

Errata

On Coppock, Hill and Vavreck (2020), *Science Advances* 6(36)

This note lists eight errors in “The small effects of political advertising are small regardless of context, message, sender, or receiver: Evidence from 59 real-time randomized experiments,” published in *Science Advances* 6(36), doi [10.1126/sciadv.abc4046](https://doi.org/10.1126/sciadv.abc4046). **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.

Items 1 through 4 are in the main text, items 5 through 7 in section A of the supplementary materials, and item 8 in Materials and Methods. Every corrected value below is computed from the deposited replication archive ([10.7910/DVN/TN7KWR](https://doi.org/10.7910/DVN/TN7KWR)) at the time this document is rendered. The maintained code is at github.com/active-maintenance-aec/coppock_hill_vavreck_2020.

1. The CATE homogeneity p-value for vote choice is misstated by a factor of ten

Published, Results:

For the set of conditional average treatment effects (CATEs), P values are 0.0002 and **0.96**.

Corrected:

For the set of conditional average treatment effects (CATEs), P values are 0.0002 and **0.0961**.

Table S2 of the same article reports 0.0961 for this test, which is what the deposited code produces. The test still fails to reject the null of homogeneity at conventional levels, so the sentence two paragraphs later, that homogeneity is not rejected in three of four opportunities, is unaffected.

2. The general election coefficient disagrees with Table 1

Published, Results:

On average, general election advertisements move candidate favorability by **0.121 scale points** more than primary advertisements, although the estimate is not precise enough to achieve statistical significance.

Corrected:

On average, general election advertisements move candidate favorability by **0.123 scale points** more than primary advertisements, although the estimate is not precise enough to achieve statistical significance.

Table 1 of the same article prints 0.123 in this cell.

Table 1 also prints the model 1 battleground coefficient as -0.00 , where the deposited code gives -0.008 ; the facing text quotes its magnitude correctly as 0.008. The table is reprinted below with that cell at the three decimals the rest of the column carries. Every other cell is unchanged.

Meta-analysis of average treatment effects of advertisements on target candidate favorability and vote choice.

	Candidate favorability (1)	Candidate favorability (2)	Vote choice (3)	Vote choice (4)
Average effect	0.056* (0.020)	0.062* (0.020)	0.007 (0.007)	0.008 (0.007)
Democratic respondent (versus Republican)	0.035 (0.035)	0.022 (0.036)	0.011 (0.010)	0.006 (0.011)
Independent respondent (versus Republican)	0.023 (0.051)	0.015 (0.052)	0.009 (0.020)	0.007 (0.020)
Battleground state (versus non-battleground)	-0.008 (0.033)	-0.007 (0.033)	-0.017 (0.010)	-0.017 (0.010)
PAC sponsor (versus campaign sponsor)	-0.012 (0.043)	0.026 (0.047)	-0.023 (0.013)	-0.016 (0.014)
Time (scaled in months)	-0.023 (0.014)	0.005 (0.010)	-0.009* (0.004)	-0.008 (0.004)
Attack advertisement (versus promotional advertisement)	-0.017 (0.046)		0.028 (0.016)	
General election (versus primary election)	0.123 (0.067)			
Pro-Trump advertisement (versus pro-Clinton advertisement)		-0.124 (0.101)		-0.016 (0.034)
Anti-Clinton advertisement (versus pro-Clinton advertisement)		-0.105 (0.070)		0.012 (0.023)
Anti-Trump advertisement (versus pro-Clinton advertisement)		-0.041 (0.058)		0.026 (0.021)
Pro-Sanders advertisement (versus pro-Clinton advertisement)		-0.075 (0.089)		
Pro-Cruz advertisement (versus pro-Clinton advertisement)		0.047 (0.116)		
Pro-Kasich advertisement (versus pro-Clinton advertisement)		-0.182 (0.145)		
Number of observations	354	354	204	204

* $p < 0.05$

Note: Table 1 as published, with the model 1 battleground coefficient printed to three decimals. Standard errors in parentheses; * denotes $p < 0.05$. Estimated from the deposited replication archive by `maintained/table_1_meta_regression.R`.

3. No estimate in the study exceeds half a scale point

Published, Results:

...one large positive result toward the end of the campaign in October (**more than +0.5 points**), and two other negative results in October.

Corrected:

...one large positive result toward the end of the campaign in October (**more than +0.4 points**), and two other negative results in October.

The largest positive weekly estimate anywhere in the study is 0.432, on the advertisement “Example” in the week of 31 October. The passage’s point, that a significance filter applied to these experiments would leave a handful of results that look large, is unaffected.

4. Standard errors do not exceed the coefficients in every partisanship cell

Published, Results:

Turning first to features of receiver, we see that the differences in treatment response across Democrats, Independents, and Republicans are small, **with SEs greater than coefficients**.

Corrected:

Turning first to features of receiver, we see that the differences in treatment response across Democrats, Independents, and Republicans are small, **with SEs greater than coefficients in 6 of the 8 cells**.

The two exceptions are the Democratic respondent coefficients in columns 1 and 3 of Table 1, and the second is visible on the printed page: 0.011 with a standard error of 0.010. Neither coefficient is distinguishable from zero, so the claim the sentence supports, that treatment response does not differ by respondent partisanship, is unaffected.

5. The advertisement counts in section A disagree with the main text

Published, supplementary materials, section A:

Including the ads we deployed more than once, we tested **28** ads that attacked Trump, **nine** attacks on Clinton, **nine** Clinton promotional ads, and three promotional Trump ads.

Corrected:

Including the ads we deployed more than once, we tested **30** ads that attacked Trump, **11** attacks on Clinton, **8** Clinton promotional ads, and three promotional Trump ads.

Counted as the section says, over advertisement-week deployments rather than distinct advertisements, the breakdown is 30, 11, 8 and 3, which is exactly what the main text states. Counted over distinct advertisements it is 24, 11, 6 and 3. The published breakdown matches neither.

6. The number of unique advertisements in section A disagrees with the abstract

Published, supplementary materials, section A:

Table S1 shows the name and target of the **50** unique advertisements and the week they were deployed.

Corrected:

Table S1 shows the name and target of the **49** unique advertisements and the week they were deployed.

The deposited data hold 49 campaign advertisements plus the placebo, and the abstract's 49 is right. Table S1 itself carries 48 distinct names, because two of the 49 advertisements share the title "Sacrifice" and the week of 5 September.

7. Section A names eight advertisements shown in multiple weeks

Published, supplementary materials, section A:

Some advertisements ("**America**," "**Love/Kindness**," "**Man of People**," "**Mirrors**," "**Quotes**," "**Real Life**," "**Role Model**," "**Speak**") were shown in multiple weeks, enabling us to estimate how the effect of the same advertisement may vary with the changing electoral context.

Corrected:

Some advertisements ("**America**," "**Love/Kindness**," "**Mirrors**," "**Quotes**," "**Real Life**," "**Role Model**," "**Speak**") were shown in multiple weeks, enabling us to estimate how the effect of the same advertisement may vary with the changing electoral context.

"Man of People" was fielded once, in the week of 14 March. Section B and Figure S1 both use the correct count of 7.

8. The manipulation check range is one point narrow at each end

Published, Materials and Methods:

Tiny fractions of the control groups claimed to have seen campaign advertisements (2 to 4%) compared with large majorities of the treatment groups (**92 to 95%**).

Corrected:

Tiny fractions of the control groups claimed to have seen campaign advertisements (2 to 4%) compared with large majorities of the treatment groups (**91 to 96%**).

Across the ten waves in which the manipulation check was asked, the unweighted share of treated respondents reporting a campaign advertisement runs from 91.4 to 95.5 per cent, which is what the deposited script written for this sentence computes. Weighting by the survey weights widens the range further, to 90.2 to 95.9 per cent. The control group range is correct as published.