

Misinformation and Morality:

Encountering Fake-News Headlines Makes Them Seem Less Unethical to Publish and Share

Supplemental Online Material – Reviewed
(SOM-R)

Daniel A. Effron
London Business School

Medha Raj
Marshall School of Business
University of Southern California

Table of Contents

Appendix 1: Sensitivity Analyses for Experiments 1-4	p. 3
Appendix 2: Pilot Experiment	p. 7
Appendix 3: Fake-News Headlines	p. 9
Appendix 4: Experiment 1's Mediation Analysis	p. 10
Appendix 5: Supplemental Analyses for Experiment 3	p. 12
Appendix 6: Experiment 4's Mediation Analysis	p. 16
Appendix 7: Verbatim Text of Materials from Experiment 1	p. 20
Appendix 8: Verbatim Text of Materials from Experiment 2	p. 28
Appendix 9: Verbatim Text of Materials from Experiment 3	p. 30
Appendix 10: Verbatim Text of Materials from Experiment 4	p. 32
References	p. 36

Figures

Figure SOM-1: Experiment 1's Mediation Analysis	p. 11
Figure SOM-2: Experiment 3's Moderated Mediation Analysis	p. 14
Figure SOM-3: Experiment 4's Mediation Analysis of Sharing Intentions	p. 18

Tables

Table SOM-1: Mixed Model Parameters from Pilot Experiment, Used in Sensitivity Analyses	p. 5
Table SOM-2: Smallest Effect-Size Detectable at Approximately 80% Power in Experiments 1, 2, and 4	p. 5
Table SOM-3: Indirect Effects of Prior Exposure on Sharing Intentions in Experiment 4	p. 18

Appendix 1

Sensitivity Analyses for Experiments 1-4

Analytic Approach

To determine our experiments' sensitivity – i.e., the smallest effect size that each experiment could detect at 80% power given the number of participants and 12 repeated measures – we conducted Monte Carlo simulations, implemented with the *powerSim* command in R (Green & MacLeod, 2015). For this command, the user inputs the parameters of a mixed regression model, specifying the size of the hypothesized effect as an unstandardized regression coefficient. Then *powerSim* tests the significance of the hypothesized effect in 1,000 simulated studies, and outputs statistical power, computed as the percentage of the simulated studies in which the hypothesis test was significant.

Although *powerSim* computes power given an effect size, we were able to use it to infer the smallest effect size detectable with given power (i.e., a sensitivity analysis). Specifically, we plugged in different effect sizes in search of the smallest one that could be detected at approximately 80% of the time. This trial-and-error method makes it difficult to find the effect size that yields *exactly* 80% power, so we considered power “approximately 80%” if it was within $\pm 1.5\%$ of 80. We report the exact power with 95% *CI* in the Results section below.

All experiments tested their main hypothesis using mixed regression models, with random intercepts for participants and fixed effects for the specific headlines rated. The sensitivity analyses required that we estimate this model's fixed intercept, fixed effects for headlines, random intercept variance, and residual variance. We made these estimates based on a large pilot study (7,075 observations from 594 people) that we conducted before any of the

experiments in the main text (see SOM-R Appendix 2, below). Table SOM-1 displays the relevant parameters.

For each study, we set the number of participants and observations to the number in our final samples after data exclusions (see main text and Results below). We tested significance with the z -distribution, setting the α level to .05, two-tailed. Although Experiments 2 and 3 pre-registered one-tailed tests, *powerSim* only allows two-tailed tests, so the estimated sensitivity for these two experiments is conservative – that is, they can detect smaller effects using one-tailed than two-tailed tests.

Results

Experiments 1, 2, and 4. Experiments 1, 2, and 4 tested the hypothesis that previously-seen headlines would be rated as less unethical on average than new headlines. Table SOM-2 displays the results of the sensitivity analysis for these three studies alongside the number of participants and observations. Each experiment had approximately 80% power to detect a small mean difference on the 100-point measure of moral condemnation. For example, the study with the smallest sample size, Experiment 1, could detect a mean difference of 2.15.

We converted these mean differences to a standardized effect size, Cohen's d_z , which is equal to the mean difference divided by its SD . We used an estimate of this SD from the pilot study mentioned above: 9.55. The results showed that each experiment could detect a small effect, e.g., $|.23|$ SD s in Experiment 1 (see Table SOM-2).

Experiment 3. The hypothesized effect in Experiment 3 was an interaction. Specifically, we expected to replicate the mean difference in moral condemnation between repeated and new headlines in the intuitive-thinking condition, and attenuate or eliminate this mean difference in the deliberative-thinking condition. To calculate the sensitivity to detect this interaction with

Experiment 3's 8,731 observations from 761 people, we first had to estimate the effect size that would be observed in the intuitive-thinking condition. We chose the effect size from Experiment 1, a mean difference of $b = -3.83$, $d_z = -.26$, because Experiment 1 was our only previous study that exposed participants to repeated headlines four times (Experiment 2 and the Pilot used only one exposure).

The sensitivity analysis also required that we estimate the fixed effect for thinking condition in the mixed regression model. Based on the dummy coding we used (see main text), this represents how much deliberative thinking increased or decreased unethicity judgments for new headlines. We had little basis for estimating this effect *a priori*. Because we conducted these analyses after we had collected Experiment 3's data, and because this effect is not relevant to the hypotheses, we simply input the effect size that Experiment 3 actually produced, $b = 3.33$. We emphasize that all other estimates were made based on data collected before Experiment 3 was run.

Given these estimates, the sensitivity analysis showed that the smallest interaction coefficient that Experiment 3 could detect at approximately 80% power was $b = 1.92$ (actual power: 79.1%, [76.45, 81.58]). This coefficient indicates that a mean difference of $b = -3.83$ scale points in the intuitive-thinking condition would be reduced by half to $b = -1.92$ in the deliberative-thinking condition. Converting these unstandardized effect-sizes to standardized effect sizes using the *SD* of the mean difference observed in Experiment 1 ($SD = 15.71$), the interaction would reduce an effect of repetition from $d_z = -.24$ ¹ in the intuitive-thinking condition to $d_z = -.12$ in the deliberative-thinking condition.

¹ Due to rounding, this is slightly different from the effect size computed in Experiment 1 mentioned above, $d_z = -.26$.

Table SOM-1

Mixed Model Parameters from Pilot Experiment, Used in Sensitivity Analyses

Parameter	Value
Random intercept	384.59
variance	
Residual variance	251.55
Fixed intercept	82.45
Fixed slopes	
<i>Headline 2</i>	−3.86
<i>Headline 3</i>	−13.81
<i>Headline 4</i>	−1.37
<i>Headline 5</i>	−7.33
<i>Headline 6</i>	−8.71
<i>Headline 7</i>	−3.71
<i>Headline 8</i>	−7.40
<i>Headline 9</i>	−8.55
<i>Headline 10</i>	−7.83
<i>Headline 11</i>	−3.14
<i>Headline 12</i>	−2.96

Table SOM-2

Smallest Effect-Size Detectable at Approximately 80% Power in Experiments 1, 2, and 4

Experiment	<i>N</i> participants	<i>N</i> observations	Power	95% CI	Smallest detectable $ M_{dif} $	Smallest detectable $ d_z $
1	138	1,648	80.10%	[77.49, 82.53]	2.15	.23
2	796	9,536	78.50%	[75.82, 81.01]	.92	.10
4	296	3,552	80.30%	[77.70, 82.72]	1.50	.15

Notes. Each experiment had 12 repeated measures (observations). Due to missing data, the number of observations is slightly less than 12 * the number of participants. d_z was computed by dividing the mean difference by the *SD* for the mean difference in a prior pilot study (9.55). M_{dif} = mean difference.

Appendix 2

Experiment 2's Pilot

Before running Experiment 2, we ran a Pilot Experiment with essentially the same procedure: Participants rated the unethicity of publishing 12 fake-news headlines, a random 6 of which they had seen earlier in the study. Consistent with our subsequent studies, the results showed that headlines seen earlier were rated as less unethical to publish than headlines not previously seen. We used these results to inform the sensitivity analyses reported for subsequent studies.

Method

Participants. In May of 2018, we posted slots for 600 American participants on Amazon Mechanical Turk (MTurk) based on our intuition that this would constitute a “large” sample for design with 12 repeated measures. Participants could only begin the study if they consented and correctly answered a reading-comprehension question, and we took precautions to prevent people from with duplicate MTurk IDs, or non-US or duplicate IP addresses from beginning. After applying our lab-standard exclusion criteria (duplicate or non-US IP addresses, duplicate MTurk IDs), 7,075 responses from 594 people remained (341 women and 253 men; $M_{\text{age}} = 33.90$, $SD = 11.76$; 327 people who identified as or leaned Democrat, and 181 who identified as or leaned Republican).

We planned to analyze the data with a mixed regression model. We did not compute the smallest effect size this analysis could detect in our design because doing so would require specifying the model's fixed intercept, fixed effects for specific headlines, random intercept variance, and residual variance. Typically, these quantities are estimated based on pilot data, but

because this was the first study we had run with these stimuli in this design, we did not have informed estimates (see Green & MacLeod, 2015).

Procedure. The procedure, measures, and stimuli were identical to Experiment 2 (see main text), with one exception. For the comprehension check at the end of the Pilot, participants indicated how many headlines they thought were true: *None*, *A few of them*, *Around half of them*, *Most of them*, or *All of them*. Unlike Experiment 2, the Pilot did not follow-up by asking which specific headlines (if any) participants thought were true.

Results and Discussion

As predicted, participants rated previously-seen headlines as less unethical to publish than previously-unseen headlines ($M_s = 75.62$ vs. 76.65 , $SD_s = 25.49$ and 25.35), $b = -1.03$, $SE_b = .38$, $z = 2.74$, $p = .006$, $d_z = -.11$. As a robustness check, we reran the analysis, excluding participants who incorrectly said that any of the headlines were true (19% of the sample). The results remained significant and in the same direction, $b = -.77$, $SE_b = .38$, $z = 2.02$, $p = .043$, $d_z = -.08$. These results provide initial evidence that a single exposure to fake-news headlines make them seem less unethical to publish. Experiment 2 replicated these findings in a pre-registered design.

The Pilot Experiment also yielded useful information for conducting sensitivity analyses in subsequent studies (i.e., in the experiments described in the main text). Table SOM-1 summarizes this information (see p. 6, above).

Appendix 3:

Fake-News Headlines

These stimuli were collected by Pennycook, Cannon, and Rand (2018). Each headline was accompanied by a relevant image and the name of the news source (e.g., “DAILYHEADLINES.COM”). Pennycook et al. have posted verbatim stimuli at <https://osf.io/w3jst/>; we used the 12 files labeled “Fake News/No warning.”

- Election night: Hillary Was Drunk, Got Physical With Mook and Podesta
- Obama Was Going To Castro’s Funeral – Until Trump Told Him This ...
- Donald Trump Sent His Own Plane To Transport 200 Stranded Marines
- Mike Pence: Gay Conversion Therapy Saved My Marriage
- Pennsylvania Federal Court Grants Legal Authority To REMOVE TRUMP After Russian Meddling
- Trump On Revamping the Military: We’re Bringing Back the Draft
- BLM Thug Protests President Trump With Selfie ... Accidentally Shoots Himself in the Face
- NYTimes David Brooks: “Trump Needs To Decide If He Prefers To Resign, Be Impeached, Or Get Assassinated” – US Politics
- Clint Eastwood Refuses to Accept Presidential Medal of Freedom From Obama, Says, “He is not my president” – Usa news
- FBI Director Comey Just Proved His Bias By Putting Trump Sign On His Front Lawn
- Trump to Ban All TV Shows that Promote Gay Activity Starting with Empire as President – The #1 Empowering Conscious Website In The World
- Sarah Palin Calls To Boycott Mall Of America Because “Santa Was Always White in the Bible”

Appendix 4

Experiment 1's Mediation Analysis

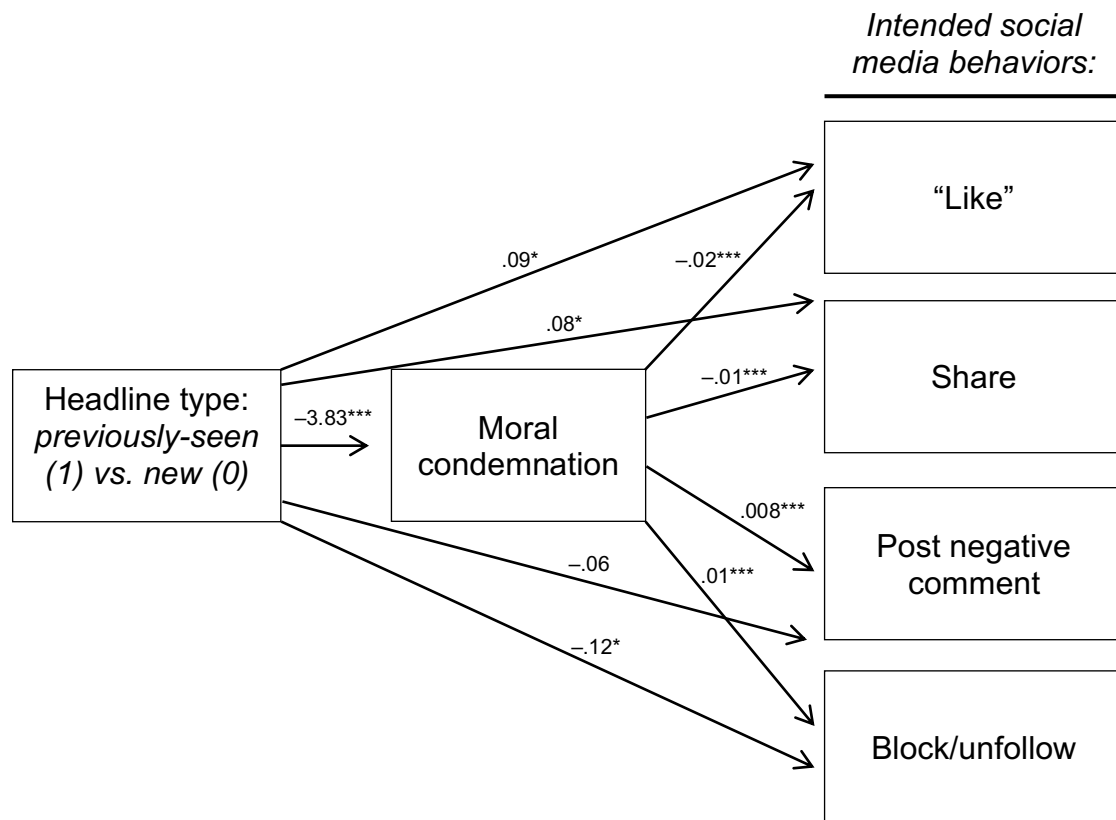
In Experiment 1, we predicted that repeated exposure to a headline, by reducing moral condemnation, would affect how participants intended to respond if an acquaintance of theirs posted the headline on social media (i.e., intentions to “like” the headline, share the headline, respond by positing a negative comment, and block/unfollow the acquaintance). To test this prediction, we computed multi-level mediation models with headline type as the independent variable (1 = *Previously-seen*; 0 = *New*), moral condemnation as the mediator, and the relevant social media measure as the dependent variable, with fixed effects for the specific headlines and random intercepts for participants. We ran the multi-level models with Stata's *gsem* command, and computed indirect effects by multiplying the *a* and the *b* paths together using the *nlcom* command.

The results revealed significant indirect effects on each of the four measures. Specifically, seeing a headline earlier in the study reduced how unethical participants thought it was to publish, which in turn predicted greater willingness to “like” and share it on social media, and less inclination to post a negative comment and to block/unfollow the person who posted it; respectively, $bs = .07, .05, -.03$, and $-.05$, $zs > 3.25$, $ps < .002$ for the indirect effects (see Figure SOM-1 on the next page).

These results are consistent with a causal model in which repeated exposure affects intended media behavior by dampening moral condemnation. However, because mediation analysis is correlational (Fiedler, Schott, & Meiser, 2011) and cannot measure all possible variables (Bullock, Green, & Ha, 2010), the results cannot rule out alternative causal models.

Figure SOM-1

Experiment 1's Mediation Analysis



Appendix 5

Supplemental Analyses of Experiment 3

Cognitive Reflection Test

To examine whether individual differences in deliberative versus intuitive thinking moderated Experiment 3's results, we included the three-item Cognitive Reflection Test (CRT; Frederick, 2005). We pre-registered this measure as exploratory because we thought our manipulation of deliberative versus intuitive thinking could swamp any individual differences in thinking style.

We first examined whether CRT moderated people's tendency to rate previously-seen headlines as less unethical to publish than new headlines. Specifically, we submitted the moral condemnation measure to a multi-level regression model with random intercepts for participants and headline-type (1 = previously-seen, 0 = new), CRT score (i.e., the number of correctly answered CRT questions, ranging from 0 to 4, with higher numbers indicating more deliberative thinking), and their interaction as fixed effects. As controls, the model also included fixed effects for the thinking manipulation (0 = intuitive, 1 = deliberative) and for the specific headlines rated. The results showed no evidence that CRT scores moderated the effect of headline type, $b = -.31$, $SE(b) = .36$, $z = .87$, $p = .386$ for the interaction (cf. De keersmaecker et al., in press).

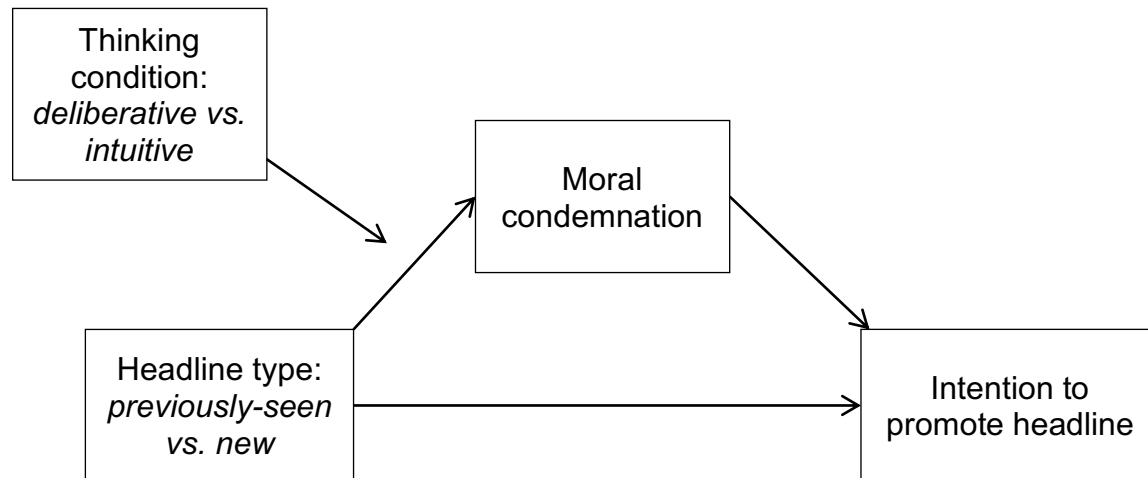
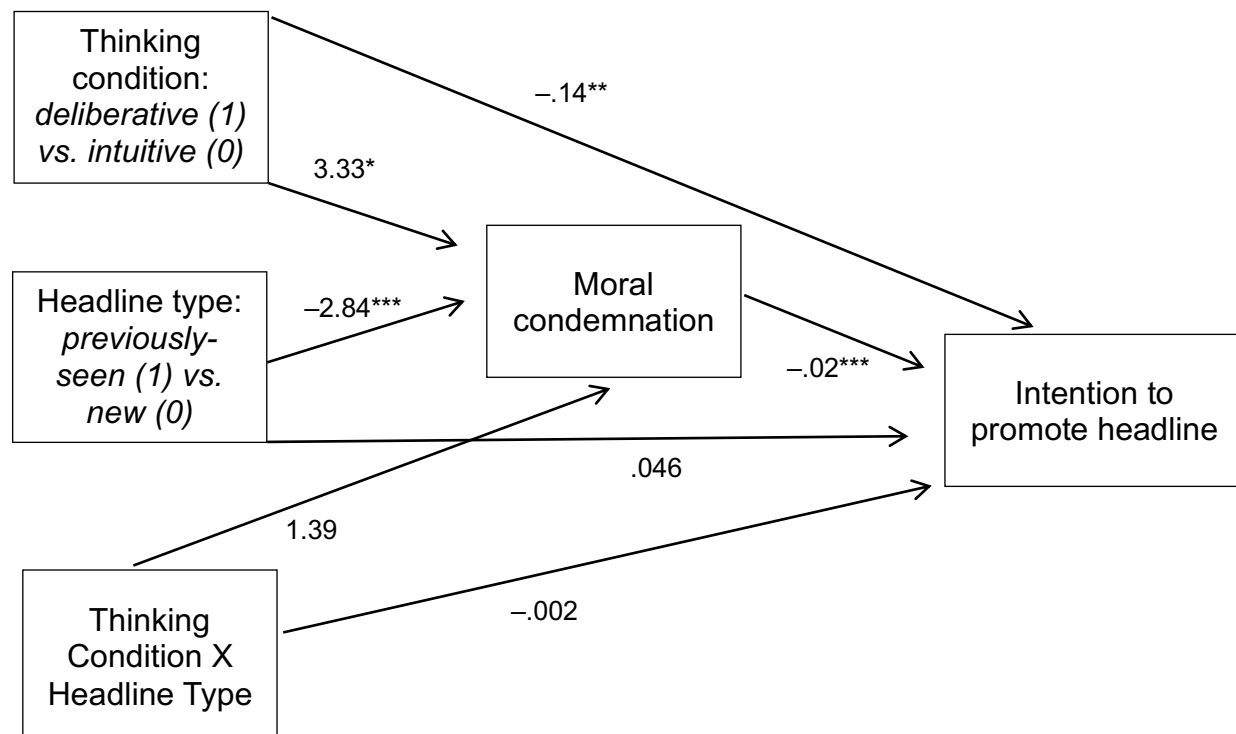
Next, we removed the interaction term from the model to explore whether there was a main effect of CRT on moral condemnation. The results revealed that the more deliberative people were (based on the results of the CRT), the more unethical they thought the headlines were to share, $b = 2.13$, $SE(b) = .44$, $z = 3.20$, $p = .001$ (cf. Pennycook & Rand, 2019).

Moderated Mediation Analysis

In Experiment 1, we observed that exposure to fake-news headlines affected participants' intended behaviors on social media by reducing how unethical participants thought it was to publish those headlines. In Experiment 3, we pre-registered the hypothesis that these indirect/mediation effects would replicate in the intuitive-thinking condition, and that these effects would be significantly weaker or absent in the deliberative-thinking condition. To test this hypothesis, we computed the moderated mediation model shown in Figure SOM-2's top panel (i.e., Model 8 from Hayes, 2013) for each of the three behavioral intentions measures (i.e., sharing, liking, and blocking/unfollowing on social media). The results were nearly identical with the same significance level for each measure, so we averaged all four the measures, reverse-coding "blocking/unfollowing," so that higher numbers indicate greater inclination to promote (vs. express disapproval of) the headlines. Because each participant submitted multiple observations, we specified random intercepts for participant and included fixed effects for the specific headline. (We mistakenly did not pre-register these fixed effects, but the results do not meaningfully change when we remove them from analysis). As a pre-registered robustness check, we repeated this analysis after excluding responses to any headlines a participant misidentified as true.

Figure SOM-2

Experiment 3's Moderated Mediation Analysis

Top Panel: Conceptual Model*Bottom Panel: Results*

Note. For simplicity, the figure shows significance levels based on two-tailed tests, even though we pre-registered some one-tailed tests. * $p < .05$, ** $p < .01$, *** $p < .001$.

In the intuitive-thinking condition, as in Experiment 1, participants rated previously-seen (vs. new) headlines as less unethical to publish, which in turn predicted stronger intentions to promote the headlines on social media – a significant indirect effect, $b = .054$, $SE(b) = .011$, $z = 4.72$, $p < .001$ in the main analysis, and $b = .051$, $SE(b) = .011$, $z = 4.73$, $p < .001$ in the robustness check. In the deliberative-thinking condition, the coefficient for this indirect effect was about half the size in the main analysis, $b = .027$, $SE(b) = .012$, $z = 2.25$, $p < .012$, and even smaller in the robustness check, $b = .017$, $SE(b) = .012$, $z = 1.47$, $p = .071$. This pattern is consistent with our hypothesis, but the difference in indirect effects was only significant in the pre-registered robustness check, $b = .034$, $SE(b) = .016$, $z = 2.15$, $p = .016$, and not in the main analysis, $b = .026$, $SE(b) = .017$, $z = 1.58$, $p = .057$.

Overall, these findings offer mixed support for the prediction that deliberative (vs. intuitive) thinking would attenuate the indirect effect of previously seeing the headlines, via moral condemnation, on social media intentions. Additionally, these results carry the usual caveats of statistical mediation: the possibility of alternative causal models (Fiedler et al., 2011) and the risk of omitted-variable bias (Bullock et al., 2010).

Appendix 6

Experiment 4's Mediation Analysis

We tested a mediation model in which previously-seen (vs. new) headlines received less condemnation, which in turn led to increased intentions to share the headlines if encountered on social media. Specifically, we computed multi-level regression models using Stata's *gsem* command, specifying fixed effects for the specific headlines and random intercepts for participants, and then computed the indirect effect by multiplying the *a* and the *b* paths together using the *nlcom* command. As noted in the main text, the results were consistent with statistical mediation, $b = .05$, $SE(b) = .010$, $z = 4.97$, $p < .001$.

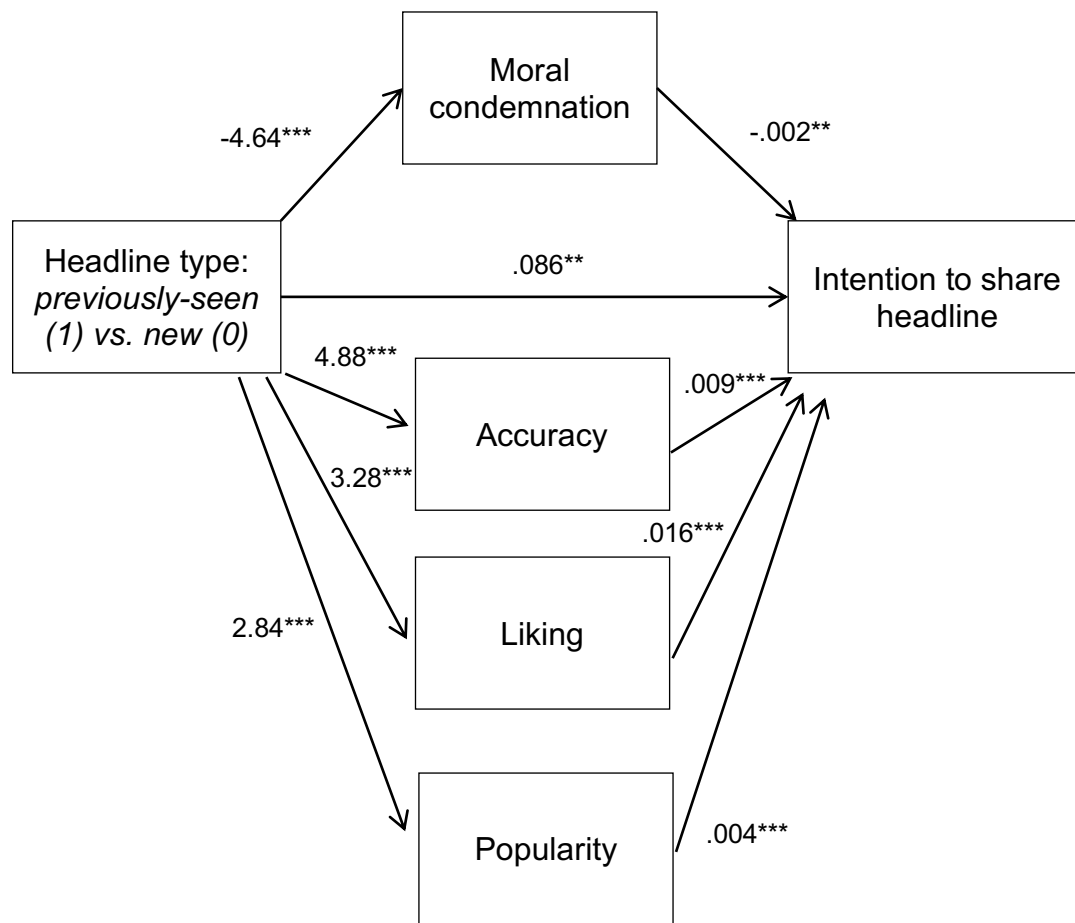
We next tested the parallel-mediator model shown in Figure SOM-3. The results were consistent with a model in which repeated exposure to the headlines increased intentions to share them through four independent routes: by increasing perceptions that the headlines were accurate, by increasing how much participants liked the headlines, by increasing estimates of the number of people who had seen them, and – most importantly for the present research – by diminishing judgments that sharing would be unethical. That is, the indirect effect through all four variables was significant, $ps < .011$ (see Table SOM-3). These results are consistent with the possibility that, above and beyond other variables that are affected by repeated exposure, moral condemnation helps explain why prior encounters with headlines increases intentions to share them.

Given the standard limitations of measurement-of-mediation analysis (e.g., Bullock et al., 2010), the results should be interpreted with caution. For example, given the correlational nature of mediation analysis, other causal models are possible (Fiedler et al., 2011). However, the results do reduce concerns about a standard limitation of measurement-of-mediation designs: that

unmeasured variables correlate with the mediator and dependent variable, thereby biasing estimates of the indirect effect (Bullock et al., 2010). Although it is never possible to prove that all omitted variables have been accounted for, judgments of accuracy, popularity, and liking are key variables that should correlate with moral condemnation and sharing intentions and feature prominently in research on repetition (Dechêne, Stahl, Hansen, & Wänke, 2010; Kwan, Yap, & Chiu, 2015; Weaver, Garcia, Schwarz, & Miller, 2007; Zajonc, 1968).

Figure SOM-3

Experiment 4's Mediation Analysis of Sharing Intentions



Note. See Table SOM-3 for indirect effects. $^{**} p < .005$, $^{***} p < .001$, two-tailed.

Table SOM-3

Indirect Effects of Prior Exposure on Sharing Intentions in Experiment 4

Mediator	<i>b</i>	<i>SE(b)</i>	<i>z</i>	<i>p</i>
Moral condemnation	0.008	0.003	2.54	0.011
Accuracy	0.043	0.008	5.77	0.000
Liking	0.053	0.014	3.78	0.000
Popularity	0.012	0.003	3.37	0.001

Note. Indirect effects were computed with Stata's *gsem* command by multiplying the *a*-path by the *b*-path; see Figure 4. The model includes fixed effects for the specific headlines and random intercepts for participants.

Appendix 7

Verbatim Text of Materials from Experiment 1

[*Note.* The study was programmed in Qualtrics. Text in square brackets was not shown to participants. “-----” indicates a page break]

In today's study, you will be asked to respond to a variety of different questions. It's important that you pay close attention and read all directions carefully. To indicate that you're able to read and understand directions please select the word below that describes something that you use to cook.

- ☐ Rock
- ☐ Car
- ☐ Stove
- ☐ Guitar
- ☐ Shoe
- ☐ Painting
- ☐ Dog
- ☐ Clock

Please enter your Prolific ID here:

[FAMILIARIZATION PHASE]

In today's study, you will complete a series of tasks. Some of these tasks are related, while others are not.

Please make sure to pay attention to each part of the study.

Please click the arrow to proceed.

In the first part of today's study you will be asked to read a series of news headlines that were recently published online.

Please make sure to read the headlines carefully, without any distractions. For each headline you read, you may be asked to answer a series of questions.

When you are ready to begin, please click the arrow below to proceed.

Now, we'll ask you how engaging each headline is.

[Note. The order of “engaging,” “funny,” “well-written,” and “interesting” were randomized]

[Note. One of the 6 headlines randomly selected to appear during Phase 1 appeared here]

How engaging is this headline?

Not at all				Somewhat			Very
1	2	3	4	5	6	7	

[Note. This question was repeated for the remaining 5 headlines that were randomly selected to appear during Phase 1]

Now, we'll ask you how well-written each headline is.

[Note. One of the 6 headlines randomly selected to appear during Phase 1 was displayed here]

How well-written is this headline?

Not at all				Somewhat			Very
1	2	3	4	5	6	7	

[Note. This question was repeated for the remaining 5 headlines that were randomly selected to appear during Phase 1]

Now, we'll ask you how interesting each headline is.

[Note. One of the 6 headlines randomly selected to appear during Phase 1 was displayed here]

How interesting is this headline?

Not at all				Somewhat			Very
1	2	3	4	5	6	7	

[Note. This question was repeated for the remaining 5 headlines that were randomly selected to appear during Phase 1]

Now, we'll ask you how funny each headline is.

[Note. One of the 6 headlines randomly selected to appear during Phase 1 was displayed here]

How funny is this headline?

Not at all				Somewhat			Very
1	2	3	4	5	6	7	

[Note. This question was repeated for the remaining 5 headlines that were randomly selected to appear during Phase 1]

[DISTRACTOR TASK]

We would now like you to answer a few questions.

Please click the arrow below to proceed.

How old are you, in years?

What is your gender?

- ☐ Male
- ☐ Female
- ☐ Other

What is the highest level of school you have completed or the highest degree you have received?

- ☐ Some high school or less
- ☐ High school diploma
- ☐ Some college (1 yr. to less than 4 yrs.)
- ☐ Two-year college degree (A.A.)
- ☐ Four-year college degree (B.A. or B.S.)
- ☐ MA/PhD, MD, MBA, Law Degree

How proficient are you in English?

Not at all		Somewhat		Very proficient
1	2	3	4	5

On social issues I am:

Strongly liberal	Somewhat liberal	Moderate	Somewhat conservative	Strongly conservative
------------------	------------------	----------	-----------------------	-----------------------

On economic issues I am:

Strongly liberal	Somewhat liberal	Moderate	Somewhat conservative	Strongly conservative
------------------	------------------	----------	-----------------------	-----------------------

In politics today, do you consider yourself:

- ☐ A Democrat
- ☐ A Republican
- ☐ An Independent
- ☐ None of the above

[Note. Displayed only if “An Independent” of “None of the above” was selected]

You said you do not consider yourself either a Democrat or a Republican. Do you lean towards one party or the other?

- ☐ I lean Democrat
- ☐ I lean Republican
- ☐ I do not lean towards either party

Did you vote in the most recent Presidential election?

- ☐ Yes
- ☐ No

[Note. Displayed if answer to previous question was “yes”]

Whom did you vote for?

- ☐ Donald Trump
- ☐ Hillary Clinton
- ☐ None of the Above

[Note. Displayed if participants indicated they did not vote]

If you had voted, whom would you have voted for?

- ☐ Donald Trump
- ☐ Hillary Clinton
- ☐ None of the Above

If you absolutely had to choose between only Clinton and Trump, who would you have preferred to be the President of the United States?

- ☐ Hillary Clinton
- ☐ Donald Trump

[JUDGMENT PHASE]

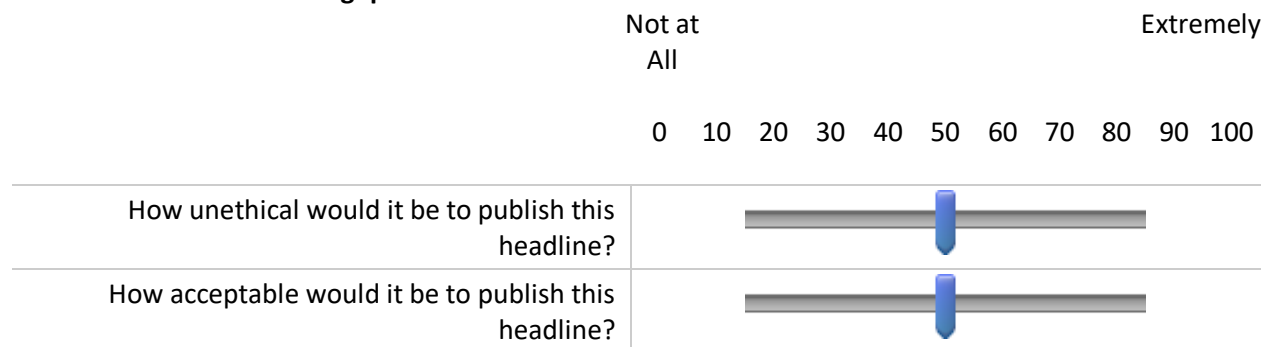
Now, you will move on to the final part of today's study. For this part of the study you will be asked to read a series of fake news headlines that were recently published online. The information in these

headline is not real. Non-partisan fact-checking websites have confirmed that these headlines describe events that did NOT happen. After reading each headline, you will be asked to answer a few short questions.

Please click the arrow below to proceed.

[Note. A randomly selected headline from our full bank of 12 appeared here.]

Please answer the following questions.



If one of your acquaintances shared this headline on social media...

[illegible]

To the best of your knowledge, how accurate is the claim in the above headline?

Not at all accurate	Not very accurate	Somewhat accurate	Very accurate
1	2	3	4

[Note. These questions appeared again for all 12 headlines, presented in randomized orders]

You just rated how justified, unethical, etc. it would be to publish each of 12 headlines.

Which statement is correct about these 12 headlines?

- ☐ All the headlines were FALSE
- ☐ All the headlines were TRUE
- ☐ Some of the headlines were FALSE and some were TRUE

[Note. If participants indicated that some were true and some were false, the following questions were displayed]

You indicated that some of the headlines were true and some were false. Now we'll ask you which ones you think are true versus false.

[Note. One of the 12 headlines was randomly selected to appear here]

Is this headline:

- ☐ TRUE
- ☐ FALSE

[Note. This question repeated for all 12 headlines, shown in randomized orders]

Did you notice any errors or typos when taking this study, or experience any technical problems? If so, please explain.

Thank you for your participation in today's study.

If you have any comments about this study, please enter them here.

Appendix 8

Verbatim Text of Materials from Experiment 2

Experiment 2's materials were identical to Experiment 1's, with the following three exceptions:

- 1) Replace the text in Appendix 2 from "For each headline you read, you may be asked to answer a series of questions." to "How funny is this headline?" with the following:

For each headline you will be asked to indicate how interesting you think the headline is.

When you are ready to begin, please click the arrow below to proceed.

[Note. One of the 6 headlines randomly selected to appear during Phase 1 was displayed here]

How interesting is this headline?

Not at all interesting

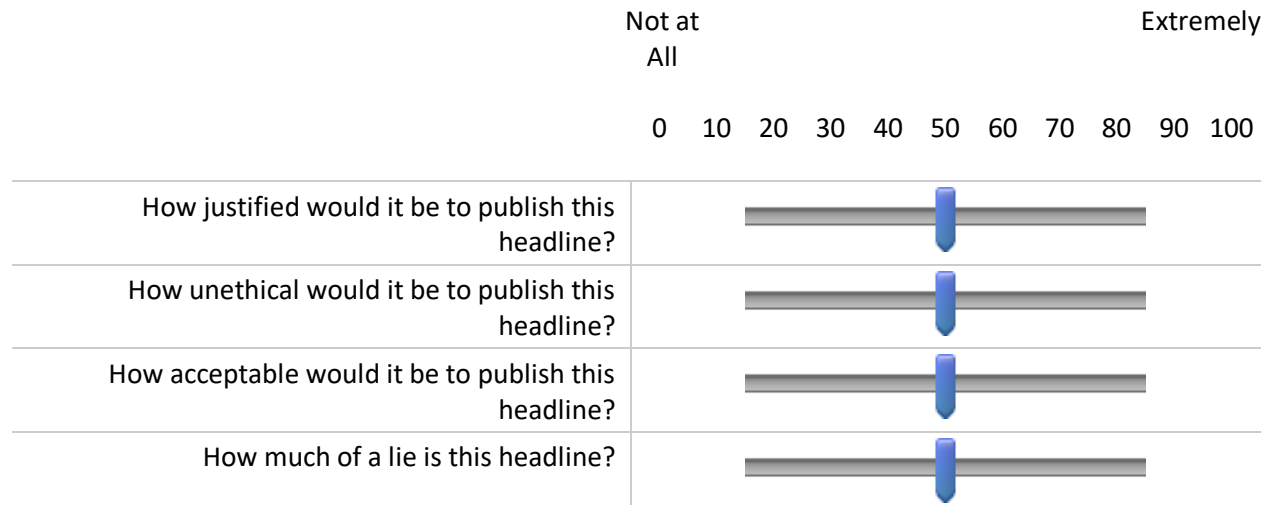
Somewhat interesting

Very interesting

[Note. This question was repeated for the remaining 5 headlines that were randomly selected to appear during Phase 1]

2) Replace the two-item measure of moral condemnation with the following measure:

Please answer the following questions.



3) Omit the questions asking “If one of your acquaintances shared this headline on social media...” and “To the best of your knowledge, how accurate is the claim in the above headline?”

Appendix 9

Verbatim Text of Materials for Experiment 3

Experiment 4's materials were identical to Experiment 1, with the following exceptions:

1) In Phase 2, insert the following text after "Please click the arrow below to proceed," just before the question about how unethical and acceptable each headline would be to publish:

*[Text displayed in the **deliberative-thinking condition** only:]*

Research shows that surveys are most informative when respondents like you **take their time** and **think carefully** about each survey question before answering.

We will now ask you to rate how unethical it would be to publish some headlines.

Before giving your ratings, you should **take time to deliberate**. **Think very hard** about the headlines, try to **ignore any gut feelings**, and **generate clear reasons** for your judgment about how unethical it is.

These reasons could focus on the **specific information** in the headline, the **principles** that the headline violates or supports, or the potential **consequences** of publishing the headline. We will ask you to write down two of the reasons.

*[Text displayed in the **intuitive-thinking condition** only:]*

Research shows that surveys are most informative when respondents like you answer each question **quickly, based on gut feelings**.

We will now ask you to rate how unethical it would be to publish some headlines.

You should make each rating based on your **first instinct**. **Pay attention to your feelings**, and **don't think too hard**. Your ratings should reflect how unethical the headline seems to you at first glance.

- 2) In the deliberative-thinking condition only, we displayed the following two questions just before we asked about how unethical and acceptable the relevant headline would be to publish:

Please provide one reason why you believe publishing this headline is ethical or unethical.

Please provide another reason why you believe publishing this headline is ethical or unethical.

- 3) Omit the question, “How likely would you be to respond by posting a negative comment?”
- 4) After the questions about whether each headline was true or false, participants completed the Cognitive Reflection Task (Frederick, 2005):

Please answer the following questions.

A bat and a ball cost \$1.10 in total. The bat costs \$1.00 more than the ball. How much does the ball cost?

If it takes 5 machines 5 minutes to make 5 widgets, how long would it take 100 machines to make 100 widgets?

In a lake, there is a patch of lily pads. Every day, the patch doubles in size. If it takes 48 days for the patch to cover the entire lake, how long would it take for the patch to cover half the lake?

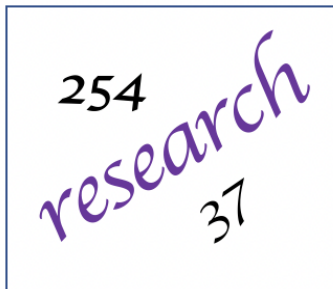
Appendix 10

Verbatim Text of Materials for Experiment 4

Experiment 4's materials were identical to Experiment 1's, with the following exceptions:

- 1) Added a CAPTCHA test just after the question beginning, "In today's study, you will respond to a variety of questions ...":

To let us know that you are not a robot, please type the letters you see in the following image into the response box below. Ignore any numbers you may see. Note that the response box is case sensitive.



- 2) In Phase 1, replace "well-written" with "surprising" in the following two sentences: "Now we'll ask you how well-written the headline is." and "How well-written is this headline?"
- 3) Omit the following text from the beginning of Phase 2:

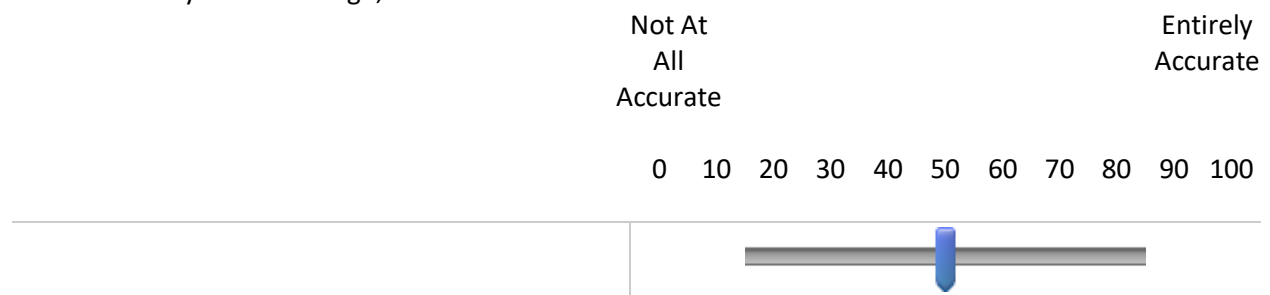
Now, you will move on to the final part of today's study. For this part of the study you will be asked to read a series of fake news headlines that were recently published online. The information in these headline is not real. Non-partisan fact-checking websites have confirmed that these headlines describe events that did NOT happen. After reading each headline, you will be asked to answer a few short questions.

Replace it with the following text:

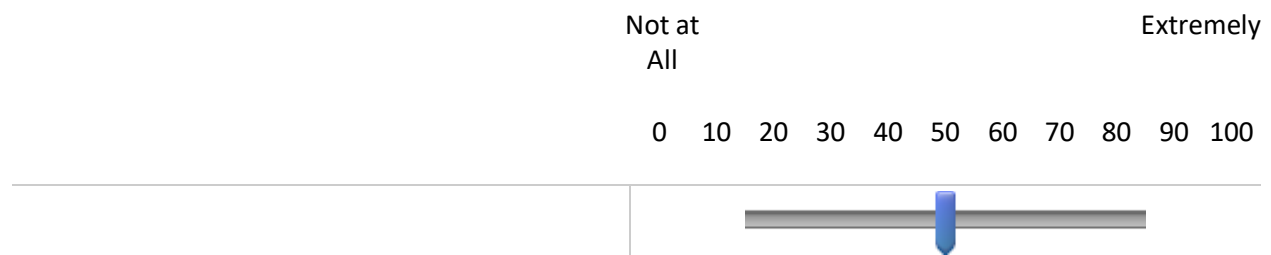
Thanks. Now we have some more questions about news headlines that were published on social media.

- 4) In Phase 2, replace all questions from “How unethical would it be to publish this headline?” to “To the best of your knowledge, how accurate is the claim in the above headline?” with the following questions, presented in randomized orders:

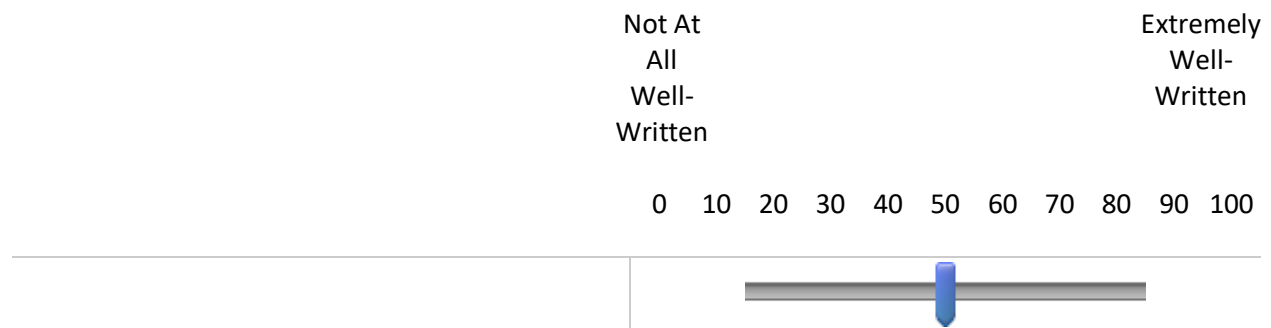
To the best of your knowledge, how **accurate** is the claim in the above headline?



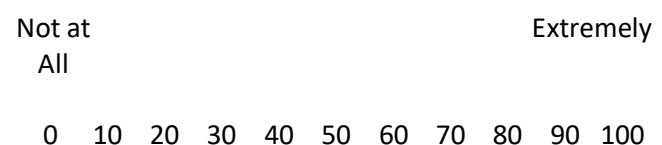
How **unethical** would it be to share the above headline on social media?



How **well-written** is the above headline?



How much do you **like** the above headline?





What **percentage** of American social media users do you think have seen this headline before?

Nobody Everyone

0 10 20 30 40 50 60 70 80 90 100



If one of your acquaintances shared this headline on social media, how likely would you be to share it?

Not at All Likely 1	2	3	Somewhat Likely 4	5	6	Extremely Likely 7
☐	☐	☐	☐	☐	☐	☐

- 5) Added a new behavioral measure after participants had rated all headlines on the measures listed in the previous point (#4):

We would like to run a study where research participants see some of the headlines you've been looking at today. These participants can choose to click on the headlines to read the full article.

The headlines shown below are the ones you've seen today. Please choose **four** of them to share with the participants in our next study.

All 12 headlines were shown in a grid in randomized orders, and participants ticked boxes next to the 4 they selected to share.

- 6) Omitted the question at the end of Phase 2 asking “Which statement is correct about these 12 headlines?” and also omitted the questions asking whether each headline was true or false.

- 7) Added debriefing information at the end of the study because participants were not previously told the headlines were false:

Thank you for your participation in today's study.

We are conducting research on how people think about fake news. **All of the news articles you saw were fake – they describe events that, according to independent fact-checkers, did not actually occur.**

References

- Bullock, J. G., Green, D. P., & Ha, S. E. (2010). Yes, but what's the mechanism? (Don't expect an easy answer). *Journal of Personality and Social Psychology*, 98(4), 550-558.
- De keersmaecker, J., Dunning, D., Pennycook, G., Rand, D. G., Sanchez, C., Unkelbach, C., & Roets, A. (in press). Investigating the robustness of the illusory truth effect across individual differences in cognitive ability, need for cognitive closure, and cognitive style. *Personality and Social Psychology Bulletin*. doi:10.1177/0146167219853844
- Dechêne, A., Stahl, C., Hansen, J., & Wänke, M. (2010). The truth about the truth: A meta-analytic review of the truth effect. *Personality and Social Psychological Review*, 14(2), 238-257. doi:10.1177/1088868309352251
- Fiedler, K., Schott, M., & Meiser, T. (2011). What mediation analysis can (not) do. *Journal of Experimental Social Psychology*, 47(6), 1231-1236.
- Frederick, S. (2005). Cognitive reflection and decision making. *Journal of Economic Perspectives*, 19(4), 25-42.
- Green, P., & MacLeod, C. J. (2015). SIMR: An R package for power analysis of generalized linear mixed models by simulation. *Methods in Ecology and Evolution*, 7(4), 493-498. doi:<https://doi.org/10.1111/2041-210X.12504>
- Hayes, A. F. (2013). *Introduction to Mediation, Moderation, and Conditional Process Analysis: A Regression-Based Approach*. New York: Guilford.
- Kwan, L. Y. Y., Yap, S., & Chiu, C.-y. (2015). Mere exposure affects perceived descriptive norms: Implications for personal preferences and trust. *Organizational Behavior and Human Decision Processes*, 129, 48-58. doi:10.1016/j.obhdp.2014.12.002

Pennycook, G., Cannon, T., & Rand, D. G. (2018). Prior exposure increases perceived accuracy of fake news. *Journal of Experimental Psychology: General*, 147, 1865–1880.

doi:10.1037/xge0000465

Pennycook, G., & Rand, D. G. (2019). Lazy, not biased: Susceptibility to partisan fake news is better explained by lack of reasoning than by motivated reasoning. *Cognition*, 188, 39-50.

doi:10.1016/j.cognition.2018.06.011

Weaver, K., Garcia, S. M., Schwarz, N., & Miller, D. T. (2007). Inferring the popularity of an opinion from its familiarity: A repetitive voice can sound like a chorus. *Journal of personality and social psychology*, 92(5), 821-833.

Zajonc, R. B. (1968). Attitudinal effects of mere exposure. *Journal of personality and social psychology*, 9, 1-27. doi:10.1037/h0025848