

**Repeated by Many Versus Repeated by One: Examining the Role of Social Consensus in
the Relationship Between Repetition and Belief**

Raunak M. Pillai

Lisa K. Fazio

Department of Psychology and Human Development, Vanderbilt University

Word count (Introduction & Discussion sections): 2,980

Author Note

This material is based upon work supported by the National Science Foundation Graduate Research Fellowship Program under Grant No. 1937963. This material is also based on work supported by the National Science Foundation under grant No. 2122640 to the second author. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

All data, materials, preregistration documents, and supplemental analyses are available at the project's Open Science Framework (OSF) site: <https://osf.io/z7bqy> (Pillai & Fazio, 2024).

Correspondence concerning this article should be addressed to Raunak Pillai, Department of Psychology and Human Development, Vanderbilt University, 230 Appleton Place, Nashville, TN 37203. Email: raunak.m.pillai@vanderbilt.edu

Abstract

Repetition makes statements seem more true. Studies of this illusory truth effect typically focus on verbatim repetition of claims without a specific source. However, in real life, we often encounter claims with different wordings from multiple sources. Prominent cognitive theories suggest this variation should not matter. Repetition should increase belief by making statements easier to process—regardless of its wording or source. However, theories of social influence suggest that repetition should be most influential when it reflects a consensus, like when it is paraphrased or repeated by different people. We evaluate these perspectives across three experiments (N = 718 US-based MTurk/Connect participants). Participants saw repeated claims that varied in their phrasing (verbatim versus paraphrased) and/or number of unique sources (one versus multiple), then rated the truth of these claims along with new ones. In line with cognitive theories, repetition increased belief, regardless of variation in the wording or source.

Keywords: truth; belief; social consensus; fluency; illusory truth; repetition

General Audience Summary

In our daily lives, we often encounter information repeatedly. For instance, our social media feeds may be abuzz with posts from different accounts reporting about the latest current event. A long line of research suggests that repeated exposure can make information seem more true—regardless of whether or not it actually is. However, past research has not looked at the whether the effects of repetition differ when it comes from a single, repetitive source versus many distinct sources. For example, imagine if you opened your social media feed to a single account repeatedly sharing the same description of an event over and over, rather than seeing many different accounts describe the event, each in their own way. Would both of these types of repetition have the same effect on your belief? Across three experiments, we address this question. In all experiments, we had US-based participants read a series of news headlines. Some of these headlines were repeated in the same wording, and others were repeated in different, paraphrased forms taken from different news outlets. Additionally, some news headlines were indicated as being shared by the same social media user multiple times, and others were indicated as being shared by different users each time. Finally, after reading these headlines, we had participants evaluate whether these headlines, along with some new ones they had not seen before, seemed true or false. Repetition increased belief in the news headlines, consistent with past research. Interestingly, however, it did not matter whether the repetition was verbatim versus paraphrased or whether it came from one source versus many. These results suggest that repetition can increase belief in information simply by making the key idea feel easier to mentally process—regardless of variation in how the information is worded or how many distinct people repeat it.

Repeated by Many Versus Repeated by One: Examining the Role of Social Consensus in the Relationship Between Repetition and Belief

Repeated exposure makes statements seem more true (Hasher et al., 1977). This “illusory truth effect” has been replicated across dozens of experiments, the vast majority of which examine how one additional verbatim exposure to a statement affects belief (Henderson et al., 2021). In real life, however, we often see information multiple times, in different wordings and from different sources. For instance, we may hear about an event separately from multiple witnesses or repeatedly from a single, incessant source. These different kinds of repetition connote different levels of consensus and agreement, but will they have different impacts on belief? This question is particularly important at a time where many people are following news via social media (e.g., Pew Research Center, 2023), a medium that allows single actors to repeatedly share information with ease. This paper examines how repetition affects belief when that repetition varies in phrasing (verbatim versus paraphrased) and in the number of unique sources (one versus multiple), testing predictions from three theoretical perspectives.

Explanations of the Illusory Truth Effect

The most prominent current account of the illusory truth effect suggests that repetition increases processing fluency, or the subjective ease with which mental operations occur (Reber & Schwarz, 1999). People may then use fluency as a cue for truth because they have implicitly learned from their environments that fluently processed information tends to be true or because they explicitly believe that fluency connotes truth (Unkelbach & Greifeneder, 2013).

Importantly, repetition affects two kinds of fluency. First, repetition increases *perceptual fluency*—the ease associated with visually or auditorily perceiving a statement. Second, repetition increases *conceptual fluency*—the ease with which the meaning of a statement is

processed. This conceptual fluency is thought to contribute more strongly to judgements of truth than perceptual fluency (Vogel et al., 2020).

Other current theories offer related explanations (see Unkelbach et al., 2019 for a review). For instance, Unkelbach and Rom (2017) suggest that repetition strengthens the coherence of peoples' semantic representations for statements in memory, and this semantic coherence enhances perceived truth. Because this theory and the conceptual fluency account make identical predictions, we frame our hypotheses in terms of fluency for simplicity.

Critically, all of the current accounts focus on low-level cognitive processes, ignoring the potential role of social information, like broad endorsement of the claim. Repetition should increase processing fluency similarly regardless of whether the repetitions come from a single person or a wide social consensus, leading to similar increases in belief.

A Possible Role of Social Consensus

In contrast to the theories highlighted above, there are good reasons to believe that social consensus may affect perceptions of truth. Classic research on conformity (Asch, 1956) and opinion formation (Festinger, 1954) suggests that people alter their thoughts and actions in line with what is socially agreed upon. For instance, people are more confident in claims supported by multiple, independent sources of evidence than claims supported by a single repetitive source (Connor Desai et al., 2022). Relatedly, people are more likely to accept that humans are causing climate change after reading about the overwhelming scientific consensus (Lewandowsky et al., 2013). Importantly, people are not only sensitive to expert consensus—consensus among peers is also influential. For example, people are more likely to agree with simulated blogs and social media posts when there is strong agreement among peers who reply, comment, or post about the topic (Lewandowsky et al., 2019; Ransom et al., 2021; Simmonds et al., 2023).

While people are clearly sensitive to social consensus, it is unclear whether this sensitivity will influence belief in repeatedly encountered statements. In past research, participants received information about social consensus in a single, uninterrupted experience, like reading a thread of comments under a post (e.g., Ransom et al., 2021), a summary statistic (e.g., 97% of experts agree; Lewandowsky et al., 2013) or a series of articles on the same topic (Connor Desai et al., 2022). In these cases, consensus is either directly communicated or can be easily comprehended by comparing adjacent pieces of information (e.g., noticing a comment agrees with the ones below it). Here, we are instead concerned with peoples' sensitivity to social consensus cues around *repeated* information, where consensus must be inferred across discrete instances of exposure to information. For example, imagine you are scrolling through a social media newsfeed, seeing the same news repeatedly, but separated by several other, unrelated posts. This situation requires that people a) track whether information was repeated by multiple sources or one in memory b) interpret this memory for multiple sources as a signal for broad agreement and c) use this perceived consensus as a cue for truth judgements. While the work reviewed above justifies the latter assumption, the first two are less certain.

The Present Research

The present experiments examine how repetition affects belief when that repetition reflects a consensus versus a single, repetitive source. In Experiment 1, participants read news headlines that were repeated verbatim or in paraphrased forms derived from different news outlets (with no source information). In Experiment 2, participants read headlines shared by one person multiple times or by multiple different people (with no wording variation). Finally, Experiment 3 integrated these manipulations, showing participants verbatim repetitions of headlines from a single source or paraphrased repetitions from multiple sources. The key

question across experiments is how participants then rate the truth of these headlines relative to each other and to new headlines that were not previously encountered.

If higher-level inferences about social consensus matter for truth judgements, varying the phrasing or source of repeated headlines should increase belief. By contrast, if repetition affects belief through lower-level cognitive processes alone, neither manipulation should increase belief. The conceptual fluency account suggests that repetition should increase belief so long as the same ideas are repeated, regardless of their wording or source. The perceptual fluency account suggests that paraphrased repetition should be *less* impactful, since perceptual features of the statement change on each exposure, and broadly suggests that source variability should not matter. Table 1 summarizes these predictions. To preview, our results are most consistent with a conceptual fluency account: repetition increases belief to a similar degree, regardless of variation in wording or source.

Table 1

Predictions Across Experiments

Account	Experiment 1	Experiment 2	Experiment 3
Fluency Alone			
Conceptual	Verbatim = Paraphrased	1 source = 3 sources	Verbatim, 1 source = Paraphrased, 3 sources
Perceptual	Verbatim > Paraphrased	1 source = 3 sources	Verbatim, 1 source > Paraphrased, 3 sources
Social Consensus	Verbatim < Paraphrased	1 source < 3 sources	Verbatim, 1 source < Paraphrased, 3 sources

Note. Key predictions about the effects of different kinds of repetition on belief. Cells indicate which conditions are predicted to result in the highest perceived truth.

Open Practices

The hypotheses, design and analysis plan for all experiments were pre-registered. The pre-registration documents are available at the project's Open Science Framework (OSF) site (<https://osf.io/z7bqy>), along with the materials, participant instructions, data, and analysis code

for each experiment. For all experiments, we report how we determined our sample size, all data exclusions, all manipulations, and all measures in the study.

Experiment 1

In Experiment 1, participants read a set of headlines repeated three times in verbatim or paraphrased form. They then rated the truth of the verbatim headlines, a fourth unique version of the paraphrased headlines, and some completely new headlines.

A few studies have examined how paraphrased repetition affects belief (Silva et al., 2017; Vogel et al., 2020), using strict linguistic rules to create paraphrased statements (e.g., “The pigeon has a lifetime superior to that of a rabbit” becomes “A rabbit has a lifetime inferior to that of a pigeon”). These studies find that verbatim and paraphrased repetitions have similar impacts on belief.

By contrast, the present experiment used more naturalistic paraphrased stimuli, drawn from different news outlets reporting on the same event. These stimuli vary more drastically in tone, details and wording (e.g., “A study found that an iPhone 12 can disable a cardiac rhythm management device” versus “Cardiologists Find Apple iPhone 12 Magnet Deactivates Implantable Cardiac Devices”). Thus, while these paraphrased versions convey the same gist meaning, they vary in their surface features (see Reyna & Brainerd, 1995 for a discussion of this distinction), creating greater perceptual variability. These headlines may also convey social information, like the fact that these different headlines come from different sources.

We predicted that headlines repeated in verbatim or paraphrased form would be rated as truer than new statements, replicating the illusory truth effect. Critically, we also predicted that participants would provide higher truth ratings for headlines repeated in verbatim versus paraphrased form. This prediction was based on current cognitive theories suggesting that both

types of repetition would increase conceptual fluency, but that verbatim repetition would most greatly increase perceptual fluency.

Method

All experiments received ethics approval under IRB #170586 at the authors' institution.

Participants

Statistical Power. Our pre-registered sample size was based on an *a priori* power analysis in G*Power (Faul et al., 2009), which revealed that 262 participants was needed to achieve 80% power to detect a an effect size of $f = .055$ in a two-group repeated measures ANOVA.

This two-group repeated measures ANOVA is equivalent to the paired *t*-test we planned to run comparing participants' mean responses to items in the repeated paraphrase and repeated verbatim conditions. We were minimally interested in a difference between these conditions of 0.1 points on our 6-point scale, and, assuming a standard deviation of these ratings of 0.88 from past studies with similar materials and numbers of items (Pillai et al., under review), this difference corresponds to a minimal effect size of interest of $f = .055$. For context, this minimal effect size of interest is about 4.5 times smaller than the overall expected effect of repetition (verbatim repetition vs. new) based on past meta-analytic evidence ($f = 0.25$; Dechêne et al., 2010). Finally, our power analysis also assumed a correlation among repeated measures of 0.8 based on prior research (Pillai, et al., under review) and no correction for non-sphericity.

Recruitment. 262 adult participants were recruited from Amazon's Mechanical Turk (MTurk) platform to complete the experiment through the CloudResearch platform (Litman et al., 2017) for a payment of \$1.81. To ensure data quality, we recruited participants from CloudResearch's approved participants list (Peer et al., 2021) and excluded participants (not

counting towards the above sample size) for failing two attention checks at the beginning of our survey (typing a response to “Puppy is to dog as kitten is to ____?” and selecting two requested responses on a 5-point multiple choice question).

Demographics. The mean age of participants was 39.69 ($SD = 11.37$; Range = 18-76). Our final sample was predominantly White (76%, 9.5% Black, 8.8% Asian, 2.3% Multiracial, 1.5% Other, 1.5% not reporting) and non-Hispanic (89%, 9.9% Hispanic, 0.8% not reporting), and 51% of participants were men (46% women, 0.8% nonbinary, 2.3% not reporting).

Design

We manipulated repetition type (new, repeated verbatim, repeated paraphrase) within-subjects. We counterbalanced repetition across participants by splitting our 36 items into three sets of 12 and rotating these sets through each level of repetition type. This created three possible counterbalancing groups to which we assigned participants.

Materials

Stimuli consisted of news headlines describing 36 different events or facts that were either confirmed by the third-party fact-checking site Snopes as “true” or were reported by various reputable mainstream news outlets (e.g., The New York Times, The Washington Post). Example headlines include “Exquisitely-preserved wolf pup mummy discovered in Yukon permafrost” and “‘Cocaine bananas’ accidentally shipped to grocers in bungled drug deal”. Note that participants were never shown the original source of the headline—only the text of the headline itself.

Each of the 36 events had four different headlines, each from a different online source. (e.g., “Cocaine found in banana shipment part of drug deal gone bad,” “Cocaine-stuffed shipments of bananas ended up at Canadian grocery stores due to a drug-trafficking mix-up.”)

“‘Cocaine bananas’ accidentally shipped to grocers in bungled drug deal,” “‘Cocaine bananas’ shipped to grocery stores in botched operation”). For each of the 36 items, we randomly selected one of the four versions as the key headline that was shown during the rating phase. For new items, this is the only version participants saw in the experiment. For repeated verbatim items, participants saw the key headline three times during the exposure phase and then again in the rating phase. For repeated paraphrased items, participants saw the other three versions of the headline during the exposure phase followed by the key headline in the rating phase. Table 2 shows an example item in each of these three conditions. Note that each participant saw only one of these three conditions (new, repeated verbatim or repeated paraphrased) for a given item.

Table 2

Sample Headline in New, Repeated Verbatim, and Repeated Paraphrase Conditions

Condition		Exposure Phase			Rating Phase
New					A Hacker Tried to Poison a Florida City’s Water Supply, Officials Say
Repeated Verbatim		A Hacker Tried to Poison a Florida City’s Water Supply, Officials Say	A Hacker tried to Poison a Florida City’s Water Supply, Officials Say	A Hacker tried to Poison a Florida City’s Water Supply, Officials Say	A Hacker tried to Poison a Florida City’s Water Supply, Officials Say
Repeated Paraphrased		In Florida City, a Hacker Tried to Poison the Drinking Water	Someone tried to poison Oldsmar, Florida’s water supply during hack, sheriff says	Feds tracking down hacker who tried to poison Florida town’s water supply	A Hacker Tried to Poison a Florida City’s Water Supply, Officials Say

Note. Participants saw a given headline in one of the three conditions.

Procedure

This experiment was administered online via Qualtrics.

Exposure Phase. After reading the information sheet and completing two attention checks, participants were instructed that we wanted to get their opinion on “various claims that have been posted online.” As all headlines were true, we did not inform participants of the truth of the headlines they were about to see. Starting on the following screen, participants saw 72 headlines, one at a time, in the center of the screen above the question “How interesting is the headline above?” Participant selected from the options *Very Uninterested*, *Uninterested*, *Slightly Uninterested*, *Slightly Interested*, *Interested*, or *Very Interested* to proceed (as in past studies of the illusory truth effect (e.g., Fazio, 2020)). The 72 headlines consisted of 12 items repeated three times verbatim, and 12 items repeated three times in paraphrased form. These headlines were presented in a random order for each participant (as were all headlines in all phases and experiments).

Rating Phase. Immediately after the exposure phase, participants began the rating phase. Participants were correctly informed “some of the headlines you will have seen in the previous section, others will be new.” Again, participants were not informed about the truth of the headlines. Participants then saw the 36 key headlines, one at a time, above the question “How true or false is the headline above?” and selected from the options *Definitely False*, *Probably False*, *Possibly False*, *Possibly True*, *Probably True*, or *Definitely True* (scored from 1 to 6 in our analyses). Twelve headlines were new (i.e., shown for the first time on the rating phase), 12 were repeated verbatim in the exposure phase, and 12 were repeated in paraphrased form in the exposure phase.

Demographics and Debrief. Participants then answered some optional demographic questions (gender, race, ethnicity, education level) and were asked a few debriefing questions (what they thought the study was about, whether they noticed the same statement multiple times

in the first phase, whether they noticed different versions of the statement in the first phase). Finally, participants were thanked for their time and informed about the purpose of the study.

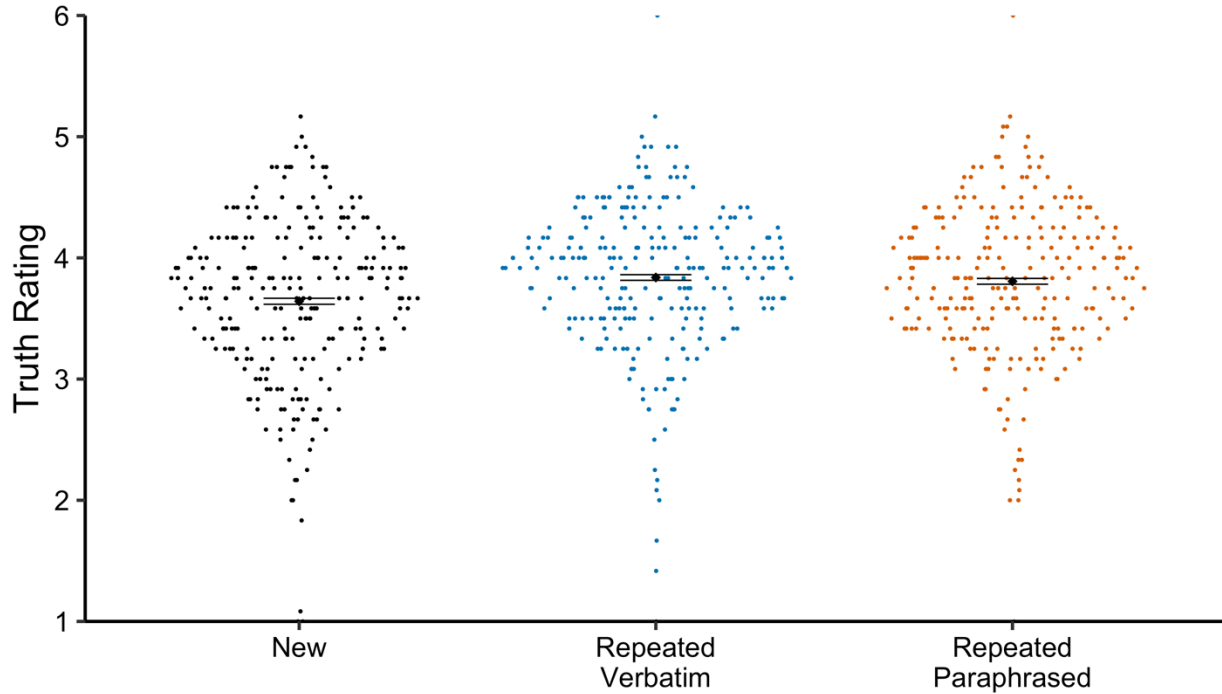
Results

For all experiments, all statistical tests are conducted at the .05 alpha level and are pre-registered unless labelled as exploratory. ANOVA tests were conducted using *rstatix* version 0.7.0 (Kassambara, 2020), frequentist *t*-tests were conducted in base R version 3.6.3 (R Core Team, 2020), and Bayesian *t*-tests were conducted using *BayesFactor* version 0.9.12 (Morey et al., 2015).

We hypothesized that, in line with the illusory truth effect, repetition (in both verbatim or paraphrased form) would increase belief. As shown in Figure 1, this was the case. Critically, we also predicted that headlines repeated in verbatim form would be perceived as more true than repeated paraphrased headlines. Contrary to our hypothesis, truth ratings were similar across the verbatim and paraphrased headlines.

Figure 1

Mean Truth Ratings for Headlines by Repetition Status (Experiment 1)



Note. Ratings were provided on a scale from 1 = *Definitely False* to 6 = *Definitely True*. Each dot represents one participant ($N = 262$) with values horizontally shifted to represent the density distribution. Black diamonds reflect group means and error bars reflect standard errors of the mean.

A one-way ANOVA revealed a significant main effect of repetition type (new, repeated verbatim, repeated paraphrase), $F(2, 522) = 17.58, p < .001, \eta_p^2 = 0.06$. Relative to new items ($M = 3.64, SD = 0.68$), participants gave higher truth ratings to items repeated verbatim ($M = 3.84, SD = 0.62$), $t(261) = 5.08, p < .001$, 95% confidence interval of the difference (CI) [0.12, 0.27], $d = 0.27$ and items repeated in paraphrased form ($M = 3.81, SD = 0.63$), $t(261) = 4.41, p < .001$, 95% CI [0.09, 0.24], $d = 0.31$. However, we did not observe a significant difference in ratings for items repeated in verbatim versus in paraphrased form, $t(261) = 1.05, p = .295$, 95% CI [-0.03, 0.09], $d = 0.06$.

To follow up on this null result, we conducted an exploratory Bayesian *t*-test comparing participant's mean ratings for items repeated in verbatim versus paraphrased form. Using the default Cauchy distribution with width 0.707 for our prior probability distribution, we calculated a Bayes factor of 8.40 in favor of the null hypothesis. That is, our data are 8.40 times more likely under the null hypothesis (that ratings are identical in the repeated verbatim and repeated paraphrased conditions) than under the alternative (that there is a difference). This Bayes Factor is in the range generally considered moderate evidence in favor of the null hypothesis (Held & Ott, 2018).

Discussion

Consistent with a conceptual fluency account, repeating an idea using the same or different wording increased belief to a similar degree. These results contradict predictions from both a perceptual fluency account (verbatim repetition should be most impactful), and a social consensus account (paraphrased headlines should indicate endorsement by different sources, increasing belief). However, Experiment 1 has two important limitations. First, our manipulation may have evoked multiple mechanisms. Perceptual fluency may have increased belief in verbatim headlines, and social consensus may have increased belief in paraphrased headlines, producing null results. Second, participants may not have inferred that the different wordings indicated different sources.

Experiment 2

Thus, Experiment 2 examines a more direct manipulation of consensus: the number of unique people sharing a headline. To our knowledge, only one prior paper has examined the role of source variability in the effects of repetition on belief (Roggeveen & Johar, 2002), finding no effect in one study (Experiment 2) and a limited effect in another (Experiment 3). However, the

former study presented sources during both the initial exposure and the rating phase. Thus, social consensus was confounded with the level of perceptual overlap between the repeated stimuli. In addition, the latter study manipulated repetition between-subjects, a design that minimizes the effects of repetition on belief (Dechêne et al., 2009).

Addressing these limitations, we showed participants news headlines three times (verbatim), alongside the same source or a different source each time. Then, participants rated the perceived truth of these headlines and new headlines without any sources.

By the social consensus account, people should be most likely to believe statements repeated by different sources. By contrast, the conceptual and perceptual fluency accounts suggest that repetition should increase belief similarly regardless of who shared it.

We again predicted that repetition from one or three sources would increase belief. However, we did not make a prediction about whether one kind of repetition would be more impactful. Finally, to examine how well participants tracked source variability, we added a memory check asking participants how many different people shared each headline.

Method

Participants

Statistical Power. Our pre-registered sample size was based on an *a priori* power analysis in G*Power, which indicated that 229 participants were needed to achieve 80% power to detect an effect size of $f = .08$ in a two-group repeated measures ANOVA (equivalent to a paired t -test comparing items repeated from one or three sources).

This power analysis is identical to that of Experiment 1 except in two ways. First, the minimal effect size of interest increased from $f = .055$ to $f = .08$, as we used the observed *SD* of truth rating data in Experiment 1 ($SD = .62$ for repeated verbatim & paraphrased items, versus

$SD = 0.88$ used in the power analysis in Experiment 1) . The minimal effect size of interest still corresponds to an absolute difference of 0.1 points on our 6-point scale. Second, we reduced the correlation among repeated measures from 0.8 to 0.63, again based on the observed correlation value for repeated verbatim and repeated paraphrased items in Experiment 1.

Recruitment. 229 adult participants were recruited in the same manner as described in Experiment 1, except participants received \$2.42 for completing this longer experiment. Four additional participants were excluded for failing our two attention checks (typing the name of the animal depicted by a black and white cartoon image and selecting a requested response on a 5-point Likert item) and did not count towards our final sample.

Demographics. The mean age of participants was 39.76 ($SD = 11.11$; Range = 19-77). Our final sample was predominantly White (81 %, 7.4% Asian, 7.0 % Black, 3.1% Multiracial, 0.4% Other, 1.3% not reporting) and non-Hispanic (94 %, 4.4% Hispanic, 1.8% not reporting), and 52% were women (46 % men, 0.9% nonbinary, 0.9 % not reporting).

Design

We manipulated repetition type (new, repeated from one source, repeated from three sources) within-subjects. As in Experiment 1, participants were assigned to one of three counterbalancing groups, created by splitting items into three sets and rotating them through the three levels of repetition type.

Materials

We used the 36 key headlines shown in the rating phase of Experiment 1. Note that only one version of each headline was used in this experiment.

During the exposure phase, headlines were paired with sources, which were full-body photographs created by Connor et al. (2021). The full set includes 454 photos of people whose

race (Asian, Black, White) and gender (male, female) were noted. We selected 48 photos for this experiment, attempting to evenly sample across all combinations of race and gender present in the database. However, as the set only consisted of three photos of Asian women, we selected all three, and then randomly selected nine photos within each of the remaining five combinations of the photo subject's gender and race.

For each participant, 12 of the 48 sources were randomly assigned to the “repeated from one source” condition and the remaining 36 to the “repeated from three sources” condition. In the one source condition, the 12 headlines were randomly paired with a single source and repeated three times with that same source during the exposure phase. In the three-source condition, each of the three instances of the 12 headlines had a different, unique source, exhausting all 36 remaining sources. No source was paired with more than one headline.

Procedure

This experiment was administered online using the gorilla.sc platform (Anwyl-Irvine et al., 2020).

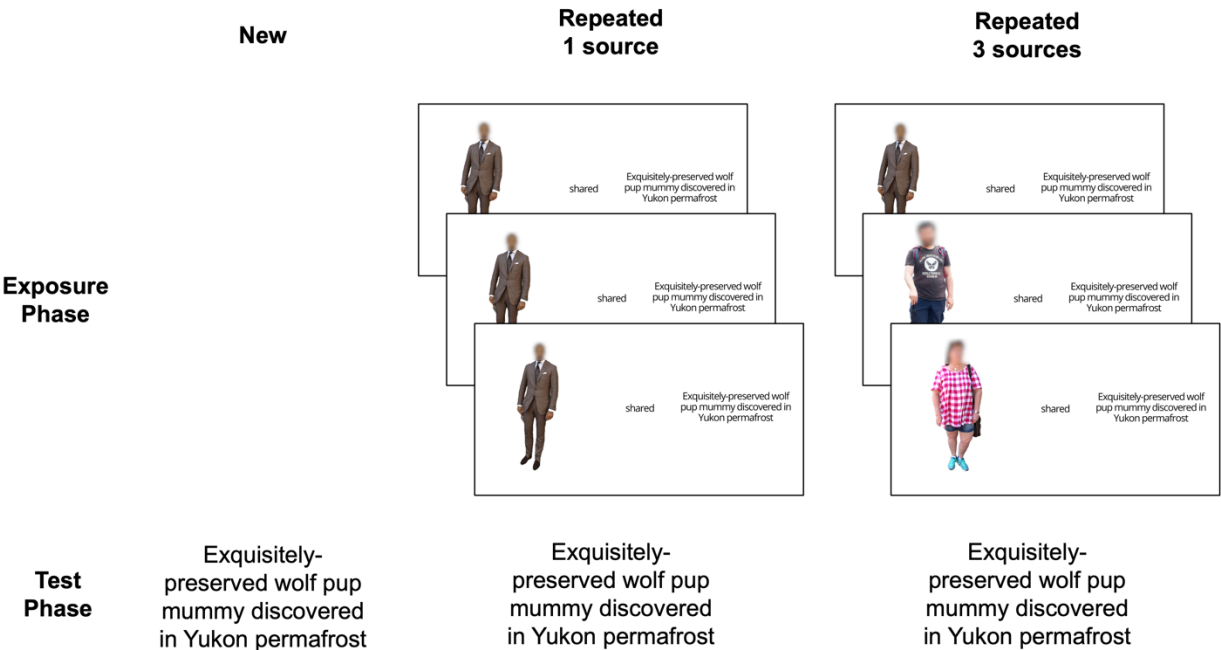
Exposure Phase. As in Experiment 1, participants began by reading an information sheet, completing two attention checks, and receiving instructions to rate their interest in “various claims that have been posted online.” Unlike in Experiment 1, participants were also told that headlines would appear next to a photo of a person who shared it online and were told “Please pay attention to who shared each headline. We will ask you some questions about who shared each headline at the end of this experiment.” We again did not inform participants of the truth of the headlines.

Then, participants saw 72 headlines, one at a time, and were asked “How interesting is the headline above?” (*Very Uninteresting, Uninteresting, Slightly Uninteresting, Slightly*

Interesting, Interesting, or Very Interesting). Unlike in Experiment 1, headlines were placed next to a full body photo of a person, as shown in Figure 2. As described above, 12 headlines were repeated three times with the same source and 12 were repeated three times with different sources each time. All headlines were repeated verbatim.

Figure 2

Sample Headline in New, Repeated (1 Source), and Repeated (3 Sources) Conditions



Note. Participants saw a given headline in one of the three conditions. The exact sources associated with any headline were randomized on a participant-by-participant basis. Faces are blurred for anonymity; participants saw images with unobstructed faces.

Rating Phase. The rating phase was identical to that of Experiment 1, except that we also told participants, “we will not indicate who shared each headline” . All participants rated the

truth of 36 headlines by responding to the question “How true or false is the headline above?” (*Definitely False, Probably False, Possibly False, Possibly True, Probably True, or Definitely True*, scored from 1 to 6).

Memory Check. After the rating phase, participants completed an additional, exploratory measure of their memory for the sources of each headline. Participants were shown 12 headlines (randomly selected for each participant), with even numbers of headlines that were new, repeated from one source, and repeated from three sources. For each of the 12 headlines, participants were asked “How many different people shared this headline?” (0, 1, 2, or 3). Participants were instructed that this task referred to the sources that shared each headline during the first part of the experiment (exposure phase).

Demographics and Debrief. Finally, participants answered optional demographic questions (gender, race, ethnicity, education), a debrief question (what they thought the purpose of the study was), and were informed about the purpose of the study.

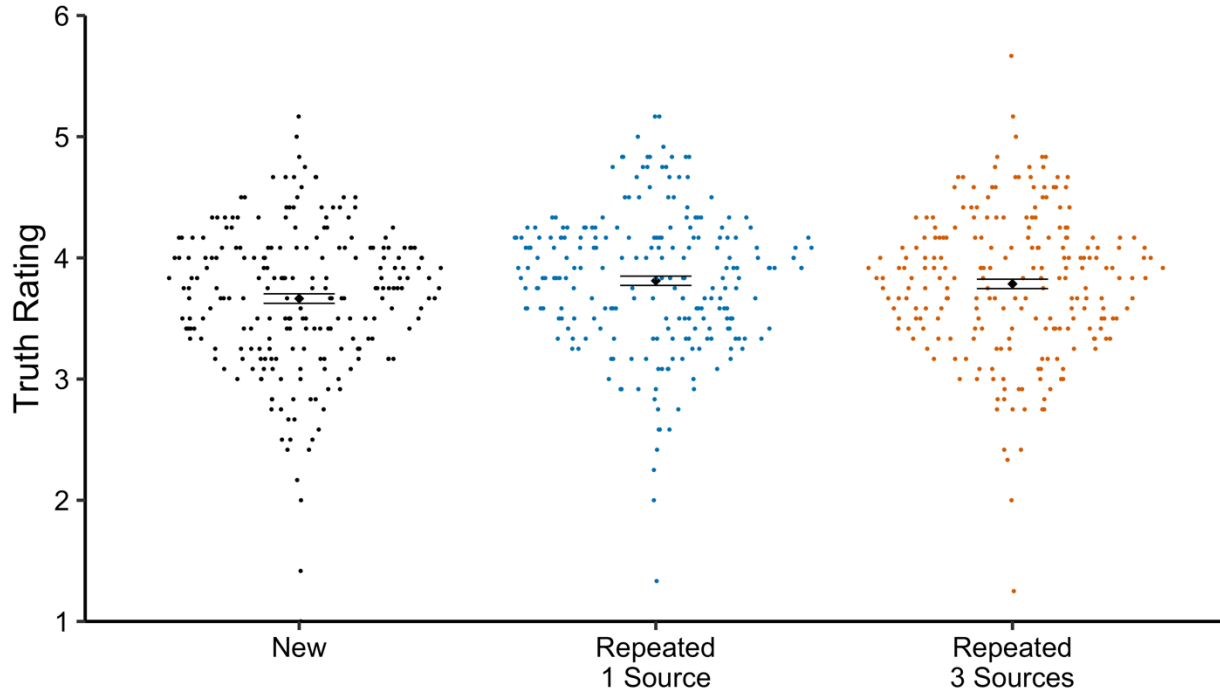
Results

Truth Ratings

We hypothesized that, in line with the illusory truth effect, repetition of headlines from a single source or from multiple sources would increase belief. As shown in Figure 3, this was this case. Our main research question was whether headlines repeated from a single source would be perceived as more or less true than headlines repeated from multiple sources. As shown in Figure 3, ratings were similar across these two conditions.

Figure 3

Mean Truth Ratings for Headlines by Repetition Status (Experiment 2)



Note. Ratings were provided on a scale from 1 = *Definitely False* to 6 = *Definitely True*. Each dot represents one participant ($N = 229$) with values horizontally shifted to represent the density distribution. Black diamonds reflect group means and error bars reflect standard errors of the mean.

A one-way ANOVA revealed a significant main effect of repetition type (new, repeated from one source, repeated from three sources), $F(2, 456) = 11.50, p < .001, \eta_p^2 = 0.05$. Relative to new items ($M = 3.66, SD = 0.59$), participants gave higher truth ratings to headlines repeated from one source ($M = 3.81, SD = 0.58$), $t(228) = 4.41, p < .001$, 95% CI [0.08, 0.21], $d = 0.29$ and headlines repeated from three sources ($M = 3.79, SD = 0.59$), $t(228) = 3.47, p < .001$, 95% CI [0.05, 0.19], $d = 0.23$. However, we did not observe a significant difference in ratings for headlines repeated from one versus three sources, $t(228) = 0.87, p = .387$, 95% CI [-0.03, 0.08], $d = 0.06$.

Like in Experiment 1, we followed up on this null result by conducting an exploratory Bayesian t -test comparing ratings for headlines repeated for one versus three sources. The Bayes factor for this t -test was 9.33, suggesting our data are 9.33 times more likely under the null hypothesis. Again, we find moderate evidence in favor of the null hypothesis that the two kinds of repetition we tested do not have different effects on belief.

Memory for Source Variability

Finally, to verify that participants were able to remember the extent to which there was social consensus around different claims, we conducted an exploratory analysis on participants' responses to the memory check ("How many different people shared this headline?" (0, 1, 2, or 3)). Overall, while estimates were not very accurate, participants were able to qualitatively distinguish between the new, single-source and three-source headlines. Participants provided higher estimates for headlines that were repeated from a single source ($M = 1.92$, $SD = 0.55$) relative to new headlines ($M = 0.68$, $SD = 0.68$), $t(228) = 20.55$, $p < .001$, 95% CI [1.12, 1.36], $d = 1.36$. In addition, participants provided higher estimates for headlines that were repeated from three sources ($M = 2.11$, $SD = 0.49$) relative to headlines that were repeated from a single source, $t(228) = -5.24$, $p < .001$, 95% CI [0.12, 0.26], $d = 0.35$.

While participants were able to qualitatively distinguish between headlines across the different conditions, memory was not particularly accurate. For instance, participants overestimated the number of unique sources for headlines repeated by a single source on average ($M = 1.92$). Given this, we were interested in examining whether the key pattern of truth ratings reported in Figure 3 would differ for participants with more accurate memory for the number of sources. Specifically, we focused on participants with the largest difference in memory responses

between the one-source and three-source conditions (top one third of the sample; $N = 62$)¹. Note that even among these participants memory was not particularly accurate. The average difference between memory ratings was 0.95 ($SD = 0.23$), while perfect memory would be a difference of 2 ($3 - 1$). We then conducted an exploratory t -test comparing truth ratings for headlines repeated by one versus three sources among these participants. As in our main analyses, we found no difference in truth ratings for headlines repeated by one ($M = 3.88$, $SD = 0.56$) versus three ($M = 3.90$, $SD = 0.60$) sources in this condition, $t(61) = -0.12$, $p = .905$, 95% CI $[-0.12, 0.10]$, $d = 0.01$, and we observed moderate Bayesian evidence in favor of this null effect ($BF_{01} = 7.02$).

Discussion

Repetition increased belief regardless of whether it came from one person or many, contrary to predictions that social consensus would magnify the effects of repetition on belief. Experiment 2 is again most consistent with fluency-based accounts in which repetition increases belief by making statements easier to process.

Still, Experiment 2 has one important limitation. Even when participants saw headlines from different sources, these sources repeated the same headline verbatim. In this way, the consensus may not have been perceived as coming from sources who independently endorsed the information. Other work suggests people are less sensitive to social consensus cues when multiple sources co-depend on the same data (Connor Desai et al., 2022; Whalen et al., 2018; but see Yousif et al., 2019). Thus, Experiment 3 has different sources convey paraphrased versions of each headline, making the consensus more likely to reflect a process by which multiple people separately decided to share the same idea. This sort of “independent” consensus may seem more

¹Note that memory difference scores for participants were fairly discrete, taking on values from -1.25 to 2 in increments of 0.25. Thus, in selecting a cutoff for the top tertile of responses, we ended up with slightly less than one third of our full sample size.

informative and thus more likely to affect judgements, providing a stronger test of the social consensus hypothesis.

Experiment 3

Experiment 3 replicates Experiment 2, with two modifications. First, headlines shared by three different sources are now paraphrased repetitions, rather than verbatim. Second, we added exploratory measures to examine how well participants tracked consensus. Experiment 2 suggested that people can differentiate claims repeated by many people versus one person in memory. Here, we examine whether this memory translates into perceptions of consensus around the repeated claims using three exploratory measures. First, we asked participants to estimate what proportion of social media users would agree with a subset of headlines. We also added two questions asking participants why they thought people shared the news headlines they saw, to see if participants thought the sources shared the news without actually *endorsing* its accuracy.

We again predicted that either kind of repetition would increase belief. However, we again did not make a prediction about whether ratings would differ between the two kinds of repetition.

Method

Participants

Statistical Power. Our pre-registered sample size was based on an *a priori* power analysis in G*Power, which indicated that 226 participants were needed to achieve 80% power to detect an effect of $d_z = 0.19$ in a matched-pairs *t*-test. Note that this power calculation is identical to that reported in Experiment 2, except that we directly used the matched-pairs *t*-test option in G*Power rather than the two-group repeated measures ANOVA option. These two

power calculations are mathematically equivalent, but due to rounding differences, produce slightly different sample size estimates (226 versus 229).

Recruitment. 227 participants were recruited via the CloudResearch Connect platform and received \$3.02 for completing the experiment. Note that this is one higher than our pre-registered sample size due to an additional participant completing the experiment before the posting formally closed on Connect. As in Experiments 1 & 2, participants were excluded (not counting towards our final sample size) if they failed two attention checks.

Demographics. The mean age of participants was 38.93 ($SD = 12.56$; Range = 18-77). Our final sample was predominantly White (76%, 12% Black, 5.7% Asian, 2.6% Multiracial, 0.88% reporting some other race, 2.2% not reporting) and non-Hispanic (90%, 6.5% Hispanic, 3.5% not reporting), and 55% were men (40% women, 1.3% nonbinary, 4.4% not reporting).

Design

We manipulated repetition type (new, repeated verbatim from one source, repeated paraphrased from three sources) within-subjects. As in Experiments 1 and 2, participants were assigned to one of three counterbalancing groups, created by splitting items into three sets and rotating them through the three levels of repetition type.

Materials

Stimuli consisted of the 36 news headlines (each with four paraphrased versions) used in Experiment 1. As in Experiment 1, one key version of each headline was used in the rating phase. Depending on the condition, participants either saw this headline three times, saw three other versions of the headline, or did not see the headline at all during the exposure phase.

In addition, we used the same 48 full-body photographs of people as in Experiment 2.

Procedure

The procedure was identical to that of Experiment 2 with a few exceptions. First, during the exposure phase, the headlines shared by different sources were all shared in paraphrased wording each time.

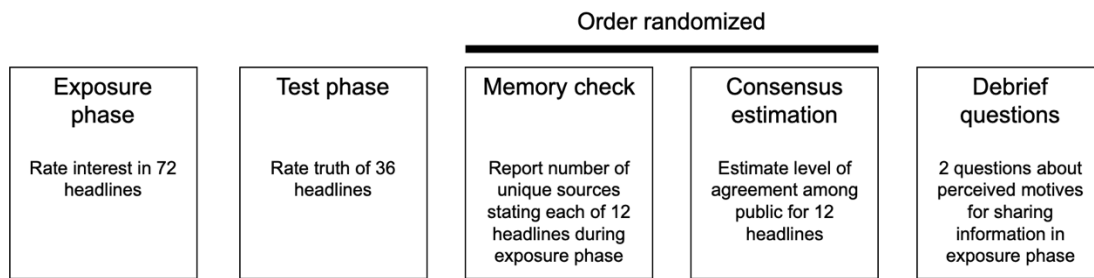
Second, we changed the wording of the memory check, as headlines were repeated not only multiple times, but also in different wordings. Thus, participants answered the question “How many different people shared this headline or a similar headline that contained the same idea” (0, 1, 2, or 3).

Third, in addition to the memory check, we added an additional “social consensus estimation” task. Participants were shown 12 headlines and asked for each “What percent of social media users do you think would believe this headline to be true?” entering their response using a slider from 0 to 100 (default set to 50). To avoid repeating headlines between the memory check and social consensus, we split our 36 key headlines into three sets of 12 (each containing even numbers of headlines per repetition condition), and randomly showed participants one set for the memory check and another for the social consensus estimation task. In addition, the order of these two tasks was randomized across participants.

Finally, after the social estimation and memory tasks, we added two exploratory debrief questions. First, we asked participants the open-ended question “What do you think the main motivation was for people to share the headlines? In other words, why did these social media users decide to share the headlines you saw?” Next, we asked participants “How important do you think it was to these social media users that the information they share is accurate?” (1-10 Likert scale with anchors “*not very important*” and “*very important*”). Figure 4 depicts all phases of this experiment.

Figure 4

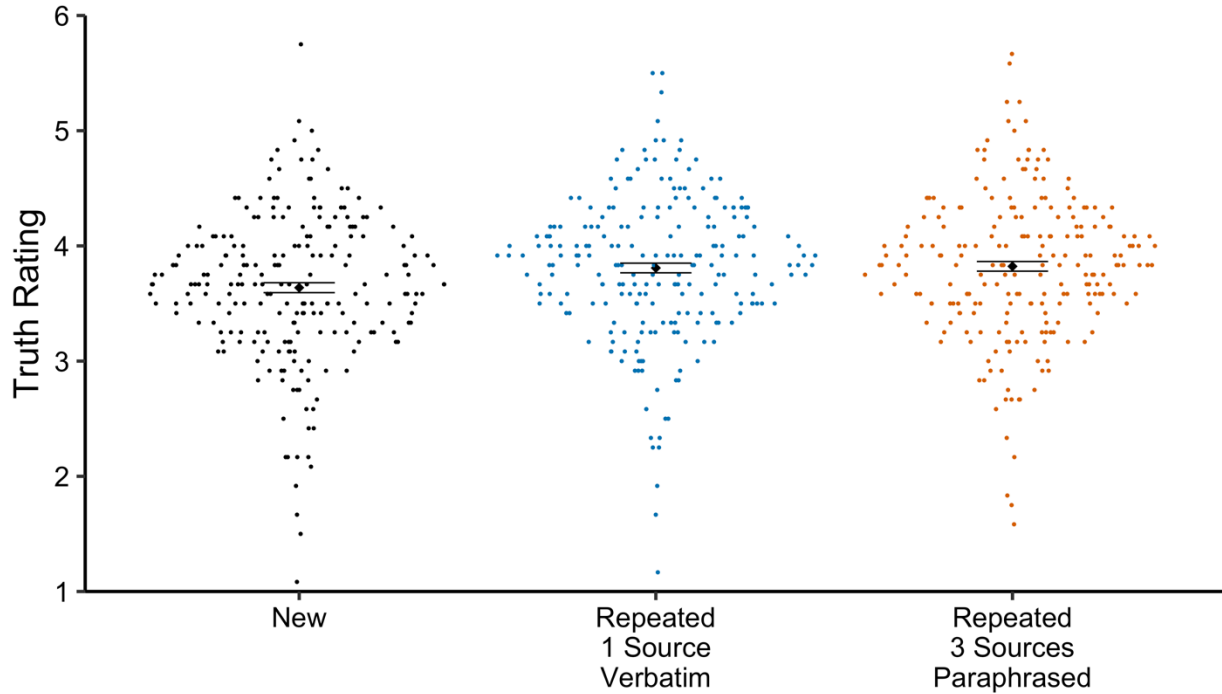
Phases of Experiment 3

**Results*****Truth Ratings***

We hypothesized that, in line with the illusory truth effect, repetition of headlines from a single source or from multiple sources would increase belief. As shown in Figure 5, this was this case. Our main research question was whether headlines repeated verbatim from a single source would be perceived as more or less true than headlines repeated in paraphrased form from multiple sources. As shown in Figure 5, ratings were similar across these two conditions.

Figure 5

Mean Truth Ratings for Headlines by Repetition Status (Experiment 3)



Note. Ratings were provided on a scale from 1 = *Definitely False* to 6 = *Definitely True*. Each dot represents one participant ($N = 227$) with values horizontally shifted to represent the density distribution. Black diamonds reflect group means and error bars reflect standard errors of the mean.

A one-way ANOVA revealed a significant main effect of repetition type (new, repeated verbatim from one source, repeated paraphrased from three sources), $F(2, 452) = 18.13, p < .001$, $\eta_p^2 = 0.07$. Relative to new items ($M = 3.64, SD = 0.65$), participants gave higher truth ratings to headlines repeated verbatim from one source ($M = 3.81, SD = 0.63$), $t(226) = 4.91, p < .001$, 95% CI [0.10, 0.24], $d = 0.33$ and headlines repeated in paraphrased form from three sources ($M = 3.82, SD = 0.63$), $t(226) = 5.12, p < .001$, 95% CI [0.11, 0.26], $d = 0.34$. However, we did not observe a significant difference in ratings across the two kinds of repeated headlines, $t(226) = -0.45, p = .650$, 95% CI [-0.08, 0.05], $d = 0.03$.

As in Experiments 1 & 2, we followed up on this null result by conducting a Bayesian t -test comparing ratings for headlines in the two repeated conditions. The Bayes factor for this t -test was 12.16 in favor of the null hypothesis, suggesting our data are 12.16 times more likely under the null hypothesis. Again, we find moderate evidence that the two kinds of repetition we tested do not have different effects on belief.

Memory for Source Variability

Like in Experiment 2, we sought to verify that participants were able to remember whether there was social consensus around different claims. Thus, we conducted exploratory analyses on participants' responses to the memory check ("How many different people shared this headline or a similar headline that contained the same idea?" (0, 1, 2, or 3)). Again, while estimates were not very accurate, participants were able to qualitatively distinguish between the new, single-source and three-source headlines. Participants provided higher estimates for headlines that were repeated from a single source ($M = 1.84$, $SD = 0.58$) relative to new headlines ($M = 0.73$, $SD = 0.78$), $t(226) = 18.94$, $p < .001$, 95% CI [0.99, 1.23], $d = 1.26$. In addition, participants provided higher estimates for headlines that were repeated from three sources ($M = 2.11$, $SD = 0.55$) relative to headlines that were repeated from a single source, $t(226) = 6.61$, $p < .001$, 95% CI [0.19, 0.35], $d = 0.44$.

We again observed that while participants qualitatively distinguished between different kinds of headlines, memory was not particularly accurate. Thus, like in Experiment 2, we conducted an exploratory t -test comparing truth ratings for headlines repeated by one versus three sources among participants with the most accurate memories (the top third of participants in terms of the difference between estimated number of sources for one vs three source items; mean difference = 1.09, $SD = 0.32$; $N = 59$). As in our main analyses, we found no difference in

truth ratings for headlines repeated by one ($M = 3.83$, $SD = 0.58$) versus three ($M = 3.81$, $SD = 0.59$) sources among these participants, $t(58) = 0.24$, $p = .811$, 95% CI $[-0.11, 0.14]$, $d = 0.03$, and we observed moderate Bayesian evidence in favor of this null effect ($BF_{01} = 6.83$).

Perceptions of Social Consensus

We next conducted a series of exploratory analyses directly examining participants' estimates of the level of social consensus around different claims. Recall that we asked participants "What percent of social media users do you think would believe this headline to be true?" For "new" headlines, participants estimated that 56.2% of social media users would believe the headline. Participants thought more social media users would believe headlines repeated verbatim from one source (60.0%), $t(226) = 3.60$, $p < .001$, 95% CI $[1.72, 5.86]$, $d = 0.24$. However, participants were no more likely to think social media users would believe a headline when it was repeated in paraphrased form by three sources (59.9%) than when it was repeated verbatim from one source, $t(226) = -0.09$, $p = .924$, 95% CI $[-2.09, 1.89]$, $d = 0.01$.

These results suggest a disconnect between memory and perceptions of social consensus—even though participants can distinguish headlines repeated by multiple versus a single person, they do not attribute these differences to different levels of endorsement. One possible explanation is that participants simply thought that the sources in the exposure phase were not sharing information they believed. Indeed, when asked an open-ended question about the sources' main motive for the sharing the headlines, only 11.5% of participants spontaneously reported that the sources shared the headlines because they believed them or thought headlines were accurate.² However, this is not to say that participants thought the sources were attempting

² We had two research assistants code responses to the open-ended question "What do you think the main motivation was for people to share the headlines? In other words, why did these social media users decide to share the headlines you saw?" Responses were coded as 1 = accuracy or belief-based or 0 = other (Cohen's kappa = 0.834). All discrepancies were resolved by the first author. The full coding scheme is available at the OSF site.

to share information they disagreed with. When directly asked, participants reported thinking that the sources placed a moderate level of importance on sharing information they thought was accurate ($M = 5.35$, $SD = 2.79$; scale from 0 = not very important to 10 = very important).

Discussion

Again, consistent with a conceptual fluency account, repetition increased belief to a similar degree whether it came from distinct sources or a single source—even though the distinct sources were made to seem independent (i.e., sharing different versions of the same information).

Our exploratory measures shed some insight into the disconnect between our manipulation of consensus and belief. While participants could distinguish between news headlines repeated by one person versus many, this memory did not translate into perceptions about the level of social consensus around a claim. Instead, participants thought claims were more likely to be agreed upon simply due to repetition—regardless of the actual number of unique sources. Thus, people seem to represent consensus around repeated claims poorly.

General Discussion

Across three experiments, we find that variation in the wording or source of repeatedly-encountered news items does not moderate the effects of repetition on belief. These results contradict theories suggesting that repetition should be most influential when it reflects a consensus among varied sources. These findings also go against a perceptual fluency account, which predicts that verbatim repetition should increase belief more than paraphrased repetition. Instead, our results were most consistent with theories highlighting the role of low-level semantic cognitive processes like conceptual fluency in judgements of truth. Repetition increases belief by making statements conceptually easier to process, regardless of the precise wording or source.

Why did information about consensus not bear on participants' judgements of truth? We argued that, for this relationship to hold, participants must 1) track, in memory, how many unique sources repeated a statement 2) interpret this source variability as indicating different degrees of consensus and 3) use this consensus as a cue for truth. In Experiments 2 & 3, exploratory source memory data suggested the first process likely held—participants could track which claims were repeated by multiple sources (although not particularly well). Instead, exploratory data from Experiment 3 suggests participants did not reliably engage in the second process—inferring broad agreement around a claim based on their memory for the sources.

Interestingly, the news headlines that participants saw did affect perceptions of consensus: repeated headlines were judged as more widely agreed upon than new headlines. However, this effect occurred similarly whether the repetition came from one person or many. These findings mirror those of Weaver et al. (2007), who asked participants to estimate the prevalence of opinions they had seen once, three times from one source, or three times from three sources. Repetition made opinions seem more prevalent, even when only one source repeated the opinion. These results, along with ours, suggest that simply hearing information multiple times, even from a single source, can make it seem more widely accepted.

This insight also helps reconcile the present experiments with past work on how social consensus affects belief. In past work, information about consensus could affect judgements because it was directly presented (e.g., through summary statistics; Lewandowsky et al., 2013) or easily inferred (e.g., sequentially presented arguments; Connor Desai et al., 2022, Ransom et al., 2021). By contrast, in the present experiment, participants would have had to retroactively infer the level of consensus around the repeated claim they saw. Instead of effortfully tracking how

many different people shared each news item, participants may have relied on cues like fluency that imperfectly track consensus.

Regardless of the precise mechanism, our data speak to the way in which people consume news in our current world. Increasingly, people are turning to social media to stay informed, but social media platforms also lower the barriers for false or misleading information to spread (Bak-Coleman et al., 2021; Vosoughi et al., 2018). Concerns have already been raised about the consequences of the repeated exposure to misinformation online (e.g., Pennycook et al., 2018). Our work adds to this, showing that even a single repetitive social media user can increase others' belief in a claim—and they can do so just as well as a group of people.

Limitations & Future Directions

While consensus cues did not affect belief in our work, such cues might be more influential in other contexts where people are better able to track social consensus. For instance, while strangers shared the posts in these experiments, people may be more attentive to sources that are more personally salient (e.g., friends, coworkers, family) or are more clearly trustworthy or untrustworthy (e.g., news outlets, irreputable blogs). Finally, it may be easier to track consensus when statements are repeated more than three times. Thus, future work may examine the effects of consensus in other contexts.

Another limitation is that we do not know exactly how participants perceived the motives users had for sharing information. Exploratory measures suggested that participants thought users placed a moderate level of importance on sharing accurate information. However, we are unable to make more fine-grained interpretations about how participants interpreted different kinds of repetition. For instance, participants may have thought repetition by a single source indicated that the source had confidence in the headline, increasing belief despite a lack of

consensus. Future work should more closely examine the role of perceived sharing intentions in the effects of consensus.

Finally, our studies used one set of true headlines about health/science. Headlines were moderately plausible (truth ratings for new statements were just above the mid-point), so we expect our results would generalize, at least, to other moderately plausible statements. We also speculate that these results would hold for other topical domains, as the illusory truth effect replicates across a number of domains (like consumer product claims and political news; Hawkins et al., 2001; Pennycook et al., 2018), but this remains an open question.

Conclusion

Despite these limitations, our data provide valuable insights into how people form beliefs in real-world settings. In our digitized world, information moves quickly and reaches us repeatedly. Our results suggest that these repetitions can make information seem more true simply by making the information easier to process, even if that repetition comes from a single source. Across our three experiments, repetition of an idea increased belief—regardless of who said it or how it was said.

Acknowledgements

Our thanks to Leigh Farah for research assistance in selecting stimuli and preparing Experiments 1 and 2 in the Qualtrics and Gorilla platforms.

References

- Anwyl-Irvine, A. L., Massonnié, J., Flitton, A., Kirkham, N., & Evershed, J. K. (2020). Gorilla in our midst: An online behavioral experiment builder. *Behavior Research Methods*, 52(1), 388–407.
- Asch, S. E. (1956). Studies of independence and conformity: I. A minority of one against a unanimous majority. *Psychological Monographs: General and Applied*, 70(9), 1–70. <https://doi.org/10.1037/h0093718>
- Bak-Coleman, J. B., Alfano, M., Barfuss, W., Bergstrom, C. T., Centeno, M. A., Couzin, I. D., Donges, J. F., Galesic, M., Gersick, A. S., Jacquet, J., Kao, A. B., Moran, R. E., Romanczuk, P., Rubenstein, D. I., Tombak, K. J., Van Bavel, J. J., & Weber, E. U. (2021). Stewardship of global collective behavior. *Proceedings of the National Academy of Sciences*, 118(27), e2025764118. <https://doi.org/10.1073/pnas.2025764118>
- Connor Desai, S., Xie, B., & Hayes, B. K. (2022). Getting to the source of the illusion of consensus. *Cognition*, 223, 105023. <https://doi.org/10.1016/j.cognition.2022.105023>
- Connor, P., Varney, J., Keltner, D., & Chen, S. (2021). Social class competence stereotypes are amplified by socially signaled economic inequality. *Personality and Social Psychology Bulletin*, 47(1), 89–105.
- Dechêne, A., Stahl, C., Hansen, J., & Wänke, M. (2009). Mix me a list: Context moderates the truth effect and the mere-exposure effect. *Journal of Experimental Social Psychology*, 45(5), 1117–1122. <https://doi.org/10.1016/j.jesp.2009.06.019>
- 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 Psychology Review*, 14(2), 238–257. <https://doi.org/10.1177/1088868309352251>

- Faul, F., Erdfelder, E., Buchner, A., & Lang, A.-G. (2009). Statistical power analyses using G*Power 3.1: Tests for correlation and regression analyses. *Behavior Research Methods*, 41(4), 1149–1160. <https://doi.org/10.3758/BRM.41.4.1149>
- Fazio, L. K. (2020). Repetition increases perceived truth even for known falsehoods. *Collabra: Psychology*, 6(1), 38. <https://doi.org/10.1525/collabra.347>
- Festinger, L. (1954). A theory of social comparison processes. *Human Relations*, 7(2), 117–140.
- Hasher, L., Goldstein, D., & Toppino, T. (1977). Frequency and the conference of referential validity. *Journal of Verbal Learning and Verbal Behavior*, 16(1), 107–112. [https://doi.org/10.1016/S0022-5371\(77\)80012-1](https://doi.org/10.1016/S0022-5371(77)80012-1)
- Hawkins, S. A., Hoch, S. J., & Meyers-Levy, J. (2001). Low-involvement learning: Repetition and coherence in familiarity and belief. *Journal of Consumer Psychology*, 11(1), 1–11. https://doi.org/10.1207/S15327663JCP1101_1
- Held, L., & Ott, M. (2018). On p -values and bayes factors. *Annual Review of Statistics and Its Application*, 5(1), 393–419. <https://doi.org/10.1146/annurev-statistics-031017-100307>
- Henderson, E. L., Westwood, S. J., & Simons, D. J. (2021). A reproducible systematic map of research on the illusory truth effect. *Psychonomic Bulletin & Review*. <https://doi.org/10.3758/s13423-021-01995-w>
- Kassambara, A. (2020). *rstatix: Pipe-friendly framework for basic statistical tests* (0.7.0) [Computer software].
- Lewandowsky, S., Cook, J., Fay, N., & Gignac, G. E. (2019). Science by social media: Attitudes towards climate change are mediated by perceived social consensus. *Memory & Cognition*, 47(8), 1445–1456. <https://doi.org/10.3758/s13421-019-00948-y>

- Lewandowsky, S., Gignac, G. E., & Vaughan, S. (2013). The pivotal role of perceived scientific consensus in acceptance of science. *Nature Climate Change*, 3(4), 399–404.
<https://doi.org/10.1038/nclimate1720>
- Litman, L., Robinson, J., & Abberbock, T. (2017). TurkPrime. Com: A versatile crowdsourcing data acquisition platform for the behavioral sciences. *Behavior Research Methods*, 49(2), 433–442.
- Morey, R. D., Rouder, J. N., Jamil, T., & Morey, M. R. D. (2015). *Package ‘bayesfactor’* (0.9.12) [Computer software].
- Peer, E., David, R., Andrew, G., Zak, E., & Ekaterina, D. (2021). Data quality of platforms and panels for online behavioral research. *Behavior Research Methods*.
<https://doi.org/10.3758/s13428-021-01694-3>
- Pennycook, G., Cannon, T. D., & Rand, D. G. (2018). Prior exposure increases perceived accuracy of fake news. *Journal of Experimental Psychology: General*, 147(12), 1865–1880. <https://doi.org/10.1037/xge0000465>
- Pew Research Center. (2023). *News Platform Fact Sheet*.
<https://www.pewresearch.org/journalism/fact-sheet/news-platform-fact-sheet/>
- Pillai, R. M., & Fazio, L. K. (2024). *Repeated by many versus repeated by one: Examining the role of social consensus in the relationship between repetition and belief [OSF page]*.
<https://osf.io/z7bqy>
- Pillai, R. M., Yang, S., Jiang, Q., & Fazio, L. K. (under review). *Repetition increases belief more for trivia than for news headlines*.
- R Core Team. (2020). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing. <https://www.R-project.org/>

- Ransom, K. J., Perfors, A., & Stephens, R. (2021). Social meta-inference and the evidentiary value of consensus. In T. Fitch, C. Lamm, H. Leder, & K. Tessmar (Eds.), *Proceedings of the Annual Meeting of the Cognitive Science Society* (Vol. 43, Issue 43, pp. 833–839). Cognitive Science Society.
- Reber, R., & Schwarz, N. (1999). Effects of perceptual fluency on judgments of truth. *Consciousness and Cognition*, 8(3), 338–342. <https://doi.org/10.1006/ccog.1999.0386>
- Reyna, V. F., & Brainerd, C. J. (1995). Fuzzy-trace theory: An interim synthesis. *Learning and Individual Differences*, 7(1), 1–75. [https://doi.org/10.1016/1041-6080\(95\)90031-4](https://doi.org/10.1016/1041-6080(95)90031-4)
- Roggeveen, A. L., & Johar, G. V. (2002). Perceived source variability versus familiarity: Testing competing explanations for the truth effect. *Journal of Consumer Psychology*, 12(2), 81–91. https://doi.org/10.1207/S15327663JCP1202_02
- Silva, R. R., Garcia-Marques, T., & Reber, R. (2017). The informative value of type of repetition: Perceptual and conceptual fluency influences on judgments of truth. *Consciousness and Cognition*, 51, 53–67. <https://doi.org/10.1016/j.concog.2017.02.016>
- Simmonds, B. P., Stephens, R., Searston, R. A., Asad, N., & Ransom, K. J. (2023). The Influence of Cues to Consensus Quantity and Quality on Belief in Health Claims. In M. Goldwater, F. K. Anggoro, B. K. Hayes, & D. C. Ong (Eds.), *Proceedings of the Annual Meeting of the Cognitive Science Society* (Vol. 45, Issue 45).
- Unkelbach, C., & Greifeneder, R. (2013). A general model of fluency effects in judgment and decision making. In C. Unkelbach & R. Greifeneder, *The experience of thinking: How the fluency of mental processes influences cognition and behaviour* (pp. 11–32). Psychology Press.

- Unkelbach, C., Koch, A., Silva, R. R., & Garcia-Marques, T. (2019). Truth by repetition: Explanations and implications. *Current Directions in Psychological Science*, 28(3), 247–253. <https://doi.org/10.1177/0963721419827854>
- Unkelbach, C., & Rom, S. C. (2017). A referential theory of the repetition-induced truth effect. *Cognition*, 160, 110–126. <https://doi.org/10.1016/j.cognition.2016.12.016>
- Vogel, T., Silva, R. R., Thomas, A., & Wänke, M. (2020). Truth is in the mind, but beauty is in the eye: Fluency effects are moderated by a match between fluency source and judgment dimension. *Journal of Experimental Psychology: General*, 149(8), 1587–1596. <https://doi.org/10.1037/xge0000731>
- Vosoughi, S., Roy, D., & Aral, S. (2018). The spread of true and false news online. *Science*, 359(6380), 1146–1151. <https://doi.org/10.1126/science.aap9559>
- 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. <https://doi.org/10.1037/0022-3514.92.5.821>
- Whalen, A., Griffiths, T. L., & Buchsbaum, D. (2018). Sensitivity to shared information in social learning. *Cognitive Science*, 42(1), 168–187. <https://doi.org/10.1111/cogs.12485>
- Yousif, S. R., Aboody, R., & Keil, F. C. (2019). The illusion of consensus: A failure to distinguish between true and false consensus. *Psychological Science*, 30(8), 1195–1204. <https://doi.org/10.1177/0956797619856844>