

Errata

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

This note lists four errors in “What Have We Learned about Gender from Candidate Choice Experiments? A Meta-Analysis of Sixty-Seven Factorial Survey Experiments,” published in *The Journal of Politics* 84(2), doi [10.1086/716290](https://doi.org/10.1086/716290). **None of them changes a conclusion of the paper.** They are recorded because the numbers themselves are wrong and a reader who quoted them would be misled.

Every corrected value below is computed from the deposited replication archive ([10.7910/DVN/R27ULT](https://doi.org/10.7910/DVN/R27ULT)) at the time this document is rendered. The maintained code is at github.com/active-maintenance-aec/schwarz_coppock_2022.

1. The share of studies with a positive effect is overstated, in two places

Published, p. 659:

Across the 67 studies we summarize, **three-fourths** find a positive effect of being described as a woman.

Corrected:

Across the 67 studies we summarize, **70%** find a positive effect of being described as a woman.

Published, p. 664:

We observe considerable study-to-study variation, although **more than three-fourths** of the studies show a positive effect.

Corrected:

We observe considerable study-to-study variation, although **70%** of the studies show a positive effect.

Of the 67 study-level estimates, 47 are positive, 19 are negative and 1 is exactly zero. The positive share is 70.1 per cent, which is below three-fourths rather than above it. The claim the two sentences make is unaffected: positive effects are the large majority.

2. The count of negative estimates omits two studies

Published, p. 662:

We see a very consistent pattern in favor of women candidates on average: 47 estimates are positive (23 significant) and **18** are negative (11 significant).

Corrected:

We see a very consistent pattern in favor of women candidates on average: 47 estimates are positive (23 significant) and **19** are negative (11 significant).

The deposited estimates divide into 47 positive, 19 negative and 1 exactly zero, which is the only division summing to 67. The published 47 plus 18 sums to 65.

The two parenthetical counts of significant estimates are left as published, because no corrected value for them can be stated. They correspond to no significance rule the deposited estimates support, and not merely to a rule with a different threshold: both counts rise as a threshold is relaxed, so no threshold can produce 23 and 11 at once. Tightening the threshold until only 23 positive estimates remain significant leaves 5 negative ones; loosening it until 11 negative estimates are significant leaves 43 positive ones. The confidence interval rule the deposited script itself applies gives 29 and 6, and a test of p below 0.05 gives 19 and 2 and cannot be what was done in any case, because 19 of the 67 studies carry no p -value. The repository records this as unresolved rather than corrected.

3. Appendix Figure A.18 prints none of its estimates

Figure A.18, Costa (2020), USA - Lucid, is the one study in the appendix whose randomized attribute levels are whole tweets. At the width the appendix uses, those row labels push the plotting panel, the points, the confidence intervals, the estimate labels and the horizontal axis off the right edge of the page. What is printed is a column of truncated sentences, an axis title cut off after four letters, and no data. None of the 10 conditional effects the figure is drawn from appears on the page.

Nothing is wrong with the estimates themselves; they are in the deposited archive and are unchanged. The figure as it should have appeared:

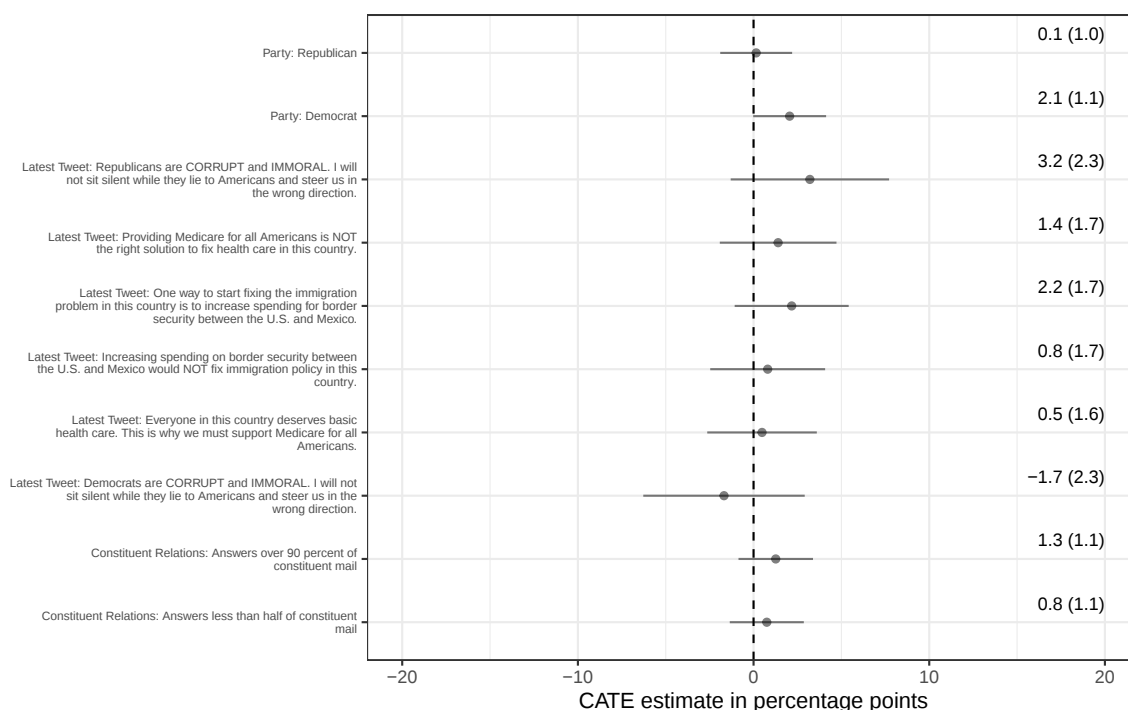


Figure A.18 (redrawn)

Note: Figure A.18 redrawn from the deposited replication archive by `maintained/figure_a2_a45_cate_by_study.R`, with the row labels wrapped and the canvas widened so the panel fits the page. Every estimate, standard error and interval is unchanged; only the layout differs. The 43 other per-study appendix figures print their estimates correctly, and all 1888 of those labels reproduce.

4. Figure 3 covers 37 experiments, not 36

Published, p. 663:

Figure 3 shows the CATEs of candidate gender, conditional on respondent gender, for the **36** experiments where we were able to identify respondent gender.

Corrected:

Figure 3 shows the CATEs of candidate gender, conditional on respondent gender, for the **37** experiments where we were able to identify respondent gender.

The deposited respondent-gender file carries 37 studies, and the published figure itself draws 37 rows in each of its two panels. Figure 3 is correct as printed and is not reproduced here; only the sentence describing it miscounts. Nothing plotted is affected and neither pooled estimate moves.