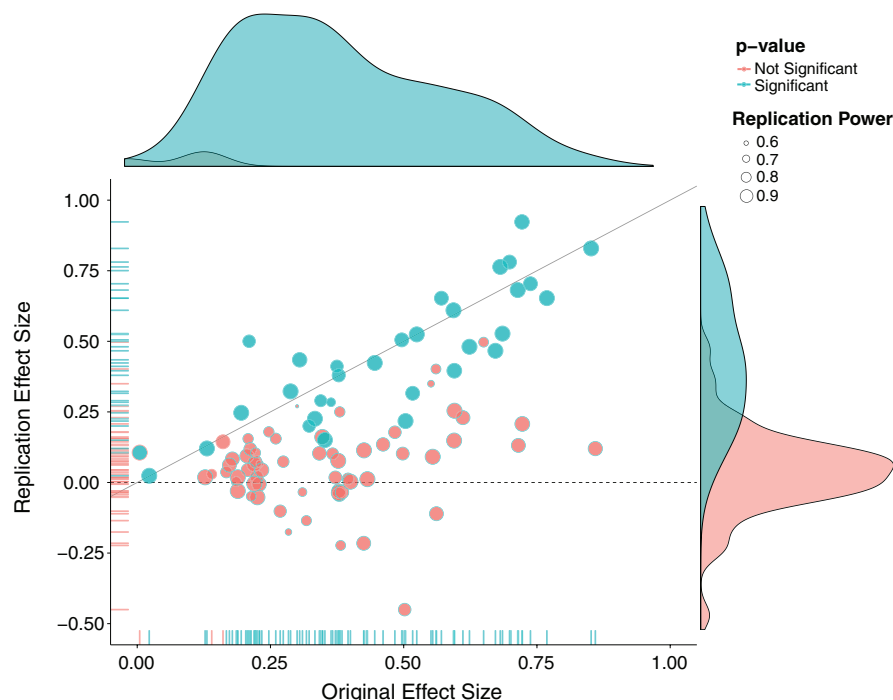
provided a subjective assessment of replication success of the study they conducted. Subjective assessments of replication success were very similar to significance testing results (39 of 100 successful replications), including evaluating “success” for two null replications when the original study reported a null result and “failure” for a  $P < 0.05$  replication when the original result was a null.

### Correlates of reproducibility

The overall replication evidence is summarized in Table 1 across the criteria described above and then separately by journal/discipline. Considering significance testing, reproducibility was stronger in studies and journals representing cognitive psychology than social psychology topics. For example, combining across journals, 14 of 55 (25%) of social psychology effects replicated by the  $P < 0.05$  criterion, whereas 21 of 42 (50%) of cognitive psychology effects did so. Simultaneously, all journals and disciplines showed substantial and similar [ $\chi^2(3) = 2.45$ ,  $P = 0.48$ ] declines in effect size in the replications compared with the original studies. The difference in significance testing results between fields appears to be partly a function of weaker original effects in social psychology studies, particularly in *JPSP*, and perhaps of the greater frequency of high-powered within-subjects manipulations and repeated measurement designs in cognitive psychology as suggested by high power despite relatively small participant samples. Further, the type of test was associated with replication success. Among original, significant effects, 23 of the 49 (47%) that tested main or simple effects replicated at  $P < 0.05$ , but just 8 of the 37 (22%) that tested interaction effects did.

Correlations between reproducibility indicators and characteristics of replication and original studies are provided in Table 2. A negative correlation of replication success with the original study  $P$  value indicates that the initial strength of evidence is predictive of reproducibility. For example, 26 of 63 (41%) original studies with  $P < 0.02$  achieved  $P < 0.05$  in the replication, whereas 6 of 23 (26%) that had a  $P$  value between  $0.02 < P < 0.04$  and 2 of 11 (18%) that had a  $P$  value  $> 0.04$  did so (Fig. 2). Almost two thirds (20 of 32, 63%) of original studies with  $P < 0.001$  had a significant  $P$  value in the replication.

Larger original effect sizes were associated with greater likelihood of achieving  $P < 0.05$  ( $r = 0.304$ ) and a greater effect size difference between original and replication ( $r = 0.279$ ). Moreover, replication power was related to replication success via significance testing ( $r = 0.368$ ) but not with the effect size difference between original and replication ( $r = -0.053$ ). Comparing effect sizes across indicators, surprisingness of the original effect, and the challenge of conducting the replication were related to replication success for some indicators. Surprising effects were less reproducible, as were effects for which it was more challenging to conduct the replication. Last, there was little evidence that perceived importance of the effect, expertise of the original or replication teams, or self-assessed quality of the replication accounted for meaningful variation



**Fig. 3. Original study effect size versus replication effect size (correlation coefficients).** Diagonal line represents replication effect size equal to original effect size. Dotted line represents replication effect size of 0. Points below the dotted line were effects in the opposite direction of the original. Density plots are separated by significant (blue) and nonsignificant (red) effects.

in reproducibility across indicators. Replication success was more consistently related to the original strength of evidence (such as original *P* value, effect size, and effect tested) than to characteristics of the teams and implementation of the replication (such as expertise, quality, or challenge of conducting study) (tables S3 and S4).

## Discussion

No single indicator sufficiently describes replication success, and the five indicators examined here are not the only ways to evaluate reproducibility. Nonetheless, collectively, these results offer a clear conclusion: A large portion of replications produced weaker evidence for the original findings (31) despite using materials provided by the original authors, review in advance for methodological fidelity, and high statistical power to detect the original effect sizes. Moreover, correlational evidence is consistent with the conclusion that variation in the strength of initial evidence (such as original *P* value) was more predictive of replication success than was variation in the characteristics of the teams conducting the research (such as experience and expertise). The latter factors certainly can influence replication success, but the evidence is that they did not systematically do so here. Other investigators may develop alternative indicators to explore further the role of expertise and quality in reproducibility on this open data set.

## Insights on reproducibility

It is too easy to conclude that successful replication means that the theoretical understanding of

the original finding is correct. Direct replication mainly provides evidence for the reliability of a result. If there are alternative explanations for the original finding, those alternatives could likewise account for the replication. Understanding is achieved through multiple, diverse investigations that provide converging support for a theoretical interpretation and rule out alternative explanations.

It is also too easy to conclude that a failure to replicate a result means that the original evidence was a false positive. Replications can fail if the replication methodology differs from the original in ways that interfere with observing the effect. We conducted replications designed to minimize a priori reasons to expect a different result by using original materials, engaging original authors for review of the designs, and conducting internal reviews. Nonetheless, unanticipated factors in the sample, setting, or procedure could still have altered the observed effect magnitudes (32).

More generally, there are indications of cultural practices in scientific communication that may be responsible for the observed results. Low-power research designs combined with publication bias favoring positive results together produce a literature with upwardly biased effect sizes (14, 16, 33, 34). This anticipates that replication effect sizes would be smaller than original studies on a routine basis—not because of differences in implementation but because the original study effect sizes are affected by publication and reporting bias, and the replications are not. Consistent with this expectation, most replication effects were

smaller than original results, and reproducibility success was correlated with indicators of the strength of initial evidence, such as lower original *P* values and larger effect sizes. This suggests publication, selection, and reporting biases as plausible explanations for the difference between original and replication effects. The replication studies significantly reduced these biases because replication preregistration and pre-analysis plans ensured confirmatory tests and reporting of all results.

The observed variation in replication and original results may reduce certainty about the statistical inferences from the original studies but also provides an opportunity for theoretical innovation to explain differing outcomes, and then new research to test those hypothesized explanations. The correlational evidence, for example, suggests that procedures that are more challenging to execute may result in less reproducible results, and that more surprising original effects may be less reproducible than less surprising original effects. Further, systematic, repeated replication efforts that fail to identify conditions under which the original finding can be observed reliably may reduce confidence in the original finding.

## Implications and limitations

The present study provides the first open, systematic evidence of reproducibility from a sample of studies in psychology. We sought to maximize generalizability of the results with a structured process for selecting studies for replication. However, it is unknown the extent to which these findings extend to the rest of psychology or other disciplines. In the sampling frame itself, not all articles were replicated; in each article, only one study was replicated; and in each study, only one statistical result was subject to replication. More resource-intensive studies were less likely to be included than were less resource-intensive studies. Although study selection bias was reduced by the sampling frame and selection strategy, the impact of selection bias is unknown.

We investigated the reproducibility rate of psychology not because there is something special about psychology, but because it is our discipline. Concerns about reproducibility are widespread across disciplines (9–21). Reproducibility is not well understood because the incentives for individual scientists prioritize novelty over replication (20). If nothing else, this project demonstrates that it is possible to conduct a large-scale examination of reproducibility despite the incentive barriers. Here, we conducted single-replication attempts of many effects obtaining broad-and-shallow evidence. These data provide information about reproducibility in general but little precision about individual effects in particular. A complementary narrow-and-deep approach is characterized by the Many Labs replication projects (32). In those, many replications of single effects allow precise estimates of effect size but result in generalizability that is circumscribed to those individual effects. Pursuing both strategies across disciplines, such as the ongoing effort in cancer biology (35), would yield insight about common and distinct

challenges and may cross-fertilize strategies so as to improve reproducibility.

Because reproducibility is a hallmark of credible scientific evidence, it is tempting to think that maximum reproducibility of original results is important from the onset of a line of inquiry through its maturation. This is a mistake. If initial ideas were always correct, then there would hardly be a reason to conduct research in the first place. A healthy discipline will have many false starts as it confronts the limits of present understanding.

Innovation is the engine of discovery and is vital for a productive, effective scientific enterprise. However, innovative ideas become old news fast. Journal reviewers and editors may dismiss a new test of a published idea as unoriginal. The claim that “we already know this” belies the uncertainty of scientific evidence. Deciding the ideal balance of resourcing innovation versus verification is a question of research efficiency. How can we maximize the rate of research progress? Innovation points out paths that are possible; replication points out paths that are likely; progress relies on both. The ideal balance is a topic for investigation itself. Scientific incentives—funding, publication, or awards—can be tuned to encourage an optimal balance in the collective effort of discovery (36, 37).

Progress occurs when existing expectations are violated and a surprising result spurs a new investigation. Replication can increase certainty when findings are reproduced and promote innovation when they are not. This project provides accumulating evidence for many findings in psychological research and suggests that there is still more work to do to verify whether we know what we think we know.

## Conclusion

After this intensive effort to reproduce a sample of published psychological findings, how many of the effects have we established are true? Zero. And how many of the effects have we established are false? Zero. Is this a limitation of the project design? No. It is the reality of doing science, even if it is not appreciated in daily practice. Humans desire certainty, and science infrequently provides it. As much as we might wish it to be otherwise, a single study almost never provides definitive resolution for or against an effect and its explanation. The original studies examined here offered tentative evidence; the replications we conducted offered additional, confirmatory evidence. In some cases, the replications increase confidence in the reliability of the original results; in other cases, the replications suggest that more investigation is needed to establish the validity of the original findings. Scientific progress is a cumulative process of uncertainty reduction that can only succeed if science itself remains the greatest skeptic of its explanatory claims.

The present results suggest that there is room to improve reproducibility in psychology. Any temptation to interpret these results as a defeat for psychology, or science more generally, must contend with the fact that this project demon-

strates science behaving as it should. Hypotheses abound that the present culture in science may be negatively affecting the reproducibility of findings. An ideological response would discount the arguments, discredit the sources, and proceed merrily along. The scientific process is not ideological. Science does not always provide comfort for what we wish to be; it confronts us with what is. Moreover, as illustrated by the Transparency and Openness Promotion (TOP) Guidelines (<http://cos.io/top>) (37), the research community is taking action already to improve the quality and credibility of the scientific literature.

We conducted this project because we care deeply about the health of our discipline and believe in its promise for accumulating knowledge about human behavior that can advance the quality of the human condition. Reproducibility is central to that aim. Accumulating evidence is the scientific community's method of self-correction and is the best available option for achieving that ultimate goal: truth.

## REFERENCES AND NOTES

- C. Hempel, Maximal specificity and lawlikeness in probabilistic explanation. *Philos. Sci.* **35**, 116–133 (1968). doi: [10.1086/288197](https://doi.org/10.1086/288197)
- C. Hempel, P. Oppenheim, Studies in the logic of explanation. *Philos. Sci.* **15**, 135–175 (1948). doi: [10.1086/286983](https://doi.org/10.1086/286983)
- I. Lakatos, in *Criticism and the Growth of Knowledge*, I. Lakatos, A. Musgrave, Eds. (Cambridge Univ. Press, London, 1970), pp. 170–196.
- P. E. Meehl, Appraising and amending theories: The strategy of Lakatosian defense and two principles that warrant it. *Psychol. Inq.* **1**, 108–141 (1990). doi: [10.1207/s15327965pli0102\\_1](https://doi.org/10.1207/s15327965pli0102_1)
- J. R. Platt, Strong Inference: Certain systematic methods of scientific thinking may produce much more rapid progress than others. *Science* **146**, 347–353 (1964). doi: [10.1126/science.146.3642.347](https://doi.org/10.1126/science.146.3642.347); pmid: [17739513](https://pubmed.ncbi.nlm.nih.gov/17739513/)
- W. C. Salmon, in *Introduction to the Philosophy of Science*, M. H. Salmon Ed. (Hackett Publishing Company, Indianapolis, 1999), pp. 7–41.
- B. A. Nosek, D. Lakens, Registered reports: A method to increase the credibility of published results. *Soc. Psychol.* **45**, 137–141 (2014). doi: [10.1027/1864-9335/a000192](https://doi.org/10.1027/1864-9335/a000192)
- S. Schmidt, Shall we really do it again? The powerful concept of replication is neglected in the social sciences. *Rev. Gen. Psychol.* **13**, 90–100 (2009). doi: [10.1037/a0015108](https://doi.org/10.1037/a0015108)
- J. P. A. Ioannidis, Why most published research findings are false. *PLoS Med.* **2**, e124 (2005). doi: [10.1371/journal.pmed.0020124](https://doi.org/10.1371/journal.pmed.0020124); pmid: [16060722](https://pubmed.ncbi.nlm.nih.gov/16060722/)
- C. G. Begley, L. M. Ellis, Drug development: Raise standards for preclinical cancer research. *Nature* **483**, 531–533 (2012). doi: [10.1038/483531a](https://doi.org/10.1038/483531a); pmid: [22460880](https://pubmed.ncbi.nlm.nih.gov/22460880/)
- F. Prinz, T. Schlange, K. Asadullah, Believe it or not: How much can we rely on published data on potential drug targets? *Nat. Rev. Drug Discov.* **10**, 712–713 (2011). doi: [10.1038/nrd3439-c1](https://doi.org/10.1038/nrd3439-c1); pmid: [21892149](https://pubmed.ncbi.nlm.nih.gov/21892149/)
- M. McNutt, Reproducibility. *Science* **343**, 229 (2014). doi: [10.1126/science.1250475](https://doi.org/10.1126/science.1250475); pmid: [24436391](https://pubmed.ncbi.nlm.nih.gov/24436391/)
- H. Pashler, E.-J. Wagenmakers, Editors' introduction to the special section on replicability in psychological science: A crisis of confidence? *Perspect. Psychol. Sci.* **7**, 528–530 (2012). doi: [10.1177/1745691612465253](https://doi.org/10.1177/1745691612465253); pmid: [26168108](https://pubmed.ncbi.nlm.nih.gov/26168108/)
- K. S. Button *et al.*, Power failure: Why small sample size undermines the reliability of neuroscience. *Nat. Rev. Neurosci.* **14**, 365–376 (2013). doi: [10.1038/nrn3475](https://doi.org/10.1038/nrn3475); pmid: [23571845](https://pubmed.ncbi.nlm.nih.gov/23571845/)
- D. Fanelli, “Positive” results increase down the hierarchy of the sciences. *PLOS ONE* **5**, e10068 (2010). doi: [10.1371/journal.pone.0010068](https://doi.org/10.1371/journal.pone.0010068); pmid: [20383332](https://pubmed.ncbi.nlm.nih.gov/20383332/)
- A. G. Greenwald, Consequences of prejudice against the null hypothesis. *Psychol. Bull.* **82**, 1–20 (1975). doi: [10.1037/h0076157](https://doi.org/10.1037/h0076157)
- G. S. Howard *et al.*, Do research literatures give correct answers? *Rev. Gen. Psychol.* **13**, 116–121 (2009). doi: [10.1037/a0015468](https://doi.org/10.1037/a0015468)
- J. P. A. Ioannidis, M. R. Munafò, P. Fusar-Poli, B. A. Nosek, S. P. David, Publication and other reporting biases in cognitive sciences: Detection, prevalence, and prevention. *Trends Cogn. Sci.* **18**, 235–241 (2014). doi: [10.1016/j.tics.2014.02.010](https://doi.org/10.1016/j.tics.2014.02.010); pmid: [24656991](https://pubmed.ncbi.nlm.nih.gov/24656991/)
- L. K. John, G. Loewenstein, D. Prelec, Measuring the prevalence of questionable research practices with incentives for truth telling. *Psychol. Sci.* **23**, 524–532 (2012). doi: [10.1177/0956797611430953](https://doi.org/10.1177/0956797611430953); pmid: [22508865](https://pubmed.ncbi.nlm.nih.gov/22508865/)
- B. A. Nosek, J. R. Spies, M. Motyl, Scientific utopia: II. Restructuring incentives and practices to promote truth over publishability. *Perspect. Psychol. Sci.* **7**, 615–631 (2012). doi: [10.1177/1745691612459058](https://doi.org/10.1177/1745691612459058); pmid: [26168121](https://pubmed.ncbi.nlm.nih.gov/26168121/)
- R. Rosenthal, The file drawer problem and tolerance for null results. *Psychol. Bull.* **86**, 638–641 (1979). doi: [10.1037/0033-2909.86.3.638](https://doi.org/10.1037/0033-2909.86.3.638)
- P. Rozin, What kind of empirical research should we publish, fund, and reward?: A different perspective. *Perspect. Psychol. Sci.* **4**, 435–439 (2009). doi: [10.1111/j.1745-6924.2009.01151.x](https://doi.org/10.1111/j.1745-6924.2009.01151.x); pmid: [26158991](https://pubmed.ncbi.nlm.nih.gov/26158991/)
- J. P. Simmons, L. D. Nelson, U. Simonsohn, False-positive psychology: Undisclosed flexibility in data collection and analysis allows presenting anything as significant. *Psychol. Sci.* **22**, 1359–1366 (2011). doi: [10.1177/0956797611417632](https://doi.org/10.1177/0956797611417632); pmid: [22006061](https://pubmed.ncbi.nlm.nih.gov/22006061/)
- Open Science Collaboration, An open, large-scale, collaborative effort to estimate the reproducibility of psychological science. *Perspect. Psychol. Sci.* **7**, 657–660 (2012). doi: [10.1177/1745691612462588](https://doi.org/10.1177/1745691612462588); pmid: [26168127](https://pubmed.ncbi.nlm.nih.gov/26168127/)
- Open Science Collaboration, in *Implementing Reproducible Computational Research (A Volume in The R Series)*, V. Stodden, F. Leisch, R. Peng, Eds. (Taylor & Francis, New York, 2014), pp. 299–323.
- R. A. Fisher, Theory of statistical estimation. *Math. Proc. Camb. Philos. Soc.* **22**, 700–725 (1925). doi: [10.1017/S0305004100009580](https://doi.org/10.1017/S0305004100009580)
- W. Viechtbauer, Conducting meta-analyses in R with the metafor package. *J. Stat. Softw.* **36**, 1–48 (2010).
- S. L. Braver, F. J. Thoemmes, R. Rosenthal, Continuously cumulating meta-analysis and replicability. *Perspect. Psychol. Sci.* **9**, 333–342 (2014). doi: [10.1177/1745691614529796](https://doi.org/10.1177/1745691614529796); pmid: [26173268](https://pubmed.ncbi.nlm.nih.gov/26173268/)
- U. Simonsohn, Small telescopes: Detectability and the evaluation of replication results. *Psychol. Sci.* **26**, 559–569 (2015). doi: [10.1177/0956797614567341](https://doi.org/10.1177/0956797614567341); pmid: [25800521](https://pubmed.ncbi.nlm.nih.gov/25800521/)
- D. Lakens, Calculating and reporting effect sizes to facilitate cumulative science: A practical primer for t-tests and ANOVAs. *Front. Psychol.* **4**, 863 (2013). doi: [10.3389/fpsyg.2013.00863](https://doi.org/10.3389/fpsyg.2013.00863); pmid: [24324449](https://pubmed.ncbi.nlm.nih.gov/24324449/)
- J. Lehrer, The truth wears off: Is there something wrong with the scientific method? *The New Yorker*, 52–57 (2010).
- R. Klein *et al.*, Investigating variation in replicability: A “many labs” replication project. *Soc. Psychol.* **45**, 142–152 (2014). doi: [10.1027/1864-9335/a000178](https://doi.org/10.1027/1864-9335/a000178)
- J. Cohen, The statistical power of abnormal-social psychological research: A review. *J. Abnorm. Soc. Psychol.* **65**, 145–153 (1962). doi: [10.1037/h0045186](https://doi.org/10.1037/h0045186); pmid: [13880271](https://pubmed.ncbi.nlm.nih.gov/13880271/)
- T. D. Sterling, Publication decisions and their possible effects on inferences drawn from tests of significance—or vice versa. *J. Am. Stat. Assoc.* **54**, 30–34 (1959).
- T. M. Errington *et al.*, An open investigation of the reproducibility of cancer biology research. *eLife* **3**, e04333 (2014). doi: [10.7554/eLife.04333](https://doi.org/10.7554/eLife.04333); pmid: [25490932](https://pubmed.ncbi.nlm.nih.gov/25490932/)
- J. K. Hartshorne, A. Schachner, Tracking replicability as a method of post-publication open evaluation. *Front. Comput. Neurosci.* **6**, 8 (2012). doi: [10.3389/fncom.2012.00008](https://doi.org/10.3389/fncom.2012.00008); pmid: [22403538](https://pubmed.ncbi.nlm.nih.gov/22403538/)
- B. A. Nosek *et al.*, Promoting an open research culture. *Science* **348**, 1422–1425 (2015). doi: [10.1126/science.aab2374](https://doi.org/10.1126/science.aab2374); pmid: [26113702](https://pubmed.ncbi.nlm.nih.gov/26113702/)

## ACKNOWLEDGMENTS

In addition to the coauthors of this manuscript, there were many volunteers who contributed to project success. We thank D. Acup, J. Anderson, S. Anzellotti, R. Araujo, J. D. Arnal, T. Bates, R. Battleday, R. Bauchwitz, M. Bernstein, B. Blohowiak, M. Boffo, E. Bruneau, B. Chabot-Hanowell, J. Chan, P. Chu, A. Dalla Rosa, B. Deen, P. DiGiacomo, C. Dogulu, N. Dufour, C. Fitzgerald, A. Foote, A. Garcia, E. Garcia, C. Gautreau, L. Germaine, T. Gill, L. Goldberg, S. D. Goldinger, H. Gweon, D. Haile, K. Hart, F. Hjorth, J. Hoening, A. Innes-Ker, B. Jansen, R. Jersakova, Y. Jie, Z. Kaldy, W. K. Kong, A. Kenney, J. Kingston, J. Koster-Hale, A. Lam, R. LeDonne, D. Lumian, E. Luong, S. Man-pui, J. Martin, A. Mauk, T. McElroy, K. McRae, T. Miller, K. Moser, M. Mullarkey, A. R. Munoz, J. Ong, C. Parks, D. S. Pate, D. Patron, H. J. M. Pennings, M. Penuliar, A. Pfammatter, J. P. Shanoltz, E. Stevenson, E. Pichler, H. Raudszus, H. Richardson, N. Rothstein, T. Scherndl, S. Schrager, S. Shah, Y. S. Tai, A. Skerry, M. Steinberg, J. Stoetrau, H. Tibboel, A. Tooley, A. Tullett, C. Vaccaro, E. Vergauwe, A. Watanabe, I. Weiss, M. H. White II,

P. Whitehead, C. Widmann, D. K. Williams, K. M. Williams, and H. Yi. Also, we thank the authors of the original research that was the subject of replication in this project. These authors were generous with their time, materials, and advice for improving the quality of each replication and identifying the strengths and limits of the outcomes. The authors of this work are listed alphabetically. This project was supported by the Center for Open Science and the Laura and John Arnold Foundation. The authors declare no financial conflict of interest with the reported research.

# The Open Science Collaboration

Alexander A. Aarts,<sup>1</sup> Joanna E. Anderson,<sup>2</sup> Christopher J. Anderson,<sup>3</sup> Peter R. Attridge,<sup>4,5</sup> Angela Attwood,<sup>6</sup> Jordan Axt,<sup>7</sup> Molly Babel,<sup>8</sup> Štěpán Bahnik,<sup>9</sup> Erica Baranski,<sup>10</sup> Michael Barnett-Cowan,<sup>11</sup> Elizabeth Bartmess,<sup>12</sup> Jennifer Beer,<sup>13</sup> Raoul Bell,<sup>14</sup> Heather Bentley,<sup>5</sup> Leah Beyan,<sup>5</sup> Grace Binion,<sup>5,15</sup> Denny Borsboom,<sup>16</sup> Annick Bosch,<sup>17</sup> Frank A. Bosco,<sup>18</sup> Sara D. Bowman,<sup>19</sup> Mark J. Brandt,<sup>20</sup> Erin Braswell,<sup>19</sup> Hilmar Brohmer,<sup>20</sup> Benjamin T. Brown,<sup>5</sup> Kristina Brown,<sup>5</sup> Jovita Bruning,<sup>21,22</sup> Ann Calhoun-Sauls,<sup>23</sup> Shannon P. Callahan,<sup>24</sup> Elizabeth Chagnon,<sup>25</sup> Jesse Chandler,<sup>26,27</sup> Christopher R. Chartier,<sup>28</sup> Felix Cheung,<sup>29,30</sup> Cody D. Christopherson,<sup>31</sup> Linda Cillessen,<sup>17</sup> Russ Clay,<sup>32</sup> Hayley Cleary,<sup>18</sup> Mark D. Cloud,<sup>33</sup> Michael Cohn,<sup>12</sup> Johanna Cohoon,<sup>19</sup> Simon Columbus,<sup>15</sup> Andreas Cordes,<sup>34</sup> Giulio Costantini,<sup>35</sup> Leslie D. Cramblet Alvarez,<sup>36</sup> Ed Cremata,<sup>37</sup> Jan Crusius,<sup>38</sup> Jamie DeCoster,<sup>7</sup> Michelle A. DeGaetano,<sup>5</sup> Nicolás Della Penna,<sup>39</sup> Bobby den Bezemer,<sup>16</sup> Marie K. Deserno,<sup>16</sup> Olivia Devitt,<sup>5</sup> Laura Dewitte,<sup>40</sup> David G. Dobolyi,<sup>7</sup> Geneva T. Dodson,<sup>7</sup> M. Brent Donnellan,<sup>41</sup> Ryan Donohue,<sup>42</sup> Rebecca A. Dore,<sup>7</sup> Angela Dorrrough,<sup>43,44</sup> Anna Dreber,<sup>45</sup> Michelle Dugas,<sup>25</sup> Elizabeth W. Dunn,<sup>8</sup> Kayleigh Easey,<sup>6</sup> Sylvia Eoibge,<sup>5</sup> Casey Eggleston,<sup>7</sup> Jo Embley,<sup>46</sup> Sacha Epskamp,<sup>16</sup> Timothy M. Errington,<sup>19</sup> Vivien Estel,<sup>47</sup> Frank J. Farach,<sup>48,49</sup> Jenelle Feather,<sup>50</sup> Anna Fedor,<sup>51</sup> Belén Fernández-Castilla,<sup>52</sup> Susann Fiedler,<sup>44</sup> James G. Field,<sup>18</sup> Stanka A. Fitneva,<sup>53</sup> Taru Flagan,<sup>13</sup> Amanda L. Forest,<sup>54</sup> Eskil Forsell,<sup>45</sup> Joshua D. Foster,<sup>55</sup> Michael C. Frank,<sup>56</sup> Rebecca S. Frazier,<sup>7</sup> Heather Fuchs,<sup>38</sup> Philip Gable,<sup>57</sup> Jeff Galak,<sup>58</sup> Elisa Maria Galliani,<sup>59</sup> Anup Gampa,<sup>7</sup> Sara Garcia,<sup>60</sup> Douglas Gazarian,<sup>61</sup> Elizabeth Gilbert,<sup>62</sup> Roger Giner-Sorolla,<sup>46</sup> Andreas Glöckner,<sup>34,44</sup> Lars Goellner,<sup>43</sup> Jin X. Goh,<sup>62</sup> Rebecca Goldberg,<sup>63</sup> Patrick T. Goodbourn,<sup>64</sup> Shauna Gordon-McKeon,<sup>65</sup> Bryan Gorges,<sup>19</sup> Jessie Gorges,<sup>19</sup> Justin Goss,<sup>66</sup> Jesse Graham,<sup>37</sup> James A. Grange,<sup>67</sup> Jeremy Gray,<sup>29</sup> Chris Hartgering,<sup>20</sup> Joshua Hartshorne,<sup>50</sup> Fred Hasselman,<sup>17,68</sup> Timothy Hayes,<sup>37</sup> Emma Heikensten,<sup>45</sup> Felix Henninger,<sup>69,44</sup> John Hodsoll,<sup>70,71</sup> Taylor Holubar,<sup>56</sup> Gea Hoogendoorn,<sup>20</sup> Denise J. Humphries,<sup>5</sup> Cathy O.-Y. Hung,<sup>30</sup> Nathali Immelman,<sup>72</sup> Vanessa C. Irsik,<sup>73</sup> Georg Jahn,<sup>74</sup> Frank Jäkel,<sup>75</sup> Marc Jekel,<sup>34</sup> Magnus Johannesson,<sup>45</sup> Larissa G. Johnson,<sup>76</sup> David J. Johnson,<sup>29</sup> Kate M. Johnson,<sup>37</sup> William J. Johnston,<sup>77</sup> Kai Jonas,<sup>66</sup> Jennifer A. Joy-Gaba,<sup>18</sup> Heather Barry Kappes,<sup>78</sup> Kim Kelso,<sup>36</sup> Mallory C. Kidwell,<sup>19</sup> Seung Kyung Kim,<sup>56</sup> Matthew Kirkhart,<sup>79</sup> Bennett Kleinberg,<sup>16,80</sup> Goran Knežević,<sup>81</sup> Franziska Maria Kolorz,<sup>17</sup> Jolanda J. Kossakowski,<sup>16</sup> Robert Wilhelm Krause,<sup>17</sup> Job Krijnen,<sup>20</sup> Tim Kuhlmann,<sup>82</sup> Yoram K. Kunkels,<sup>16</sup> Megan M. Kyc,<sup>33</sup> Calvin K. Lai,<sup>7</sup> Aamir Laique,<sup>83</sup> Daniel Lakens,<sup>84</sup> Kristin A. Lane,<sup>61</sup> Bethany Lassetter,<sup>85</sup> Ljiljana B. Lazarevic,<sup>81</sup> Etienne P. LeBel,<sup>86</sup> Key Jung Lee,<sup>56</sup> Minha Lee,<sup>7</sup> Kristi Lemm,<sup>87</sup> Carmel A. Levitan,<sup>88</sup> Melissa Lewis,<sup>89</sup> Lin Lin,<sup>30</sup> Stephanie Lin,<sup>56</sup> Matthias Lippold,<sup>34</sup> Darren Loureiro,<sup>25</sup> Ilse Luteijn,<sup>17</sup> Sean Mackinnon,<sup>90</sup> Heather N. Mainard,<sup>5</sup> Denise C. Marigold,<sup>91</sup> Daniel P. Martin,<sup>7</sup> Tylar Martinez,<sup>36</sup> E.J. Masicampo,<sup>92</sup> Josh Matacotta,<sup>93</sup> Maya Mathur,<sup>56</sup> Michael May,<sup>44,94</sup> Nicole Mechin,<sup>57</sup> Pranjal Mehta,<sup>15</sup> Johannes Meixner,<sup>21,95</sup> Alissa Melinger,<sup>96</sup> Jeremy K. Miller,<sup>97</sup> Mallorie Miller,<sup>83</sup> Katherine Moore,<sup>42,98</sup> Marcus Möschl,<sup>99</sup> Matt Motyl,<sup>100</sup> Stephanie M. Müller,<sup>47</sup> Marcus Munafo,<sup>6</sup> Koen I. Neijenhuijs,<sup>17</sup> Taylor Nervi,<sup>28</sup> Gandalf Nicolas,<sup>101</sup> Gustav Nilssonne,<sup>102,103</sup> Brian A. Nosek,<sup>73</sup> Michèle B. Nuijten,<sup>20</sup> Catherine Olsson,<sup>50,104</sup> Colleen Osborne,<sup>7</sup> Lutz Ostkamp,<sup>75</sup> Misha Pavel,<sup>62</sup> Ian S. Penton-Voak,<sup>65</sup> Olivia Perna,<sup>28</sup> Cyril Pernet,<sup>105</sup> Marco Perugini,<sup>35</sup> R. Nathan Pipitone,<sup>36</sup> Michael Pitts,<sup>89</sup> Franziska Plessow,<sup>99,106</sup> Jason M. Prenoveau,<sup>79</sup> Rima-Maria Rahal,<sup>44,16</sup> Kate A. Ratliff,<sup>107</sup> David Reinhard,<sup>7</sup> Frank Renkewitz,<sup>47</sup> Ashley A. Ricker,<sup>10</sup> Anastasia Rigney,<sup>13</sup> Andrew M. Rivers,<sup>24</sup> Mark Roebke,<sup>108</sup> Abraham M. Rutnick,<sup>109</sup> Robert S. Ryan,<sup>110</sup> Onur Sahin,<sup>16</sup> Anondah Saide,<sup>10</sup> Gillian M. Sandstrom,<sup>6</sup> David Santos,<sup>111,112</sup> Rebecca Saxe,<sup>50</sup> René Schlegelmilch,<sup>44,47</sup> Kathleen Schmidt,<sup>113</sup> Sabine Scholz,<sup>114</sup> Larissa Seibel,<sup>17</sup> Dylan Faulkner Selterman,<sup>25</sup> Samuel Shaki,<sup>115</sup> William B. Simpson,<sup>7</sup> H. Colleen Sinclair,<sup>63</sup> Jeanine L. M. Skorinko,<sup>116</sup> Agnieszka Slowik,<sup>117</sup> Joel S. Snyder,<sup>73</sup> Courtney Soderberg,<sup>19</sup> Carina Sonnleitner,<sup>117</sup> Nick Spencer,<sup>36</sup> Jeffrey R. Spies,<sup>19</sup>

Sara Steegen,<sup>40</sup> Stefan Stieger,<sup>82</sup> Nina Strohminger,<sup>118</sup> Gavin B. Sullivan,<sup>119</sup> Thomas Talhelm,<sup>119</sup> Megan Tapia,<sup>36</sup> Annieke Te Dorsthorst,<sup>17</sup> Manuela Thomas,<sup>72,120</sup> Sarah L. Thomas,<sup>7</sup> Pia Tio,<sup>16</sup> Frits Traets,<sup>40</sup> Steve Tsang,<sup>121</sup> Francis Tuerlinckx,<sup>40</sup> Paul Turchan,<sup>122</sup> Milan Valášek,<sup>105</sup> Anna E. van 't Veer,<sup>20,123</sup> Robbie Van Aert,<sup>20</sup> Marcel van Assen,<sup>20</sup> Riet van Bork,<sup>16</sup> Mathijs van de Ven,<sup>17</sup> Don van den Bergh,<sup>16</sup> Marije van der Hulst,<sup>17</sup> Roel van Dooren,<sup>17</sup> Johnny van Doorn,<sup>40</sup> Daan R. van Renswoude,<sup>16</sup> Hedderik van Rijn,<sup>114</sup> Wolf Vanpaemel,<sup>40</sup> Alejandro Vázquez Echeverría,<sup>124</sup> Melissa Vazquez,<sup>5</sup> Natalia Velez,<sup>56</sup> Marieke Vermue,<sup>17</sup> Mark Verschoor,<sup>20</sup> Michelangelo Vianello,<sup>59</sup> Martin Voracek,<sup>117</sup> Gina Vuu,<sup>7</sup> Eric-Jan Wagenmakers,<sup>15</sup> Joanneke Weerdmeester,<sup>17</sup> Ashlee Welsh,<sup>36</sup> Erin C. Westgate,<sup>7</sup> Joeri Wissink,<sup>20</sup> Michael Wood,<sup>72</sup> Andy Woods,<sup>125,6</sup> Emily Wright,<sup>36</sup> Sining Wu,<sup>63</sup> Marcel Zeelenberg,<sup>20</sup> Kellylynn Zuni<sup>36</sup>

<sup>1</sup>Open Science Collaboration, Nuenen, Netherlands. <sup>2</sup>Defense Research and Development Canada, Ottawa, ON, Canada. <sup>3</sup>Department of Psychology, Southern New Hampshire University, Schenectady, NY 12305, USA. <sup>4</sup>Mercer School of Medicine, Macon, GA 31207, USA. <sup>5</sup>Georgia Gwinnett College, Lawrenceville, GA 30043, USA. <sup>6</sup>School of Experimental Psychology, University of Bristol, Bristol BS8 1TH, UK. <sup>7</sup>University of Virginia, Charlottesville, VA 22904, USA. <sup>8</sup>University of British Columbia, Vancouver, BC V6T 1Z4 Canada. <sup>9</sup>Department of Psychology II, University of Würzburg, Würzburg, Germany. <sup>10</sup>Department of Psychology, University of California, Riverside, Riverside, CA 92521, USA. <sup>11</sup>University of Waterloo, Waterloo, ON N2L 3G1, Canada. <sup>12</sup>University of California, San Francisco, San Francisco, CA 94118, USA. <sup>13</sup>Department of Psychology, University of Texas at Austin, Austin, TX 78712, USA. <sup>14</sup>Department of Experimental Psychology, Heinrich Heine University Düsseldorf, Düsseldorf, Germany. <sup>15</sup>Department of Psychology, University of Oregon, Eugene, OR 97403, USA. <sup>16</sup>Department of Psychology, University of Amsterdam, Amsterdam, Netherlands. <sup>17</sup>Radboud University Nijmegen, Nijmegen, Netherlands. <sup>18</sup>Virginia Commonwealth University, Richmond, VA 23284, USA. <sup>19</sup>Center for Open Science, Charlottesville, VA 22902, USA. <sup>20</sup>Department of Social Psychology, Tilburg University, Tilburg, Netherlands. <sup>21</sup>Humboldt University of Berlin, Berlin, Germany. <sup>22</sup>Charité—Universitätsmedizin Berlin, Berlin, Germany. <sup>23</sup>Belmont Abbey College, Belmont, NC 28012, USA. <sup>24</sup>Department of Psychology, University of California, Davis, Davis, CA 95616, USA. <sup>25</sup>University of Maryland, Washington, DC 20012, USA. <sup>26</sup>Institute for Social Research, University of Michigan, Ann Arbor, MI 48104, USA. <sup>27</sup>Mathematica Policy Research, Washington, DC 20002, USA. <sup>28</sup>Ashland University, Ashland, OH 44805, USA. <sup>29</sup>Michigan State University, East Lansing, MI 48824, USA. <sup>30</sup>Department of Psychology, University of Hong Kong, Pok Fu Lam, Hong Kong. <sup>31</sup>Department of Psychology, Southern Oregon University, Ashland, OR 97520, USA. <sup>32</sup>College of Staten Island, City University of New York, Staten Island, NY 10314, USA. <sup>33</sup>Department of Psychology, Lock Haven University, Lock Haven, PA 17745, USA. <sup>34</sup>University of Göttingen, Göttingen, Germany. <sup>35</sup>University of Milan-Bicocca, Milan, Italy. <sup>36</sup>Department of Psychology, Adams State University, Alamosa, CO 81101, USA. <sup>37</sup>University of Southern California, Los Angeles, CA 90089, USA. <sup>38</sup>University of Cologne, Cologne, Germany. <sup>39</sup>Australian National University, Canberra, Australia. <sup>40</sup>University of Leuven, Leuven, Belgium. <sup>41</sup>Texas A & M University, College Station, TX 77845, USA. <sup>42</sup>Elmhurst College, Elmhurst, IL 60126, USA. <sup>43</sup>Department of Social Psychology, University of Siegen, Siegen, Germany. <sup>44</sup>Max Planck Institute for Research on Collective Goods, Bonn, Germany. <sup>45</sup>Department of Economics, Stockholm School of Economics, Stockholm, Sweden. <sup>46</sup>School of Psychology, University of Kent, Canterbury, Kent, UK. <sup>47</sup>University of Erfurt, Erfurt, Germany. <sup>48</sup>University of Washington, Seattle, WA 98195, USA. <sup>49</sup>Prometheus Research, New Haven, CT 06510, USA. <sup>50</sup>Department of Brain and Cognitive Sciences, Massachusetts Institute of Technology, Cambridge, MA 02139, USA. <sup>51</sup>Parmenides Center for the Study of Thinking, Munich, Germany. <sup>52</sup>Universidad Complutense de Madrid, Madrid, Spain. <sup>53</sup>Department of Psychology, Queen's University, Kingston, ON, Canada. <sup>54</sup>Department of Psychology, University of Pittsburgh, Pittsburgh, PA 15260, USA. <sup>55</sup>Department of Psychology, University of South Alabama, Mobile, AL 36688, USA. <sup>56</sup>Stanford University, Stanford, CA 94305, USA. <sup>57</sup>Department of Psychology, University of Alabama, Tuscaloosa, AL 35487, USA. <sup>58</sup>Carnegie Mellon University, Pittsburgh, PA 15213, USA. <sup>59</sup>Department FISPPA—Applied Psychology Unit, University of Padua, Padova, Italy. <sup>60</sup>Universidad Nacional De Asunción, Asunción, Paraguay. <sup>61</sup>Bard College, Providence, RI 02906, USA. <sup>62</sup>Northeastern University, Boston, MA 02115, USA. <sup>63</sup>Department of Counseling and Educational Psychology, Mississippi State University,

Mississippi State, MS 39762, USA. <sup>64</sup>School of Psychology, University of Sydney, Sydney, Australia. <sup>65</sup>Hampshire College, Amherst, MA 01002, USA. <sup>66</sup>Colorado State University-Pueblo, Pueblo, CO 81001, USA. <sup>67</sup>School of Psychology, Keele University, Keele, Staffordshire, UK. <sup>68</sup>Behavioral Science Institute, Nijmegen, Netherlands. <sup>69</sup>University of Koblenz-Landau, Landau, Germany. <sup>70</sup>Department of Biostatistics, Institute of Psychiatry, Psychology, and Neuroscience, IHR Biomedical Research Centre for Mental Health, South London, London, UK. <sup>71</sup>Maudsley NHS Foundation Trust, King's College London, London, UK. <sup>72</sup>Department of Psychology, University of Winchester, Winchester, UK. <sup>73</sup>University of Nevada, Las Vegas, Las Vegas, NV 89154, USA. <sup>74</sup>Institute for Multimedia and Interactive Systems, University of Lübeck, Lübeck, Germany. <sup>75</sup>Institute of Cognitive Science, University of Osnabrück, Osnabrück, Germany. <sup>76</sup>University of Birmingham, Northampton, Northamptonshire NN1 3NB, UK. <sup>77</sup>University of Chicago, Chicago, IL 60615, USA. <sup>78</sup>London School of Economics and Political Science, London WC2A 2AE, UK. <sup>79</sup>Loyola University Maryland, Baltimore, MD 21210, USA. <sup>80</sup>Department of Security and Crime Science, University College London, London WC1H 9EZ, UK. <sup>81</sup>Department of Psychology, University of Belgrade, Belgrade, Serbia. <sup>82</sup>Department of Psychology, University of Konstanz, Konstanz, Germany. <sup>83</sup>Open Science Collaboration, Saratoga, CA, USA. <sup>84</sup>School of Innovation Sciences, Eindhoven University of Technology, Eindhoven, Netherlands. <sup>85</sup>Department of Psychology, University of Iowa, Iowa City, IA 52242, USA. <sup>86</sup>Department of Psychology, Western University, London, ON N6A 5C2, Canada. <sup>87</sup>Department of Psychology, Western Washington University, Bellingham, WA 98225, USA. <sup>88</sup>Department of Cognitive Science, Occidental College, Los Angeles, CA 90041, USA. <sup>89</sup>Department of Psychology, Reed College, Portland, OR 97202, USA. <sup>90</sup>Department of Psychology and Neuroscience, Dalhousie University, Halifax, NS, Canada. <sup>91</sup>Renison University College at University of Waterloo, Waterloo, ON N2L 3G4, Canada. <sup>92</sup>Department of Psychology, Wake Forest University, Winston-Salem, NC 27109, USA. <sup>93</sup>Counseling and Psychological Services, California State University, Fullerton, Fullerton, CA 92831, USA. <sup>94</sup>University of Bonn, Bonn, Germany. <sup>95</sup>University of Potsdam, Potsdam, Germany. <sup>96</sup>School of Psychology, University of Dundee, Dundee, Scotland. <sup>97</sup>Willamette University, Salem, OR 97301, USA. <sup>98</sup>Arcadia University, Glenside, PA 19038, USA. <sup>99</sup>Department of Psychology, Technische Universität Dresden, Dresden, Germany. <sup>100</sup>Department of Psychology, University of Illinois at Chicago, Chicago, IL 60607, USA. <sup>101</sup>William and Mary, Williamsburg, VA 23185, USA. <sup>102</sup>Stockholm University, Stockholm, Sweden. <sup>103</sup>Karolinska Institute, Stockholm, Sweden. <sup>104</sup>Center for Neural Science, New York University, New York, NY 10003, USA. <sup>105</sup>Centre for Clinical Brain Sciences, The University of Edinburgh, Edinburgh EH16 4SB, UK. <sup>106</sup>Department of Neurology, Beth Israel Deaconess Medical Center, Harvard Medical School, Boston, MA 02215, USA. <sup>107</sup>Department of Psychology, University of Florida, Gainesville, FL 32611, USA. <sup>108</sup>Department of Psychology, Wright State University, Dayton, OH 45435, USA. <sup>109</sup>California State University, Northridge, Northridge, CA 91330, USA. <sup>110</sup>Kutztown University of Pennsylvania, Kutztown, PA 19530, USA. <sup>111</sup>Department of Psychology, Universidad Autónoma de Madrid, Madrid, Spain. <sup>112</sup>IE Business School, Madrid, Spain. <sup>113</sup>Department of Social Sciences, University of Virginia's College at Wise, Wise, VA 24230, USA. <sup>114</sup>University of Groningen, Groningen, Netherlands. <sup>115</sup>Department of Behavioral Sciences, Ariel University, Ariel, 40700 Israel. <sup>116</sup>Department of Social Science, Worcester Polytechnic Institute, Worcester, MA 01609, USA. <sup>117</sup>Department of Basic Psychological Research and Research Methods, University of Vienna, Vienna, Austria. <sup>118</sup>Duke University, Durham, NC 27708, USA. <sup>119</sup>Centre for Research in Psychology, Behavior and Achievement, Coventry University, Coventry CV1 5FB, UK. <sup>120</sup>The Open University, Buckinghamshire MK7 6AA, UK. <sup>121</sup>City University of Hong Kong, Shamsuipo, KLN, Hong Kong. <sup>122</sup>Jacksonville University, Jacksonville, FL 32211, USA. <sup>123</sup>TIBER (Tilburg Institute for Behavioral Economics Research), Tilburg, Netherlands. <sup>124</sup>Universidad de la República Uruguay, Montevideo 11200, Uruguay. <sup>125</sup>University of Oxford, Oxford, UK.

## SUPPLEMENTARY MATERIALS

[www.sciencemag.org/content/349/6251/aac4716/suppl/DC1](http://www.sciencemag.org/content/349/6251/aac4716/suppl/DC1)

Figures S1 to S7

Tables S1 to S4

References (38–41)

29 April 2015; accepted 28 July 2015

10.1126/science.aac4716

## Estimating the reproducibility of psychological science

Open Science Collaboration

*Science* **349** (6251), aac4716.  
DOI: 10.1126/science.aac4716

### Empirically analyzing empirical evidence

One of the central goals in any scientific endeavor is to understand causality. Experiments that seek to demonstrate a cause/effect relation most often manipulate the postulated causal factor. Aarts *et al.* describe the replication of 100 experiments reported in papers published in 2008 in three high-ranking psychology journals. Assessing whether the replication and the original experiment yielded the same result according to several criteria, they find that about one-third to one-half of the original findings were also observed in the replication study.

*Science*, this issue 10.1126/science.aac4716

#### ARTICLE TOOLS

<http://science.sciencemag.org/content/349/6251/aac4716>

#### SUPPLEMENTARY MATERIALS

<http://science.sciencemag.org/content/suppl/2015/08/26/349.6251.aac4716.DC1>

#### RELATED CONTENT

<http://science.sciencemag.org/content/sci/349/6251/999.2.full>  
<http://science.sciencemag.org/content/sci/349/6251/910.full>  
<http://science.sciencemag.org/content/sci/348/6242/1422.full>  
<http://science.sciencemag.org/content/sci/343/6168/229.full>  
<http://science.sciencemag.org/content/sci/348/6242/1403.full>  
<http://science.sciencemag.org/content/sci/351/6277/1037.2.full>  
<http://science.sciencemag.org/content/sci/351/6277/1037.3.full>

#### REFERENCES

This article cites 37 articles, 3 of which you can access for free  
<http://science.sciencemag.org/content/349/6251/aac4716#BIBL>

#### PERMISSIONS

<http://www.sciencemag.org/help/reprints-and-permissions>

Use of this article is subject to the [Terms of Service](#)



# Supplementary Material for

## Estimating the reproducibility of psychological science

Open Science Collaboration\*

\*Corresponding author. E-mail: [nosek@virginia.edu](mailto:nosek@virginia.edu)

Published 28 August 2015, *Science* **349**, aac4716 (2015)  
DOI: 10.1126/science.aac4716

**This PDF file includes:**

Materials and Methods  
Figs. S1 to S7  
Tables S1 to S4

**Supplemental Information for  
Estimating the Reproducibility of Psychological Science**

**Open Science Collaboration**

**Table of Contents**

1. [Method](#)
  - a. [Replication Teams](#)
  - b. [Replication Protocol](#)
2. [Measures and Moderators](#)
  - a. [Characteristics of Original Study](#)
  - b. [Characteristics of Replication](#)
3. [Guide to the Information Commons](#)
4. [Results](#)
  - a. [Preliminary Analyses](#)
  - b. [Evaluating replication against null hypothesis](#)
  - c. [Evaluating replication against original effect size](#)
  - d. [Comparing original and replication effect sizes](#)
  - e. [Combining original and replication effect sizes for cumulative evidence](#)
  - f. [Subjective assessment: Did it replicate?](#)
  - g. [Meta-analysis of all original study effect, and of all replication study effects](#)
  - h. [Meta-analysis of difference of effect size between original and replication study](#)
  - i. [Moderator analyses](#)

## Method

Two articles have been published on the methodology of the Reproducibility Project: Psychology.

1. Open Science Collaboration, An open, large-scale, collaborative effort to estimate the reproducibility of psychological science. *Perspect. Psychol. Sci.* **7**, 657-660 (2012).
2. Open Science Collaboration, The Reproducibility Project: A Model of Large-Scale Collaboration for Empirical Research on Reproducibility. In *Implementing Reproducible Computational Research (A Volume in The R Series)*, V. Stodden, F. Leisch, R. Peng, Eds. (Taylor & Francis, New York, 2014) pp. 299-323.

The first introduced the project aims and basic design. The second provided detail on the methodology and mechanisms for maintaining standards and quality control. The methods sections in the main text and below summarize the key aspects of the methodology and provide additional information, particularly concerning the latter stages of the project that were not addressed in the prior articles.

### Replication Teams

RPP was introduced publicly as a crowdsourcing research project in November 2011. Interested researchers were invited to get involved to design the project, conduct a replication, or provide other kinds of research support such as coding articles. A total of 270 individuals contributed sufficiently to earn co-authorship on this report.

Of the 100 replications completed, 85 unique senior members were identified—several of whom led multiple replications. Among those senior members, 72 had a PhD or equivalent, 9 had a master's degree or equivalent, 1 had some graduate school, and 3 had or were near completing a bachelor's degree or equivalent. By occupation, 62 were faculty members or equivalent, 8 were post-docs, 13 were graduate students, 1 was an undergraduate student, and 1 was a private sector researcher. By domain, 36 identified social psychology as their primary domain, 22 identified cognitive psychology, 6 identified quantitative psychology, and 21 identified other domains.

### Replication Protocol

Sloppy or underpowered replication attempts would provide uninteresting reasons for irreproducibility. Replication teams followed an extensive protocol to maximize quality, clarity, and standardization of the replications. Full documentation of the protocol is available at <https://osf.io/ru689/>.

**Power analysis.** After identifying the key effect, power analyses estimated the sample sizes needed to achieve 80%, 90%, and 95% power to detect the originally reported effect size. Teams were required to propose a study design that would achieve at least 80% power and were encouraged to obtain higher power if feasible to do so. All protocols proposed 80% power

or greater, however, after corrections to power analyses, three fell short in their planning, with 56%, 69%, and 76% power. On average, 92% power was proposed (median = 95%). Three replication teams were unable to conduct power analyses based on available data—their method for planning sample size is detailed in their replication reports. Following data collection, 90 of the 97 achieved 80% or greater power to detect the original effect size. Post-hoc calculations showed an average of 92% power to detect an effect size equivalent to the original studies'. The median power was 95% and 57 had 95% power or better. Note that these power estimates do not account for the possibility that the published effect sizes are overestimated because of publication biases. Indeed, this is one of the potential challenges for reproducibility.

**Obtaining or creating materials.** Project coordinators or replication teams contacted original authors for study materials in order to maximize the consistency between the original and replication effort. Of the completed replications, 89 were able to obtain some or all of the original materials. In 8 cases, the original materials were not available, and in only 3 cases the original authors did not share materials or provide information about where the materials could be obtained. Replication teams prepared materials, adapting or creating them for the particular data collection context. If information available from the original report or author contacts was insufficient, teams noted deviations or inferences in their written protocols.

**Writing study protocols.** The protocols included a brief introduction explaining the main idea of the study, the key finding for replication, and any other essential information about the study. Then, they had a complete methods section describing the power analysis, sampling plan, procedure, materials, and analysis plan. Analysis plans included details of data exclusion rules, data cleaning, inclusion of covariates in the model, and the inferential test/model that would be used. Finally, the protocol listed known differences from the original study in sampling, setting, procedure, and analysis plan. The objective was to minimize differences that are expected to alter the effect, but report transparently about them to provide a means of identifying possible reasons for variation in observed effects, and to identify factors for establishing generalizability of the results when similar effects are obtained. All replication teams completed a study protocol in advance of data collection.

Replication teams were encouraged to apply for funding for the replication to the Center for Open Science (<http://cos.io/>). A grants committee comprised of members of the collaboration reviewed study protocols made award recommendations.

**Reviewing study protocols.** The written protocols were shared with original authors for critique prior to initiating data collection. Also, protocols were reviewed by another member of the RPP team for quality assurance and consistency with the reporting template. Feedback from the original authors was incorporated into the study design. If the replication team could not address the feedback, the original author comments were included in the protocol so that readers could identify the *a priori* comments by original authors about the design. Replication teams recorded whether the original authors endorsed the design (69 replications), maintained concerns based on informed judgment/speculation (8 replications), maintained concerns based on published empirical evidence for constraints on the effect (3 replications), or did not respond (18 replications). Two replications did not seek and receive feedback prior to data collection.

**Uploading the study protocol.** Once finalized, the protocol and shareable materials were posted publicly on the Open Science Framework (OSF; <https://osf.io/ezcui/>) following a standard format. If the original author requested to keep materials private, replication teams

noted this and indicated how to contact the original author to obtain the materials. After upload, the replication team could begin data collection.

**Reporting.** Following data collection, teams initiated report writing and data sharing. If there were any deviations from the registered protocol, teams noted those in the final report. Also, teams posted anonymized datasets and a codebook to the OSF project page. Teams conducted the planned data analysis from the protocol as a confirmatory analysis. Following completion of the confirmatory analysis phase, teams were encouraged to conduct follow-up exploratory analysis if they wished and report both—clearly distinguished—in their final report.

After writing the results section of the final report, teams added discussion with open-ended commentary about insights gained from exploratory analysis, an overall assessment of the outcome of the replication attempt, and discussion of any objections or challenges raised by the original authors' review of the protocol. At least one other RPP member then conducted a review of the final report to maximize consistency in reporting format, identify errors, and improve clarity. Following review, replication teams shared their report directly with the original authors and publicly on the OSF project page. If additional issues came up following posting of the report, teams could post a revision of the report. The OSF offers version control so all prior versions of posted reports can be retrieved in order to promote transparent review of edits and improvements.

## **Measures and Moderators**

### **Characteristics of Original Study**

**Original study effect size,  $p$ -value, and sample size.** Qualities of the original statistical evidence may predict reproducibility. All else being equal, results with larger effect sizes and smaller  $p$ -values ought to be more reproducible than others. Also, larger sample sizes are a factor for increasing the precision of estimating effects; all else being equal, larger sample sizes should be associated with more reproducible results. A qualification of this expectation is that some study designs use very few participants and gain substantial power via repeated measurements.

**Importance of the result.** Some effects are more important than others. This variable was the aggregate of the citation impact of the original article and coder ratings of the extent to which the article was exciting and important. Effect importance could be a positive predictor of reproducibility because findings that have a strong impact on the field do so, in part, because they are reproducible and spur additional innovation. If they were not reproducible, then they may not have a strong impact on the field. On the other hand, exciting or important results are appealing because they advance an area of research, but they may be less reproducible than mundane results because true advances are difficult and infrequent, and theories and methodologies employed at the fringe of knowledge are often less refined or validated making them more difficult to reproduce.

*Citation impact of original article.* Project coordinators used Google Scholar data to calculate the citation impact of the original article at the time of conducting the project analysis (March 2015).

**Exciting/important effect.** Coders independent from the replication teams reviewed the methodology for the replication studies and answered the following prompt: “To what extent is the key effect an exciting and important outcome?” To answer this question, coders read the pre-data collection reports that the replication teams had created. These reports included a background on the topic, a description of the effect, a procedure, and analysis plan. Responses were provided on a scale from 1 = Not at all exciting and important, 2 = Slightly exciting and important, 3 = Somewhat exciting and important, 4 = Moderately exciting and important, 5 = Very exciting and important, 6 = Extremely exciting and important. One-hundred twenty nine coders were presented effect reports and these questions for 112 studies (100 replications reported in the main text + others for which data collection was in progress) in a random order, and coders rated as many as they wished. Each effect was rated an average of 4.52 times (median = 4). Ratings were averaged across coders.

**Surprising result.** Counterintuitive results are appealing because they violate one’s priors, but they may be less reproducible if priors are reasonably well-tuned to reality. The same coders that rated the extent to which the effect was exciting/important reviewed the methodology for the replication studies and answered the following prompt: “To what extent is the key effect a surprising or counterintuitive outcome?” Responses were provided on a scale from 1 = Not at all surprising, 2 = Slightly surprising, 3 = Somewhat surprising, 4 = Moderately surprising, 5 = Very surprising, 6 = Extremely surprising.

**Experience and expertise of original team.** Higher quality teams may produce more reproducible results. Quality is multi-faceted and difficult to measure. In the present study, after standardizing we averaged four indicators of quality - the rated prestige of home institutions of the 1st and senior authors, and the citation impact of the 1st and senior authors. Other means of assessing quality could reveal results quite distinct from those obtained by these indicators.

**Institution prestige of 1st author and senior author.** Authors were coded as being 1st and most senior; their corresponding institutions were also recorded. The resulting list was presented to two samples (Mechanical Turk participants  $n = 108$ ; Project team members  $n = 70$ ) to rate institution prestige on a scale from 7 = never heard of this institution, 6 = not at all prestigious, 5 = slightly prestigious, 4 = moderately prestigious, 3 = very prestigious, 2 = extremely prestigious, 1 = one of the few most prestigious. MTurk participants rated institution prestige in general. Project team members were randomly assigned to rate institution prestige *in psychology* ( $n = 33$ ) or *in general* ( $n = 37$ ). Correlations of prestige ratings among the three samples were very high ( $r$ ’s range .849 to .938). As such, before standardizing, we averaged the three ratings for a composite institution prestige score.

**Citation impact of 1st author and senior author.** Project members used Google Scholar data to estimate the citation impact of first authors and senior authors. These indicators identified citation impact at the time of writing this report, not at the time the original research was conducted.

## Characteristics of Replication

**Replication power and sample size.** All else equal, lower power and smaller sample tests ought to be less likely to reproduce results than higher power and larger sample tests.

The caveat above on sample size for original studies is the same as for replication studies. Replications were required to achieve at least 80% power based on the effect size of the original study. This narrows the range of actual power in replication tests to maximize likelihood of obtaining effects, but nonetheless offers a range that could be predictive of reproducibility. A qualification of this expectation is that power estimates are based on original effects. If publication bias or other biases produce exaggerated effect sizes in the original studies, then the power estimates would be less likely to provide predictive power for reproducibility.

**Challenge of conducting replication.** Reproducibility depends on effective implementation and execution of the research methodology. However, some methodologies are more challenging or prone to error and bias than others. As a consequence, variation in the challenges of conducting replications may be a predictor of reproducibility. This indicator includes coders' assessments of expertise required, opportunity for experimenter expectations to influence outcomes, and opportunity for lack of diligence to influence outcomes. Of course these issues apply to conducting the original study and interpreting its results, but we treated these as characteristics of the replication for the present purposes.

For these variables, a small group of coders were trained on evaluating original reports and a single coder evaluated each study.

*Perceived expertise required.* Reproducibility might be lower for study designs that require specialized expertise. Coders independent from the replication teams reviewed the methodology for the replication studies and answered the following prompt: "To what extent does the methodology of the study require specialized expertise to conduct effectively? [Note: This refers to data collection, *not* data analysis]" Responses were provided on a scale from 1 = no expertise required, 2 = slight expertise required, 3 = moderate expertise required, 4 = strong expertise required, 5 = extreme expertise required.

*Perceived opportunity for expectancy biases.* The expectations of the experimenter can influence study outcomes (38). Study designs that provide opportunity for researchers' beliefs to influence data collection may be more prone to reproducibility challenges than study designs that avoid opportunity for influence. Coders independent from the replication teams reviewed the methodology for the replication studies and answered the following prompt: "To what extent does the methodology of the study provide opportunity for the researchers' expectations about the effect to influence the results? (i.e., researchers belief that the effect will occur could elicit the effect, or researchers belief that the effect will not occur could eliminate the effect) [Note: This refers to data collection, *not* data analysis]." Responses were provided on a scale from 1 = No opportunity for researcher expectations to influence results, 2 = Slight opportunity for researcher expectations to influence results, 3 = Moderate opportunity for researcher expectations to influence results, 4 = Strong opportunity for researcher expectations to influence results, 5 = Extreme opportunity for researcher expectations to influence results.

*Perceived opportunity for impact of lack of diligence.* Studies may be less likely to be reproducible if they are highly reliant on experimenters' diligence to conduct the procedures effectively. Coders independent from the replication teams reviewed the methodology for the replication studies and answered the following prompt: "To what extent could the results be affected by lack of diligence by experimenters in collecting the data? [Note: This refers to data collection, *not* creating the materials]." Responses were provided on a scale from 1 = No opportunity for lack of diligence to affect the results, 2 = Slight opportunity for lack of diligence to

affect the results, 3 = Moderate opportunity for lack of diligence to affect the results, 4 = Strong opportunity for lack of diligence to affect the results, 5 = Extreme opportunity for lack of diligence to affect the results.

**Experience and expertise of replication team.** Just as experience and expertise may be necessary to obtain reproducible results, expertise and experience may be important for conducting effective replications. We focused on the senior member of the replication team and created an aggregate by standardizing and averaging scores on 7 characteristics: position (undergraduate to professor), highest degree (high school to PhD or equivalent), self-rated domain expertise, self-rated method expertise, total number of publications, total number of peer-reviewed empirical articles, and citation impact.

*Position of senior member of replication team.* Reproducibility may be enhanced by having more seasoned researchers guiding the research process. Replication teams reported the position of the senior member of the team from: 7 = Professor (or equivalent), 6 = Associate Professor (or equivalent), 5 = Assistant Professor (or equivalent), 4 = Post-doc, Research Scientist, or Private Sector Researcher, 3 = Ph.D. student, 2 = Master's student, 1 = Undergraduate student, or other.

*Highest degree of replication team's senior member.* Replication teams reported the highest degree obtained by the senior member of the team from 4 = PhD/equivalent, 3 = Master's/equivalent, 2 = some graduate school, 1 = Bachelor's/equivalent.

*Replication team domain expertise.* Reproducibility may be stronger if the replication team is led by a person with high domain expertise in the topic of study. Replication teams self-rated the domain expertise of the senior member of the project on the following scale: 1 = No expertise - No formal training or experience in the topic area, 2 = Slight expertise - Researchers exposed to the topic area (e.g., took a class), but without direct experience researching it, 3 = Some expertise - Researchers who have done research in the topic area, but have not published in it, 4 = Moderate expertise - Researchers who have previously published in the topic area of the selected effect, and do so irregularly, 5 = High expertise - Researchers who have previously published in the topic area of the selected effect, and do so regularly.

*Replication team method expertise.* Reproducibility may be stronger if the replication team is led by a person with high expertise in the methodology used for the study. Replication teams self-rated the domain expertise of the senior member of the project on the following scale: 1 = No expertise - No formal training or experience with the methodology, 2 = Slight expertise - Researchers exposed to the methodology, but without direct experience using it, 3 = Some expertise - Researchers who have used the methodology in their research, but have not published with it, 4 = Moderate expertise - Researchers who have previously published using the methodology of the selected effect, and use the methodology irregularly, 5 = High expertise - Researchers who have previously published using the methodology of the selected effect, and use the methodology regularly.

*Replication team senior member's total publications and total number of peer-reviewed articles.* All else being equal, more seasoned researchers may be better prepared to reproduce research results than more novice researchers. Replication teams self-reported the total number of publications and total number of peer-reviewed articles by the senior member of the team.

*Institution prestige of replication 1st author and senior author.* We followed the same methodology for computing institution prestige for replication teams as we did for original author teams.

*Citation impact of replication 1st author and senior author.* Researchers who have conducted more research that has impacted other research via citation may have done so because of additional expertise and effectiveness in conducting reproducible research. Project members calculated the total citations of the 1st author and most senior member of the team via Google Scholar.

**Self-assessed quality of replication.** Lower quality replications may produce results less similar to original effects than higher quality replications. Replication teams are in the best position to know the quality of project execution, but are also likely to be ego invested in reporting high quality. Nonetheless, variation in self-assessed quality across teams may provide a useful indicator of quality. Also, some of our measures encouraged variation in quality reports by contrasting directly with the original study, or studies in general. After standardizing, we created an aggregate score by averaging four variables: self-assessed quality of implementation, self-assessed quality of data collection, self-assessed similarity to original, and self-assessed difficulty of implementation. Future research may assess additional quality indicators from the public disclosure of methods to complement this assessment.

*Self-assessed implementation quality of replication.* Sloppy replications may be less likely to reproduce original results because of error and inattention. Replication teams self-assessed the quality of the replication study methodology and procedure design in comparison to the original research by answering the following prompt: "To what extent do you think that the replication study materials and procedure were designed and implemented effectively? Implementation of the replication materials and procedure..." Responses were provided on a scale from 1 = was of much higher quality than the original study, 2 = was of moderately higher quality than the original study, 3 = was of slightly higher quality than the original study, 4 = was about the same quality as the original study, 5 = was of slightly lower quality than the original study, 6 = was of moderately lower quality than the original study, 7 = was of much lower quality than the original study.

*Self-assessed data collection quality of replication.* Sloppy replications may be less likely to reproduce original results because of error and inattention. Replication teams self-assessed the quality of the replication study data collection in comparison to the average study by answering the following prompt: "To what extent do you think that the replication study data collection was completed effectively for studies of this type?" Responses were provided on a scale from 1 = Data collection quality was much better than the average study, 2 = Data collection quality was better than the average study, 3 = Data collection quality was slightly better than the average study, 4 = Data collection quality was about the same as the average study, 5 = Data collection quality was slightly worse than the average study, 6 = Data collection quality was worse than the average study, 7 = Data collection quality was much worse than the average study.

*Self-assessed replication similarity to original.* It can be difficult to reproduce the conditions and procedures of the original research for a variety of reasons. Studies that are more similar to the original research may be more reproducible than those that are more dissimilar. Replication teams self-evaluated the similarity of the replication with the original by

answering the following prompt: “Overall, how much did the replication methodology resemble the original study?” Responses were provided on a scale from 1 = Not at all similar, 2 = Slightly similar, 3 = Somewhat similar, 4 = Moderately similar, 5 = Very similar, 6 = Extremely similar, 7 = Essentially identical.

*Self-assessed difficulty of implementation.* Another indicator of adherence to the original protocol is the replication team’s self-assessment of how challenging it was to conduct the replication. Replication teams responded to the following prompt: “How challenging was it to implement the replication study methodology?” Responses were provided on a scale from 1 = Extremely challenging, 2 = Very challenging, 3 = Moderately challenging, 4 = Somewhat challenging, 5 = Slightly challenging, 6 = Not at all challenging.

**Other variables.** Some additional variables were collected and appear in the tables not aggregated with other indicators, or are not reported at all in the main text. They are nonetheless available for additional analysis. Below are highlights and a comprehensive summary of additional variables is available in the [Master Data File](#).

*Replication team surprised by outcome of replication.* The replication team rated the extent to which they were surprised by the results of their replication. Teams responded to the following prompt: “To what extent was the replication team surprised by the replication results?” Responses were provided on a scale from 1 = Results were exactly as anticipated, 2 = Results were slightly surprising, 3 = Results were somewhat surprising, 4 = Results were moderately surprising, 5 = Results were extremely surprising. Results are reported in Table S5. Across reproducibility criteria, there was a moderate relationship such that greater surprise with the outcome was associated with weaker reproducibility.

*Effect similarity.* In addition to the subjective “yes/no” assessment of replication in the main text, replication teams provided another rating of the extent to which the key effect in the replication was similar to the original result. Teams responded to the following prompt: “How much did the key effect in the replication resemble the key effect in the original study?” Responses were provided on a scale from: 7 = virtually identical (12), 6 = extremely similar (16), 5 = very similar (8), 4 = moderately similar (12), 3 = somewhat similar (14), 2 = slightly similar (9), 1 = not at all similar (28). Replication results of key effects were deemed between somewhat and moderately similar to the original results,  $M = 3.60$ ,  $SD = 2.18$ .

*Findings similarity.* Replication teams assessed the extent to which the overall findings of the study, not just the key result, were similar to the original study findings. Teams responded to the following prompt: “Overall, how much did the findings in the replication resemble the findings in the original study?” Responses were provided on a scale from: 7 = virtually identical (5), 6 = extremely similar (13), 5 = very similar (21), 4 = moderately similar (20), 3 = somewhat similar (13), 2 = slightly similar (13), 1 = not at all similar (15). Replication results of overall findings were deemed between somewhat and moderately similar to the original results,  $M = 3.78$ ,  $SD = 1.78$ .

*Internal conceptual and direct replications.* Original articles may have contained replications of the key effect in other studies. Coders evaluated whether other studies contained replications of the key result, and whether those replications were direct or conceptual. There were few of both ( $M = .91$  for conceptual replications,  $M = .06$  for direct replications).

## Guide to the Information Commons

There is a substantial collection of materials comprising this project that is publicly accessible for review, critique, and reuse. The following list of links are a guide to the major components.

1. [RPP OSF Project](https://osf.io/ezcuj/): The main repository for all project content is here (<https://osf.io/ezcuj/>)
2. [RPP Information Commons](https://osf.io/ezcuj/wiki/home/): The project background and instructions for replication teams is in the wiki of the main OSF project (<https://osf.io/ezcuj/wiki/home/>)
3. [RPP Researcher Guide](https://osf.io/ru689/): Protocol for replications teams to complete a replication (<https://osf.io/ru689/>)
4. [Master Data File](https://osf.io/5wup8/): Aggregate data across replication studies (<https://osf.io/5wup8/>)
5. Master Analysis Scripts: R script for reproducing analyses for each replication (<https://osf.io/fkmwg/>); R script for reproducing Reproducibility Project: Psychology findings (<https://osf.io/vdnrb/>)
6. Appendices: [Text summaries of analysis scripts](#)

All reports, materials, and data for each replication are available publicly. In a few cases, research materials could not be made available because of copyright. In those cases, a note is available in that project's wiki explaining the lack of access and how to obtain the materials. The following table provides quick links to the projects (with data and materials), final reports, and the R script to reproduce the key finding for all replication experiments.

Two of the articles available for replication were replicated twice (39, 40). The first (39), was replicated in an in lab setting, and online as a secondary replication. The second, experiment 7 of Albarracín et al. (2008) was replicated in a lab setting and a secondary replication of experiment 5 was conducted online. These two supplementary replications bring the total number of replications pursued to 113 and total completed to 100.

## Results

### Preliminary analyses

The input of our analyses were the  $p$ -values (DH and DT in the [Master Data File](#)), their significance (columns EA and EB), effect sizes of both original and replication study (columns DJ and DV), which effect size was larger (column EC), direction of the test (column BU), and whether the sign of both studies' effects was the same or opposite (column BT). First, we checked the consistency of  $p$ -value and test statistics whenever possible (i.e., when all were provided), by recalculating the  $p$ -value using the test statistics. We used the recalculated  $p$ -values in our analysis, with a few exceptions (see Appendix [A1] for details on the recalculation of  $p$ -values). These  $p$ -values were used to code the statistical (non)significance of the effect, with the exception of four effects with  $p$ -values slightly larger than .05 that were interpreted as significant; these studies were treated as significant. We ended up with 99 study-pairs with complete data on  $p$ -values, and 100 study-pairs with complete data on the significance of the replication effect.

*Table S1.* Statistical results (statistically significant or not) of original and replication studies.

## Results

|          |                | Replication    |             |
|----------|----------------|----------------|-------------|
|          |                | Nonsignificant | Significant |
| Original | Nonsignificant | 2              | 1           |
|          | Significant    | 62             | 35          |

The effect sizes (“correlation per df”) were computed using the test statistics (see Appendix [A3] for details on the computation of effect sizes), taking the sign of observed effects into account. Because effect size could not be computed for three study-pairs, we ended up with 97 study-pairs with complete data on effect size. Of the three missing effect sizes, for two could be determined which effect size was larger, hence we ended up with 99 study-pairs with complete data on the comparison of the effect size. Depending on the assessment of replicability, different study-pairs could be included. Seventy-three study-pairs could be included in subset MA, 75 (73+2) could be used to test if the study-pair’s meta-analytic estimate was larger than zero, and 94 (75+19) could be used to determine if the CI of the replication contained the effect size of the original study (see end of Appendix [A3] for an explanation).

### **Evaluating replication effect against null hypothesis of no effect.**

See Appendix [A2] for details. Table S1 shows the statistical significance of original and replication studies. Of the original studies, 97% were statistically significant, as opposed to 36.0% (CI = [26.6%, 46.2%]) of replication studies, which corresponds to a significant change (McNemar test,  $\chi^2(1) = 59.1$ ,  $p < .001$ ).

Proportions of statistical significance of original and replication studies for the three journals JPSP, JEP, PSCI were .969 and .219, .964 and .464, .975 and .4, respectively. Of 97 significant original studies, 36.1% were statistically significant in the replication study. The hypothesis that all 64 statistically non-significant replication studies came from a population of true negatives can be rejected at significance level .05 ( $\chi^2(128) = 155.83$ ,  $p = 0.048$ ).

The density and cumulative  $p$ -value distributions of original and replication studies are presented in Figures S1 and S2 respectively. The means of the two  $p$ -value distributions (.028 and .302) were different from each other ( $t(98) = -8.21$ ,  $p < .001$ ;  $W = 2438$ ,  $p < .001$ ). Quantiles are .00042, .0069, .023 for the original, and .0075, .198, .537 for the replication studies.

*Figure S1:* Cumulative  $p$ -value distributions of original and replication studies.

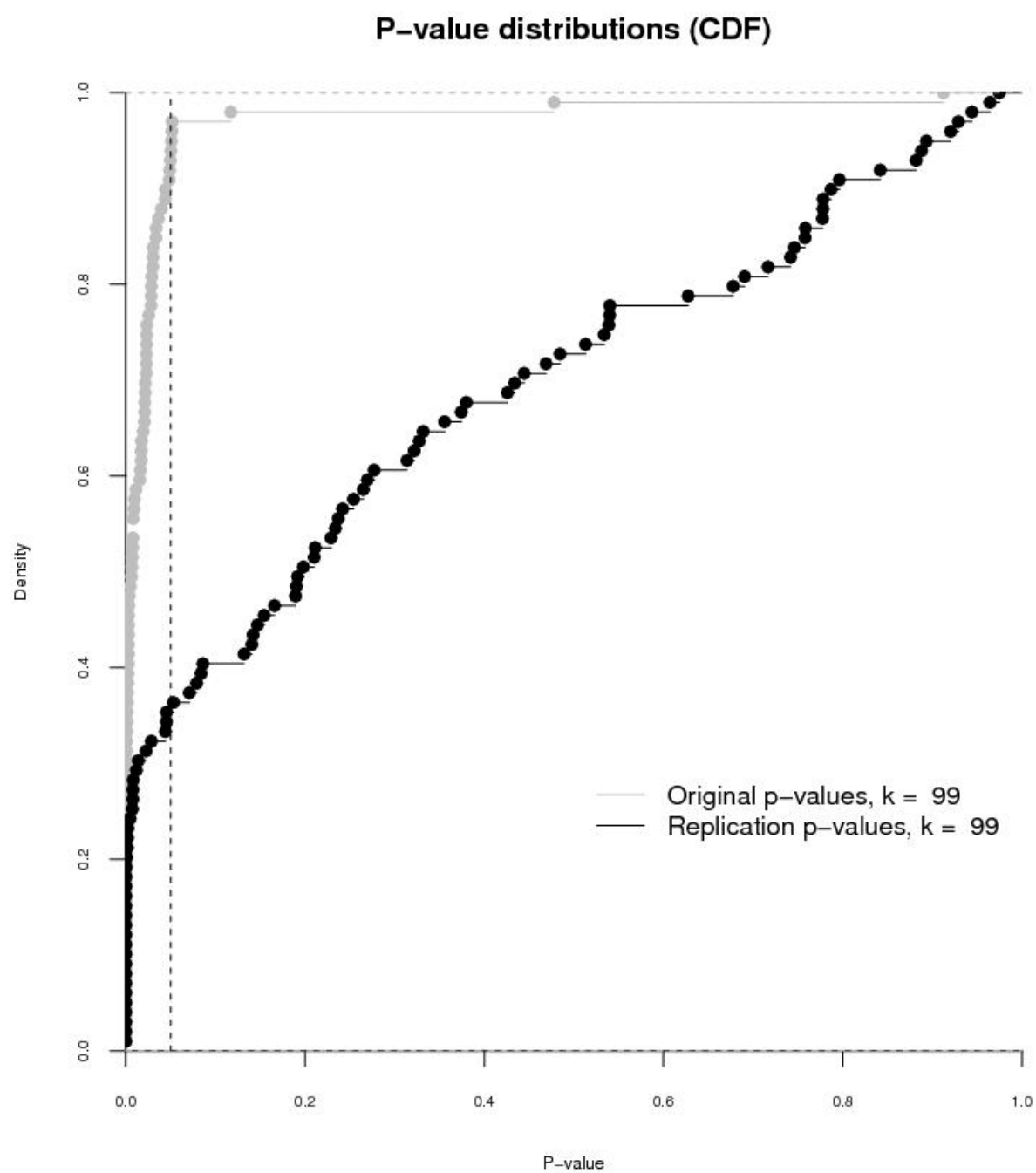
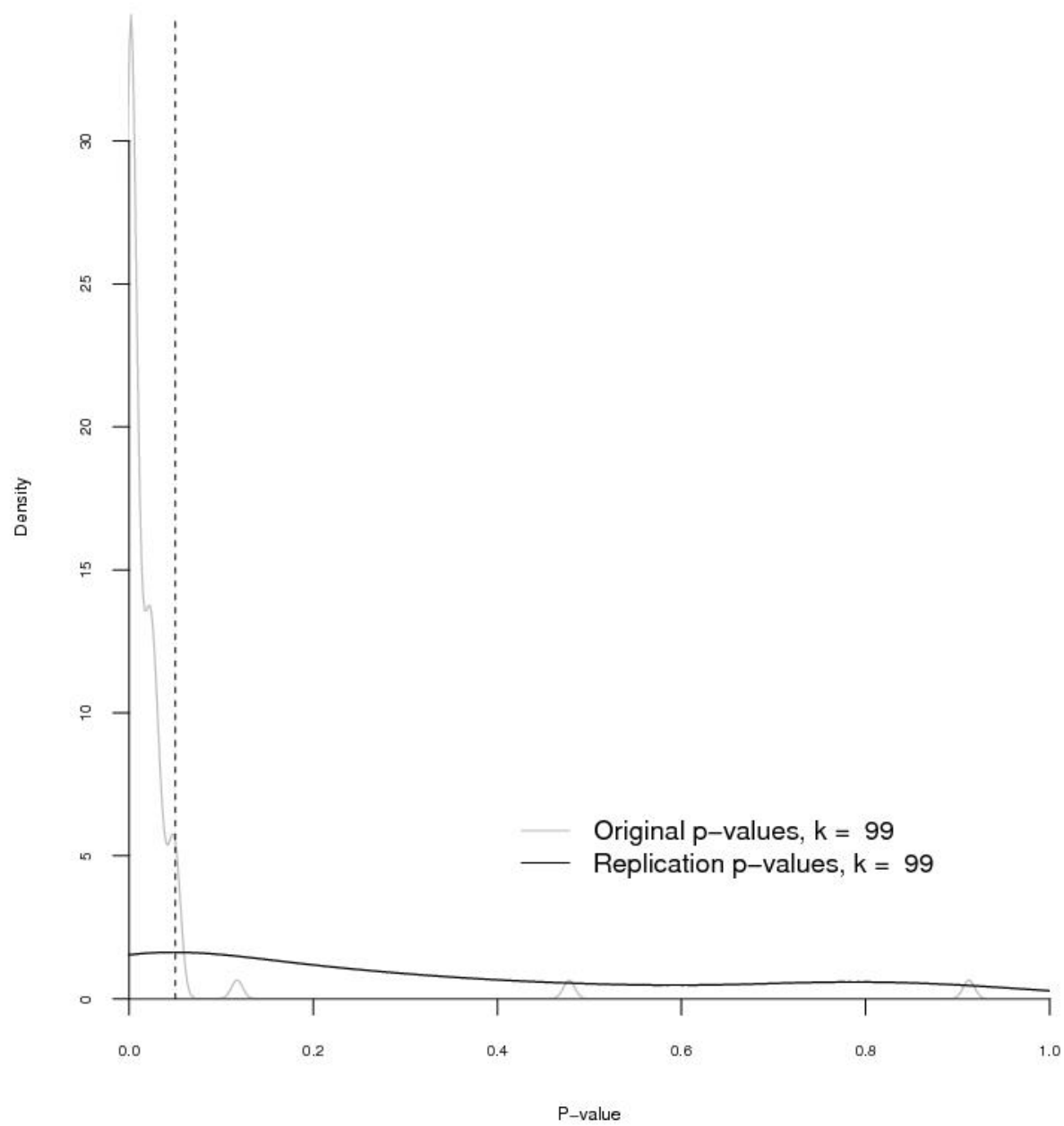


Figure S2: Density  $p$ -value distributions of original and replication studies

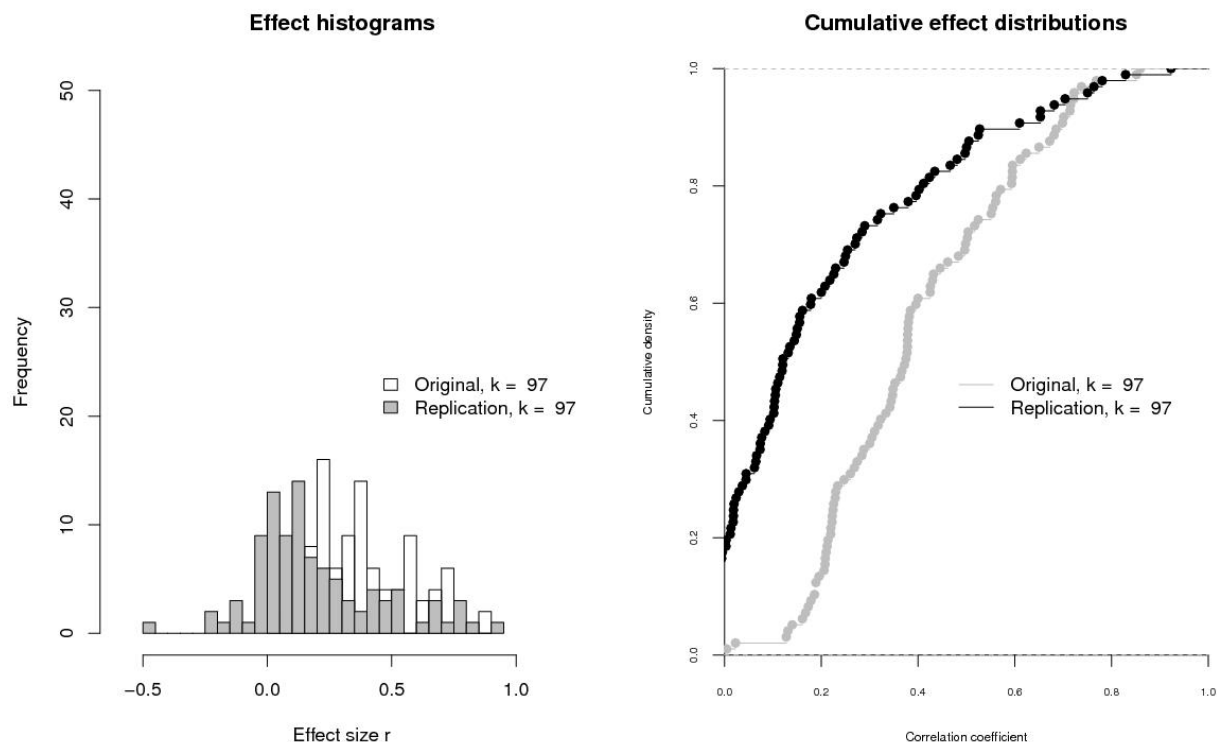
## P-value distributions (PDF)



### Comparing original and replication effect sizes.

See Appendix [A3] and Appendix [A6] for details. For 97 study pairs effect size correlations could be computed. Figure S3 (left) shows the distribution of effect sizes of original and replication studies, and the corresponding cumulative distribution functions (right). The mean effect sizes of both distributions ( $M = .403$  [ $SD = .188$ ];  $M = .197$  [ $SD = .257$ ]) were different from each other ( $t(96) = 9.36, p < .001$ ;  $W = 7137, p < .001$ ). Of those 99 studies that reported an(y) effect size in both original and replication study, 82 reported a larger effect size in the original study (82.8%;  $p < .001$ , binomial test). Original and replication effect sizes were positively correlated (Spearman's  $r = .51, p < .001$ ).

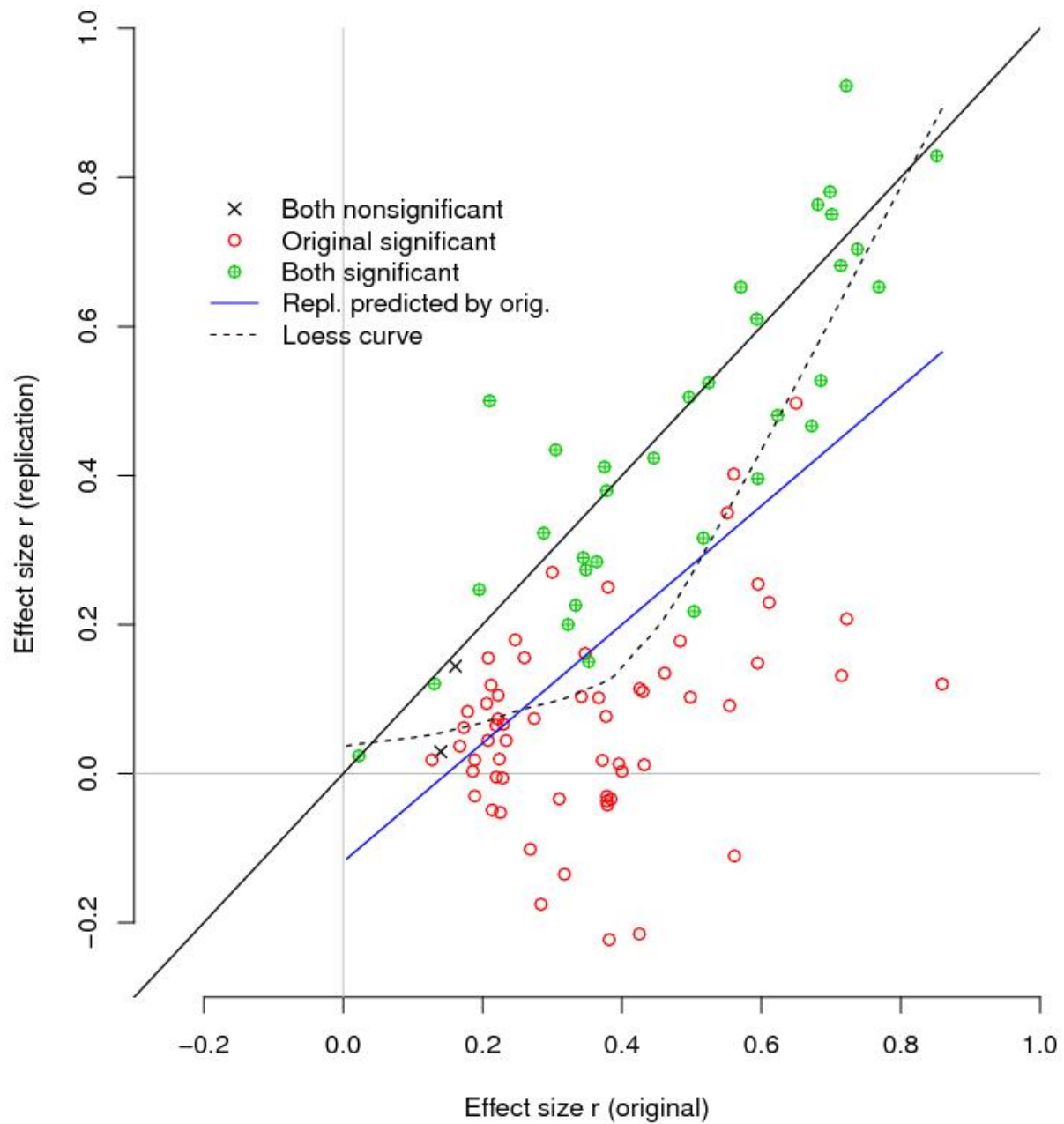
*Figure S3: Distributions (left) and cumulative distribution functions of effect sizes of original and replication studies.*



### Evaluating replication effect against original effect size.

For the subset of 73 studies where the standard error of the correlation could be computed, it was expected that 78.5% of CIs of the replication study contained the effect size of the original study; however, only 41.1% (30 out of 73) of CIs contained the original effect size ( $p < .001$ ) (see [A4] for details). For the subset of 18 and 4 studies with test statistics  $F(df_1 > 1, df_2)$  and  $\chi^2$ , respectively, 68.2% of the confidence intervals contained the effect size of the original study (see [A5] for details). This results in an overall success rate of 47.4%. Figure S4 depicts effect sizes of study-pairs for which correlations could be calculated, and codes significance of effect sizes as well.

Figure S4: Correlations of both original and replication study, coded by statistical significance. Identical values are indicated by the black diagonal line, whereas the blue and dotted line show the replication correlations as predicted by a linear model and loess, respectively.



### **Combining original and replication effect sizes for cumulative evidence.**

See Appendix [A7] for details. For 75 study-pairs a meta-analysis could be conducted on the Fisher-transformed correlation scale. In 51 out of 75 pairs the null-hypothesis of no effect was rejected (68%). The average correlation, after transforming back the Fisher-transformed estimate, was .310 ( $SD = .223$ ). However, the results differed across discipline; average effect size was smaller for JPSP ( $M = .138$ ,  $SD = .087$ ) than for the other four journal/discipline categories, and the percentage of meta-analytic effects rejecting the null-hypothesis was also lowest for JPSP (42.9%; see Table 1). As noted in the main text, the interpretability of these meta-analytic estimates is qualified by the possibility of publication bias inflating the original effect sizes.

### **Subjective assessment of “Did it replicate?”**

Replication teams provided a dichotomous yes/no assessment of whether the effect replicated or not (Column BX). Assessments were very similar to evaluations by significance testing ( $p < .05$ ) including two original null results being interpreted as successful replications when the replication was likewise null, and one original null result being interpreted as a failed replication when the replication showed a significant effect. Overall, there were 39 assessments of successful replication (39 of 100; 39%).

There are three subjective variables assessing replication success. Additional analyses can be conducted on replication teams' assessments of the extent to which key effect and overall findings resemble the original results (Columns CR and CQ).

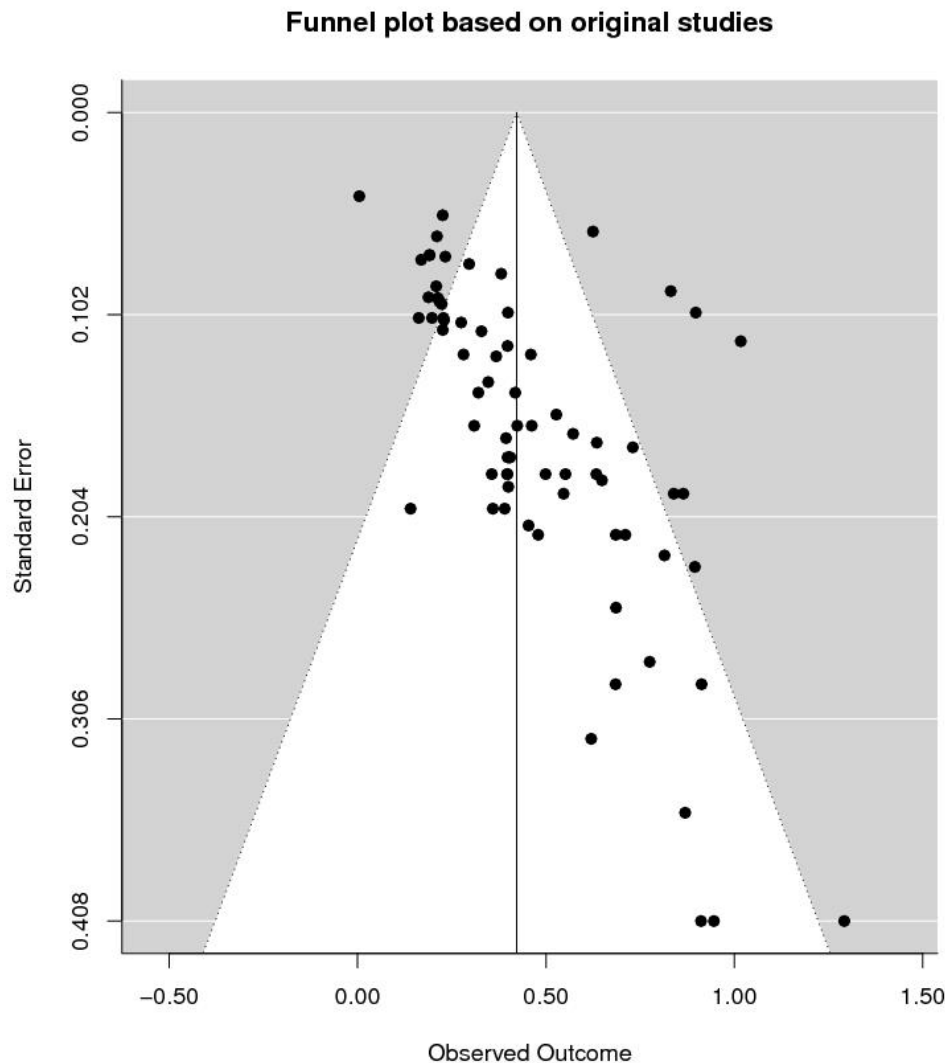
### **Meta-analysis of all original study effects, and of all replication study effects.**

Two random-effects meta-analyses were run (on studies in set MA) using REML estimation for estimating the amount of heterogeneity, one on effect sizes of original and one on effect sizes of replication studies. We ran four models; one without any predictor, one with discipline as predictor, one with studies' standard error as predictor, and one with standard error and discipline as predictor. Discipline is a categorical variable with categories JPSP-social (= reference category), JEP:LMC-cognitive, PSCI-social, and PSCI-cognitive. Standard error was added to examine small-study effects. A positive effect of standard error on effect size indicates that studies' effect sizes are positively associated with their sample sizes. The results of this one-tailed test, also known as Egger's test, is often used as test of publication bias. However, a positive effect of standard error on effect size may also indicate the use of power analysis or using larger sample sizes in fields where smaller effect sizes are observed.

See Appendix [A7] for details. The meta-analysis on all original study effect sizes showed significant ( $Q(72) = 302.67$ ,  $p < .001$ ) and large heterogeneity ( $\hat{\tau} = .19$ ,  $I^2 = 73.3\%$ ), with average effect size equal to .42 ( $z = 14.74$ ,  $p < .001$ ). The average effect size differed across disciplines ( $Q_M(3) = 14.70$ ,  $p = .0021$ ), with effect size in JPSP (.29) being significantly smaller than in JEP:LMC (.52;  $z = 3.17$ ,  $p = .0015$ ) and PSCI-Cog (.57;  $z = 3.11$ ,  $p = .0019$ ), but not PSCI-Soc (.40;  $z = 1.575$ ,  $p = .12$ ). The effect of the original studies' standard error on effect size was

large and highly significant ( $b = 2.24$ ,  $z = 5.66$ ,  $p < .001$ ). Figure S5 shows the funnel plot of the meta-analysis without predictors. After controlling for study's standard error, there was no longer an effect of discipline on effect size ( $\chi^2(3) = 5.36$ ,  $p = .15$ ); at least part of the differences in effect sizes across disciplines was associated with studies in JEP:LMC and PSCI-Cog using smaller sample sizes than JPSP and PSCI-Soc.

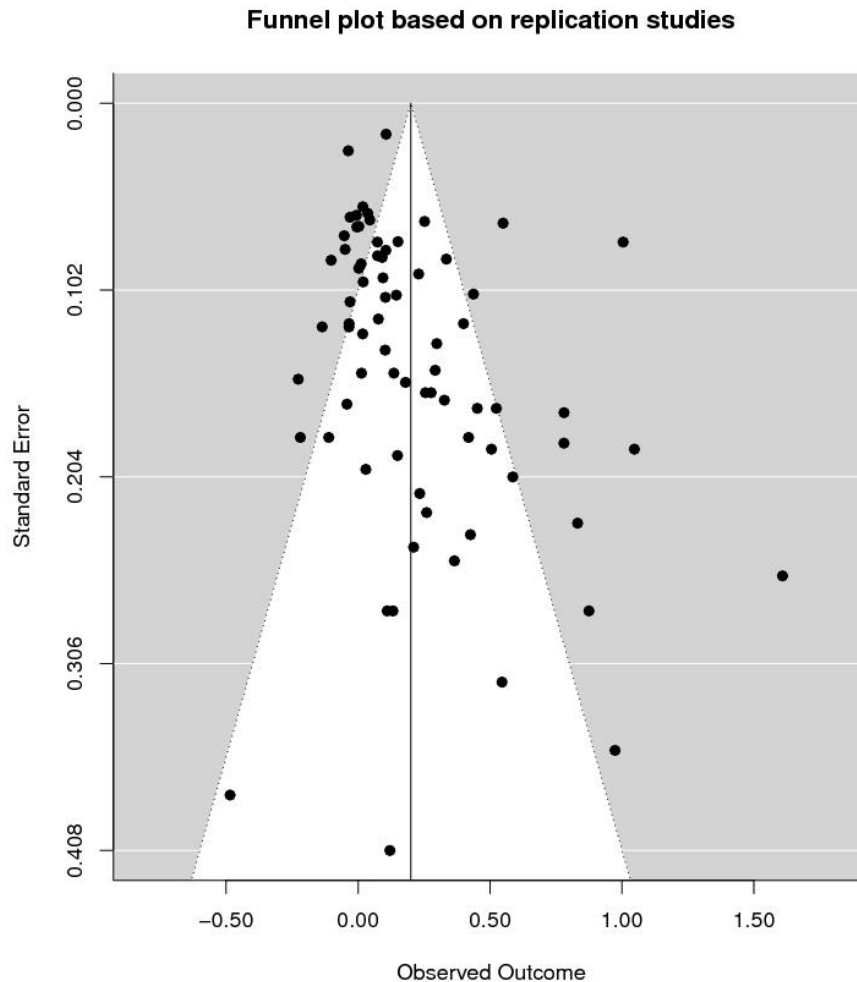
Figure S5: Funnel plot of the meta-analysis on the original study's effect size.



The same meta-analysis on replication studies' effect sizes showed significant ( $Q(72) = 454.00$ ,  $p < .001$ ) and large heterogeneity ( $\hat{\tau} = .26$ ,  $I^2 = 90.1\%$ ), with average effect size equal to .20 ( $z = 5.77$ ,  $p < .001$ ). The average effect size again differed across disciplines ( $Q_M(3) = 12.78$ ,  $p = .0051$ ). Average effect size in JPSP did not differ from 0 (.036;  $z = .63$ ,  $p = .53$ ), and was significantly smaller than average effect size in JEP:LMC (.28;  $z = 2.91$ ,  $p = .0036$ ), PSCI-Cog (.35;  $z = 2.95$ ,  $p = .0032$ ), and PSCI-Soc (.22;  $z = 2.23$ ,  $p = .026$ ). The effect of the standard error of the replication study was large and highly significant ( $b = 1.62$ ,  $z = 3.47$ ,  $p < .001$ ). Because publication bias was absent, this positive effect of standard error was likely caused by

using power analysis for replication studies, i.e., generally larger replication samples were used for smaller true effects. Figure S6 shows the corresponding funnel plot. The effect of discipline did not remain statistically significant after controlling for the standard error of the replication study ( $\chi^2(3) = 6.488, p = .090$ ); similar to the results of original studies, at least part of the differences in effect sizes across disciplines was associated with studies in JEP:LMC and PSCI-Cog using smaller sample sizes than JPSP and PSCI-Soc.

*Figure S6:* Funnel plot of the meta-analysis on the replication study's effect size.



### **Meta-analysis of difference of effect size between original and replication study**

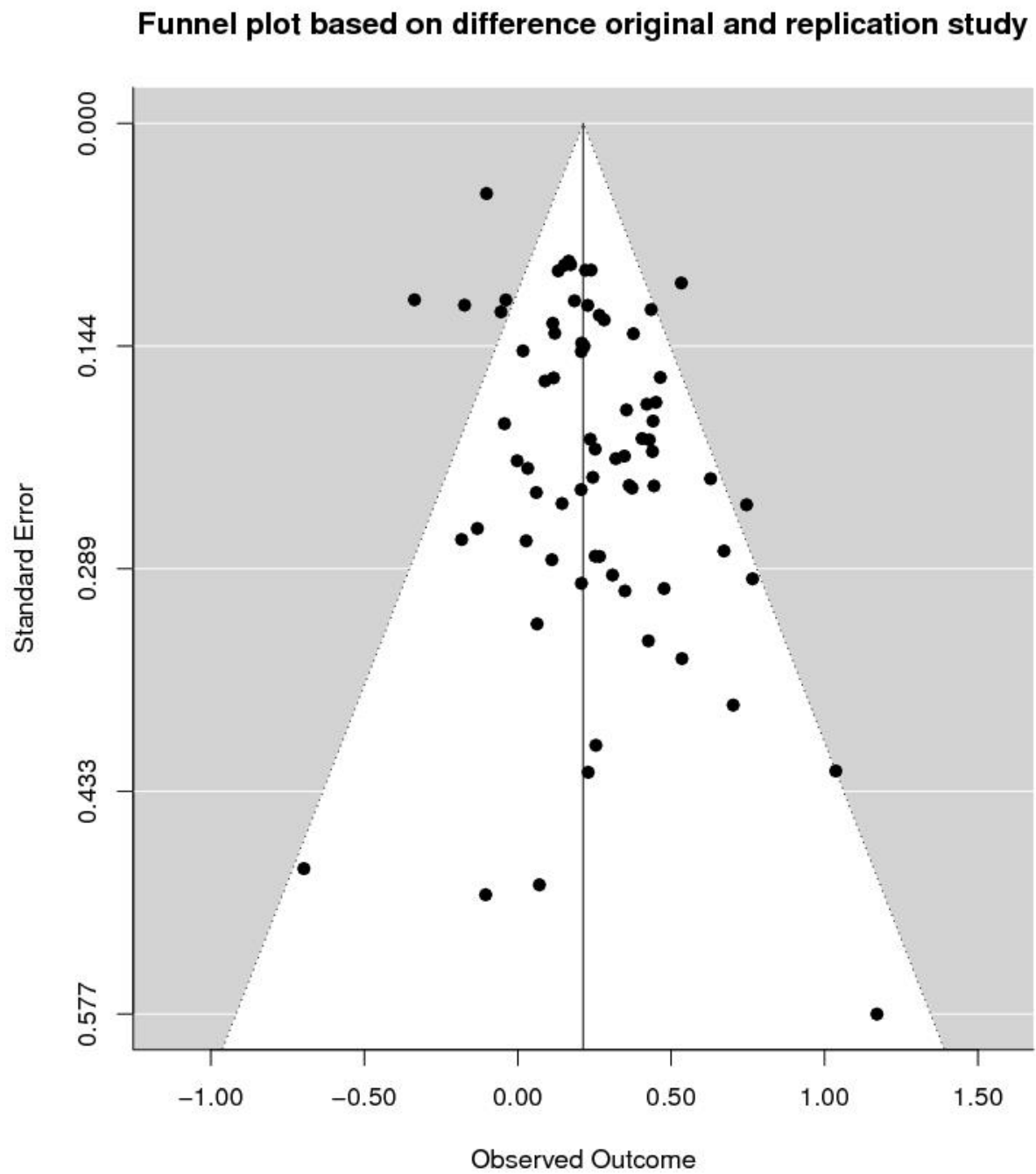
The dependent variable was the difference of Fisher-transformed correlations (original – replication), with variance equal to the sum of variances of the correlation of the original and of the replication study. Several random-effect meta-analyses were run using REML estimation for estimating the amount of heterogeneity in metafor. First, the intercept-only model was estimated; the intercept denotes the average difference effect size between original and replication study. Second, to test for small study effects, we added the standard error of the

original study as a predictor, akin to Egger's test; a positive effect is often interpreted as evidence for publication bias. Our third model tested the effect of discipline.

The null-model without predictors yielded an average estimated difference in effect size equal to .21 ( $z = 7.55, p < .001$ ) in favor of the original study. The null-hypothesis of homogeneous difference in effect sizes was rejected ( $Q(72) = 152.39, p < .001$ ), with medium observed heterogeneity ( $\hat{\tau} = .149, I^2 = 47.8\%$ ). Via Egger's test, precision of the original study was associated with the difference in effect size ( $b = .85, z = 1.88, \text{one-tailed } p = .030$ ), hence imprecise original studies (large standard error) yielded larger differences in effect size between original and replication study. This is confirmed by the funnel plot in Figure S7. Discipline was not associated with the difference in effect size,  $\chi^2(3) = 2.451, p = .48$ , (i.e., the average difference in effect size was equal for JPSP, JEP:LMC, PSCI-soc, and PSCI-cog). Also, after controlling for the effect of the standard error of the original study, no differences between disciplines were observed ( $\chi^2(3) = 2.807, p = .42$ ). No moderating effects were observed for: importance of the effect ( $b = -.010, p = .77$ ), surprising effect ( $b = .001, p = .97$ ), experience and expertise of original team ( $b = -.0017, p = .96$ ), challenge of conducting replication ( $b = 0.026, p = .45$ ), and self-assessed quality of replication ( $b = -.037, p = .51$ ). However, a positive effect of experience and expertise of replication team was observed ( $b = .13, p = .0063$ ), meaning that the difference between original and replication effect size was *higher* for replication teams with more experience and expertise.

The results from the three meta-analyses tentatively suggest that the journals/disciplines are similarly influenced by publication bias leading to overestimated effect sizes, and that cognitive effects are larger than social effects on average -- possibly because of the target of study or the sensitivity of the research designs (e.g., within-subject designs reducing error and increasing sensitivity).

Figure S7: Funnel plot of meta-analysis on difference in effect size (original – replication).



## Moderator Analyses

The main text reports correlations between five reproducibility indicators and aggregate variables of original and replication study characteristics. Below are correlations among the five reproducibility indicators (Table S3), correlations of individual characteristics of original studies with reproducibility indicators (Table S4), and correlations of individual characteristics of replication studies with reproducibility indicators (Table S5).

*Table S2.* Spearman's rank order correlations among reproducibility indicators

|                                                   | Replications<br>p < .05 in<br>original<br>direction | Effect Size<br>Difference | Meta-analytic<br>Estimate | original<br>effect size<br>within<br>replication<br>95% CI | subjective<br>"yes" to "Did<br>it replicate?" |
|---------------------------------------------------|-----------------------------------------------------|---------------------------|---------------------------|------------------------------------------------------------|-----------------------------------------------|
| Replications p < .05 in<br>original direction     | .                                                   |                           |                           |                                                            |                                               |
| Effect Size Difference                            | -0.619                                              | .                         |                           |                                                            |                                               |
| Meta-analytic Estimate                            | 0.592                                               | -0.218                    | .                         |                                                            |                                               |
| original effect size within<br>replication 95% CI | 0.551                                               | -0.498                    | 0.515                     | .                                                          |                                               |
| subjective "yes" to "Did it<br>replicate?"        | 0.956                                               | -0.577                    | 0.565                     | 0.606                                                      | .                                             |

Notes: Effect size difference (original - replication) computed after converting  $r$ 's to Fisher's  $z$ . Notes: Four original results had  $p$ -values slightly higher than .05, but were considered positive results in the original article and are treated that way here. Exclusions (see SI [A3] for explanation): "replications p < .05" (3 excluded;  $n = 97$ ), "effect size difference" (3 excluded;  $n = 97$ ); "meta-analytic mean estimates" (27 excluded;  $n = 73$ ); and, "% original effect size within replication 95% CI" (5 excluded,  $n=95$ ).

**Table S3.** Descriptive statistics and spearman's rank-order correlations of reproducibility indicators with individual original study characteristics

|                                                | M      | SD     | Median | Range             | Replications<br>p < .05 in<br>original<br>direction | Effect Size<br>Difference | Meta-<br>analytic<br>Estimate | original<br>effect size<br>within<br>replication<br>95% CI | subjective<br>"yes" to<br>"Did it<br>replicate?" |
|------------------------------------------------|--------|--------|--------|-------------------|-----------------------------------------------------|---------------------------|-------------------------------|------------------------------------------------------------|--------------------------------------------------|
| Original<br>effect size                        | 0.3942 | 0.2158 | 0.3733 | .0046 to<br>.8596 | 0.304                                               | 0.279                     | 0.793                         | 0.121                                                      | 0.277                                            |
| Original p-<br>value                           | 0.0283 | 0.1309 | 0.0069 | 0 to<br>.912      | -0.327                                              | -0.057                    | -0.468                        | 0.032                                                      | -0.260                                           |
| Original df/N                                  | 2409   | 22994  | 55     | 7 to<br>230025    | -0.150                                              | -0.194                    | -0.502                        | -0.221                                                     | -0.185                                           |
| Institution<br>prestige of<br>1st author       | 3.78   | 1.49   | 3.45   | 1.28 to<br>6.74   | -0.026                                              | 0.012                     | -0.059                        | -0.132                                                     | -0.002                                           |
| Institution<br>prestige of<br>senior<br>author | 3.97   | 1.54   | 3.65   | 1.28 to<br>6.74   | -0.057                                              | -0.062                    | 0.019                         | -0.104                                                     | -0.019                                           |
| Citation<br>impact of<br>1st author            | 3074   | 5341   | 1539   | 54 to<br>44032    | 0.117                                               | -0.111                    | 0.090                         | 0.004                                                      | 0.117                                            |
| Citation<br>impact of<br>senior<br>author      | 13656  | 17220  | 8475   | 240 to<br>86172   | -0.093                                              | -0.060                    | -0.189                        | -0.054                                                     | -0.092                                           |
| Article<br>citation<br>impact                  | 84.91  | 72.95  | 56     | 6 to 341          | -0.013                                              | -0.059                    | -0.172                        | -0.081                                                     | 0.016                                            |
| Internal<br>conceptual<br>replications         | 0.91   | 1.21   | 0      | 0 to 5            | -0.164                                              | 0.036                     | -0.185                        | -0.058                                                     | -0.191                                           |
| Internal<br>direct<br>replications             | 0.06   | 0.32   | 0      | 0 to 3            | 0.061                                               | 0.023                     | 0.071                         | 0.116                                                      | 0.047                                            |
| Surprising<br>original<br>result               | 3.07   | 0.87   | 3      | 1.33 to<br>5.33   | -0.244                                              | 0.102                     | -0.181                        | -0.113                                                     | -0.241                                           |

|                                     |      |      |      |              |        |       |        |        |        |
|-------------------------------------|------|------|------|--------------|--------|-------|--------|--------|--------|
| Importance<br>of original<br>result | 3.36 | 0.71 | 3.28 | 1 to<br>5.33 | -0.105 | 0.038 | -0.205 | -0.133 | -0.074 |
|-------------------------------------|------|------|------|--------------|--------|-------|--------|--------|--------|

Notes: Effect size difference computed after converting r's to Fisher's z. df/N refers to the information on which the test of the effect was based (e.g., df of t-test, denominator df of F-test, sample size - 3 of correlation, and sample size for z and chi2). Four original results had p-values slightly higher than .05, but were considered positive results in the original article and are treated that way here. Exclusions (see SI [A3] for explanation): "replications  $p < .05$ " (3 original nulls excluded;  $n = 97$ ), "effect size difference" (3 excluded;  $n = 97$ ); "meta-analytic mean estimates" (27 excluded;  $n = 73$ ); and, "% original effect size within replication 95% CI" (5 excluded,  $n=95$ ).

*Table S4.* Descriptive statistics and spearman's rank-order correlations of reproducibility indicators with individual replication study characteristics

|                                                        | M     | SD    | Median | Range           | Replications<br>p < .05 in<br>original<br>direction | Effect Size<br>Difference | Meta-<br>analytic<br>Estimate | original<br>effect size<br>within<br>replication<br>95% CI | subjective<br>"yes" to<br>"Did it<br>replicate?" |
|--------------------------------------------------------|-------|-------|--------|-----------------|-----------------------------------------------------|---------------------------|-------------------------------|------------------------------------------------------------|--------------------------------------------------|
| Institution<br>prestige of 1st<br>author               | 3.04  | 1.42  | 2.53   | 1.31 to<br>6.74 | -0.224                                              | 0.114                     | -0.436                        | -0.267                                                     | -0.243                                           |
| Institution<br>prestige of<br>senior author            | 3.03  | 1.4   | 2.61   | 1.31 to<br>6.74 | -0.231                                              | 0.092                     | -0.423                        | -0.307                                                     | -0.249                                           |
| Citation count of<br>1st author                        | 570   | 1280  | 91     | 0 to<br>6853    | 0.064                                               | -0.114                    | -0.045                        | 0.220                                                      | 0.058                                            |
| Citation count of<br>senior author                     | 1443  | 2573  | 377    | 0 to<br>15770   | -0.078                                              | 0.104                     | -0.070                        | 0.038                                                      | -0.067                                           |
| Position of<br>senior member<br>of replication<br>team | 2.91  | 1.89  | 2      | 1 to 7          | -0.157                                              | 0.087                     | -0.241                        | -0.195                                                     | -0.159                                           |
| Highest degree<br>of senior<br>member                  | 1.24  | 0.62  | 1      | 1 to 4          | -0.034                                              | -0.029                    | -0.040                        | -0.155                                                     | -0.025                                           |
| Senior member's<br>total publications                  | 44.81 | 69.01 | 18     | 0 to 400        | -0.021                                              | 0.079                     | 0.037                         | 0.054                                                      | -0.004                                           |
| Domain<br>expertise                                    | 3.22  | 1.07  | 3      | 1 to 5          | 0.042                                               | 0.022                     | 0.130                         | 0.180                                                      | 0.101                                            |
| Method<br>expertise                                    | 3.43  | 1.08  | 3      | 1 to 5          | -0.057                                              | 0.151                     | 0.214                         | 0.009                                                      | -0.026                                           |
| Perceived<br>expertise<br>required                     | 2.25  | 1.2   | 2      | 1 to 5          | -0.114                                              | 0.042                     | -0.054                        | -0.077                                                     | -0.044                                           |
| Perceived<br>opportunity for<br>expectancy bias        | 1.74  | 0.8   | 2      | 1 to 4          | -0.214                                              | 0.117                     | -0.355                        | -0.109                                                     | -0.172                                           |
| Perceived<br>opportunity for                           | 2.21  | 1.02  | 2      | 1 to 5          | -0.194                                              | 0.086                     | -0.333                        | -0.037                                                     | -0.149                                           |

|                                                      |       |       |      |             |        |        |        |        |        |
|------------------------------------------------------|-------|-------|------|-------------|--------|--------|--------|--------|--------|
| impact of lack of diligence                          |       |       |      |             |        |        |        |        |        |
| Implementation quality                               | 3.85  | 0.86  | 4    | 1 to 6      | -0.058 | 0.093  | -0.115 | 0.043  | -0.023 |
| Data collection quality                              | 3.60  | 1.00  | 4    | 1 to 6      | -0.103 | 0.038  | 0.230  | 0.026  | -0.106 |
| Replication similarity                               | 5.72  | 1.05  | 6    | 3 to 7      | 0.015  | -0.075 | -0.005 | -0.036 | 0.044  |
| Difficulty of implementation                         | 4.06  | 1.44  | 4    | 1 to 6      | -0.072 | 0.000  | -0.059 | -0.116 | -0.073 |
| Replication df/N                                     | 4804  | 4574  | 68.5 | 7 to 455304 | -0.085 | -0.224 | -0.692 | -0.257 | -0.164 |
| Replication power                                    | 0.921 | 0.086 | 0.95 | .56 to .99  | 0.368  | -0.053 | 0.142  | -0.056 | 0.285  |
| Replication team surprised by outcome of replication | 2.51  | 1.07  | 2    | 1 to 5      | -0.468 | 0.344  | -0.323 | -0.362 | -0.498 |

Notes: Effect size difference computed after converting r's to Fisher's z. df/N refers to the information on which the test of the effect was based (e.g., df of t-test, denominator df of F-test, sample size - 3 of correlation, and sample size for z and chi2). Four original results had p-values slightly higher than .05, but were considered positive results in the original article and are treated that way here. Exclusions (see SI [A3] for explanation): "replications  $p < .05$ " (3 original nulls excluded;  $n = 97$ ), "effect size difference" (3 excluded;  $n = 97$ ); "meta-analytic mean estimates" (27 excluded;  $n = 73$ ); and, "% original effect size within replication 95% CI" (5 excluded,  $n=95$ ).