

Errata

On Coppock and Green (2022), *The Journal of Politics* 84(2)

This note lists nine misstated claims in “Do Belief Systems Exhibit Dynamic Constraint?”, published in *The Journal of Politics* 84(2), 725–738, doi [10.1086/716294](https://doi.org/10.1086/716294). **None of them changes a conclusion of the article.** The finding that treatments which move a target attitude leave nontarget attitudes untouched is unaffected: the pooled indirect effect is small and statistically indistinguishable from zero on either the published figure or the corrected one, and every estimate in tables 1 through 4 reproduces exactly.

The first three concern the same bootstrap, whose resample count and random seed the deposited code never fixes. The next two attach real estimates to the wrong sample. The last four are counts and roundings that disagree with the article’s own tables.

Every corrected value below is computed from the deposited data at the time this document is rendered; bold marks the tokens that change. Quantities the article states at whole percentage points and which come out of the unseeded bootstrap are left at that precision, because rendering a draw to one decimal would be precise about a number the procedure does not pin down. The code and output supporting each correction are at github.com/active-maintenance-aec/coppock_green_2022.

1. The sign of the pooled ratio

If we pool the three studies together, our estimate is 1.4% (95% CI: -8%, 6%).

Corrected:

If we pool the three studies together, our estimate is **-1.7%** (95% CI: -8%, 6%).

The three inputs are the sample-specific ratios of the precision-weighted average non-target effect to the precision-weighted average target effect: 0.8% on MTurk, 6.4% among elites and -9.4% on Lucid. The deposited `in_text_numbers.R` pools them with `rmeta::meta.summaries`, which passes no `method` argument and is therefore an inverse-variance fixed-effect meta-analysis. Two of the three inputs are near zero and the third is negative with the second-smallest standard error, so no pooling at these weights can land above zero; running the deposited code today gives -2.8%. The published interval endpoints are left as printed because they fall inside the range the unseeded bootstrap

produces, which the point estimate does not: over repeated runs at the resample count the deposit ships, the pooled estimate lies between -5.1% and 0.4%.

2. The same ratio in the discussion

But in only nine of 144 opportunities did it induce significant changes on nontarget issues, and the average spillover effect is just 1.4% of the first-stage direct effect.

Corrected:

But in only nine of 144 opportunities did it induce significant changes on nontarget issues, and the average spillover effect is just **-1.7%** of the first-stage direct effect.

The same quantity as item 1, restated. The counts in the sentence are correct.

3. The upper end of the Lucid interval

On Lucid, the fraction is in fact negative, at -9.4%, with a confidence interval from -24% to 15%.

Corrected:

On Lucid, the fraction is in fact negative, at -9.4%, with a confidence interval from -24% to **2%**.

Neither the archive nor the article sets a seed, so no interval endpoint in this paper is a fixed quantity: each is one draw of a bootstrap quantile. Measured by resampling the draws, this endpoint runs from -6.3% to 6.9% at the ten resamples the deposit ships and from 1.1% to 3.6% at the thousand this repository uses. The published 15% lies outside both. The lower endpoint is left as printed: it falls inside its own range, which runs from -27.5% to -12.3% at ten resamples, as do the four endpoints the article gives for the MTurk and elite samples.

4. The sample the inconsistency rates come from

As is well documented (Ahler and Broockman 2018; Sears and Citrin 1982), inconsistency between preferences for smaller government and increased social spending is common. Among the MTurk sample, 30% of the respondents were inconsistent; this figure was 24% on Lucid.

Corrected:

As is well documented (Ahler and Broockman 2018; Sears and Citrin 1982), inconsistency between preferences for smaller government and increased social spending is common. Among the **study 2 original** sample, 30% of the respondents were inconsistent; this figure was 24% **in the replication**.

Both figures come from study 2, and both study 2 samples were recruited on Lucid: the body text says the 1,087 subjects for that experiment “were also recruited via Lucid”, and appendix table C.3 heads its two columns “Lucid Original” and “Lucid Replication”. The indicator that defines inconsistency appears in no MTurk data file in the deposit. Table 4 carries the same mislabel in its row heading “N (MTurk)”; its cells are correct and it is not reprinted here.

5. The sample the correction effects come from

The results are shown in table 4. When induced to become more consistent, people who prefer smaller government prefer less spending (-0.172 scale points on MTurk, -0.153 scale points on Lucid).

Corrected:

The results are shown in table 4. When induced to become more consistent, people who prefer smaller government prefer less spending (-0.172 scale points **in the original Lucid sample**, -0.153 scale points **in the Lucid replication**).

The same mislabel as item 4. Table 4’s top panel is fitted on the study 2 original sample, the same 1,087 Lucid respondents table 2’s top panel uses and appendix table C.3’s first column describes.

6. The number of nontarget questions in study 2

Responses to these nontarget outcomes are positively correlated among each other and with the target outcome of the manipulation (correlations range from 0.3 to 0.7), so it is reasonable to imagine that subjects would apply the religious conviction consideration when answering the three further questions.

Corrected:

Responses to these nontarget outcomes are positively correlated among each other and with the target outcome of the manipulation (correlations range from 0.3 to 0.7), so it is reasonable to imagine that subjects would apply the religious conviction consideration when answering the **four** further questions.

Table 2 reports four nontarget outcomes, and the results paragraph two paragraphs later calls them “the four additional outcome measures”. The design sentence before this one groups them into three topics rather than counting the questions, describing one measure that raises the religious objection in another domain and “the other two outcome measures” covering same-sex marriage and school prayer; school prayer is measured with two questions, not one. The stated correlation range is correct: the ten pairwise correlations in the original sample run from 0.26 to 0.65.

7. The MTurk within-domain correlation

On MTurk, the average correlation between attitudes in the same domain was smaller at 0.29, and the average correlation between attitudes in different domains was very small at 0.06.

Corrected:

On MTurk, the average correlation between attitudes in the same domain was smaller at **0.30**, and the average correlation between attitudes in different domains was very small at 0.06.

The quantity is 0.2966. Table 1 prints .30 for it, so the table rounds it correctly and the sentence does not.

8. The graduate degree shares

A full 96% of the elite sample hold a college degree, compared with 51% of the MTurk sample and 44% of the Lucid sample; 71% of the elites hold a graduate degree, compared with 13% for the MTurk sample and 15% for Lucid.

Corrected:

A full 96% of the elite sample hold a college degree, compared with 51% of the MTurk sample and 44% of the Lucid sample; **72%** of the elites hold a graduate degree, compared with **12%** for the MTurk sample and 15% for Lucid.

Appendix table B.2 prints 71.6% and 12.4% for the same two categories, so the sentence truncates the first and rounds the second the wrong way. The three college-degree shares, which appear in no table, are 96.5%, 51.4% and 44.5% and hold as printed.

9. A standard error in study 3

By contrast, the effects on economic system justification are small and close to zero: -0.18 standard units in the original study (SE: 0.10) and -0.07 units in the replication (SE: 0.06).

Corrected:

By contrast, the effects on economic system justification are small and close to zero: -0.18 standard units in the original study (SE: 0.10) and -0.07 units in the replication (SE: **0.07**).

Table 3 prints 0.066 for the same standard error. As in item 8, the sentence truncates where the table rounds.