

Some misinformation is more easily countered: An experiment on the continued influence effect

Saoirse Connor Desai (saoirse.connor-desai@city.ac.uk) & Stian Reimers (stian.reimers.1@city.ac.uk)

City, University of London, Department of Psychology, Whiskin Street

London, EC1R 0JD

Abstract

Information initially presented as a likely cause of an event but turns out to be incorrect can affect people's reasoning despite being clearly corrected – a phenomenon known as the *continued influence effect* of misinformation. The present work extends previous findings showing that misinformation that implies a likely cause of an adverse outcome is more resistant to correction than misinformation that explicitly states a likely cause. Participants either read a report describing a fire or a crash. The difference between implied and explicitly stated misinformation was replicated with the fire scenario, which has been commonly used in continued influence research. There was little evidence of a continued influence of misinformation for the (novel) crash scenario. The results constrain the generalizability of the continued influence effect and suggest that corrections that clearly invalidate initial misinformation can be effective.

Keywords: Misinformation; Continued Influence; Correction; Reasoning; Inference; Memory

Introduction

Misinformation often has a lasting effect on people's judgments and decisions despite being unequivocally retracted. A prime example of this is the widespread belief in the discredited claim that the MMR vaccination causes autism (Horne, Powell, Hummel, & Holyoak, 2015). The harmful effects of misinformation and ineffectiveness of attempts to correct mistaken beliefs have become a great concern for contemporary society (see Cook, Ecker, & Lewandowsky, 2015; Lewandowsky, Ecker, & Cook, 2017 for recent discussions).

Decades of laboratory research have shown that corrections often fail to eliminate the effects of misinformation (see Lewandowsky, Ecker, Seifert, Schwarz, & Cook, 2012 for review). The *continued influence effect* of misinformation refers to the consistent finding that information initially presented as true but later shown to be false¹ continues to influence beliefs and reasoning despite clear and credible corrections (e.g., Ecker, Lewandowsky, Swire, & Chang, 2011; Ecker, Lewandowsky, Swire, et al., 2011; Johnson & Seifert, 1994; Rich & Zaragoza, 2016).

In a typical experimental task, participants read a fictional account of an unfolding event containing a series of individually presented statements (e.g., a fire at a warehouse). A piece of target (mis)information that provides a likely cause for the outcome of the event (e.g., carelessly stored oil paint and gas cylinders were on the

premises), is presented but later corrected (i.e., the initial information is identified as being incorrect), or remains uncorrected. After reading all the statements, participants' inferential reasoning and verbatim memory for the story are assessed through a series of open-ended questions. Responses to inference questions are coded according to whether they are consistent with the explanation implied by the target (mis)information (e.g., "exploding gas cylinders"), or not (e.g., "faulty wiring"). The sequential nature of the experimental task (i.e., serial presentation of misinformation) resembles rolling news coverage.

The typical finding from continued influence studies is that misinformation continues to influence people's inferential reasoning even though they clearly understand and remember that the information was corrected (Johnson & Seifert, 1994), even when given prior warnings about the persistence of misinformation (Ecker et al., 2010). Although the correction usually does have some impact, reliance on misinformation is typically halved compared to uncorrected control group, the misinformation is still referred to significantly. The fact that corrections are often ineffective at 'removing' misinformation from people's understanding of events, emphasizes the need to identify factors that contribute to the *continued influence effect*.

There are two leading cognitive explanations for the persistence of misinformation following a correction. According to the selective retrieval account, the *continued influence effect* occurs when both the correct (i.e., the correction) and incorrect (i.e., misinformation) are stored in memory simultaneously, and misinformation is activated but insufficiently suppressed (Ecker, Lewandowsky, Swire, et al., 2011). Conversely, the model updating account maintains that people construct a mental event model 'on the fly' that is continually updated when new information becomes available. Invalidating a central piece of information (i.e., a likely cause of the event) leaves people with a gap in their model. People prefer a coherent but incorrect model to a correct but incomplete one and therefore maintain the invalidated information (Ecker, Lewandowsky, & Apai, 2011; Johnson & Seifert, 1994). Following this line of reasoning, maintaining invalid information is sensible; an incomplete model has no inferential power (i.e., it is a busted flush) whereas a complete – but erroneous – model at least allows inferences to be made (see Mercier & Sperber, 2017 for a related view on the purpose of reasons).

Support for the model updating account comes from the well-established finding that providing an alternative explanation for the outcome of the event (e.g., arson

¹ The term *misinformation* is used throughout to remain consistent with the literature on this topic.

materials were found in the warehouse) can help people revise and update their initial mistaken account of what happened by replacing the ‘gap’ left by invalidating misinformation (e.g., Ecker, Lewandowsky, & Apai, 2011; Ecker et al., 2010; Johnson & Seifert, 1994; Rich & Zaragoza, 2016). Identifying factors that promote the continuing influence of misinformation can contribute to both to scientific theory and public policy.

Implied vs. Explicitly Stated Misinformation

One factor that has recently been shown to contribute to the *continued influence effect* is whether misinformation explicitly states or merely implies the cause of an adverse outcome. Rich and Zaragoza (2016) gave participants a report describing a theft of valuable jewelry from a couple’s home while they were on vacation. Participants either initially learned that the couple had asked their son to check on the property while they were away or that police suspected the couple’s son had taken the jewelry from the house. In the former case the initial misinformation implied the son’s involvement allowing participants to infer the cause of the theft. In the latter case the cause was explicitly stated. Later in the story, participants in the correction condition learned that the son had actually been out of town the theft occurred thereby invalidating the initial misinformation.

When misinformation remained uncorrected, participants made a similar number of references to implied and explicitly stated misinformation. Although a correction reduced reliance on misinformation, the correction was more effective following explicit misinformation. There was also a larger effect of correction for explicit misinformation when the correction was paired with an alternative explanation informing participants that the actual thief had been caught.

Rich and Zaragoza (*ibid.*) argue that a likely explanation for the findings is that participants who received implied misinformation had to go beyond the available information and *infer* a likely cause of the outcome because causal relations between elements of the story were not explicitly stated. Previous research has shown that readers generate inferences between elements in the story when causal relations between pieces of information are not explicitly stated (e.g., Myers, Shinjo, & Duffy, 1987; Pennington & Hastie, 1988). There is also some evidence that evidential discrediting is less effective when people self-generate explanations than when explanations are externally provided (e.g., Davies, 1997).

Rich and Zaragoza acknowledged the limitations of the findings because they were only obtained with a single news story. Story content may interact with an individuals’ pre-existing knowledge and beliefs moderating the effects of implied and explicitly stated misinformation. For example, Ecker, Lewandowsky, Fenton, and Martin (2014) found that participants’ pre-existing attitudes (racial prejudice toward an ethnic minority group) influenced how they used race related information – although not processing of a

correction. For these reasons it is important to replicate these findings with different stories in order to establish the boundary conditions of the continued influence effect. Accordingly, this paper’s aim is to replicate the finding that implicit misinformation is more resistant to correction than explicit misinformation with two different news stories (event reports).

The present study used a ‘rolling news’ format to situate the new stories in a familiar context. Rolling news coverage is just one medium of information dissemination which can proliferate the spread of misinformation. The format of rolling news reporting aims to deliver developments in news stories in real-time. This can result in piecemeal reporting. As a result, news can be based on incomplete, mistaken, or inaccurate information. Misinformation propagated in quickly through live TV and internet coverage is not corrected as quickly as it spreads (e.g., Starbird, Maddock, Orand, Achterman, & Mason, 2014).

The structure of the event reports used here differed from previous continued influence studies. The story in Rich and Zaragoza’s study included additional information which could be interpreted as diagnostic of cause implied or explicitly stated by initial misinformation (e.g., Police are still attempting to determine whether other valuables are missing from the home. The television and home computer, however, had not been disturbed) that the son broke into the house. If the story includes information that lends credence to misinformation this might increase the perceived veracity of the misinformation and decrease the perceived veracity of the correction. The fact that the additional event information included in the present stories was designed to be non-diagnostic of misinformation could reduce reliance on misinformation relative to Rich and Zaragoza’s results. In addition to this change, the statement provided immediately before initial misinformation included background or ‘base rate’ information about common causes of the outcome described in the story. This information was included in order to increase the availability of other possible causes of the outcome. If story content is related to misinformation type, then the effectiveness of a correction to implied or explicit misinformation will differ between event narratives.

Method

Participants

In total 168 were recruited but 53 participants who answered any one of 3 instruction attention check questions incorrectly were excluded prior to analysis. One-hundred and fifteen participants recruited from Amazon Mechanical Turk (67 female; age 38.06 ± 11.25) were included in the final analysis. Participants were paid \$1.50 and took an average of 18 minutes to complete the experiment. In addition to the standard reward, participants were given the opportunity to earn an additional \$10 based on high accuracy scores across instruction check and fact recall questions.

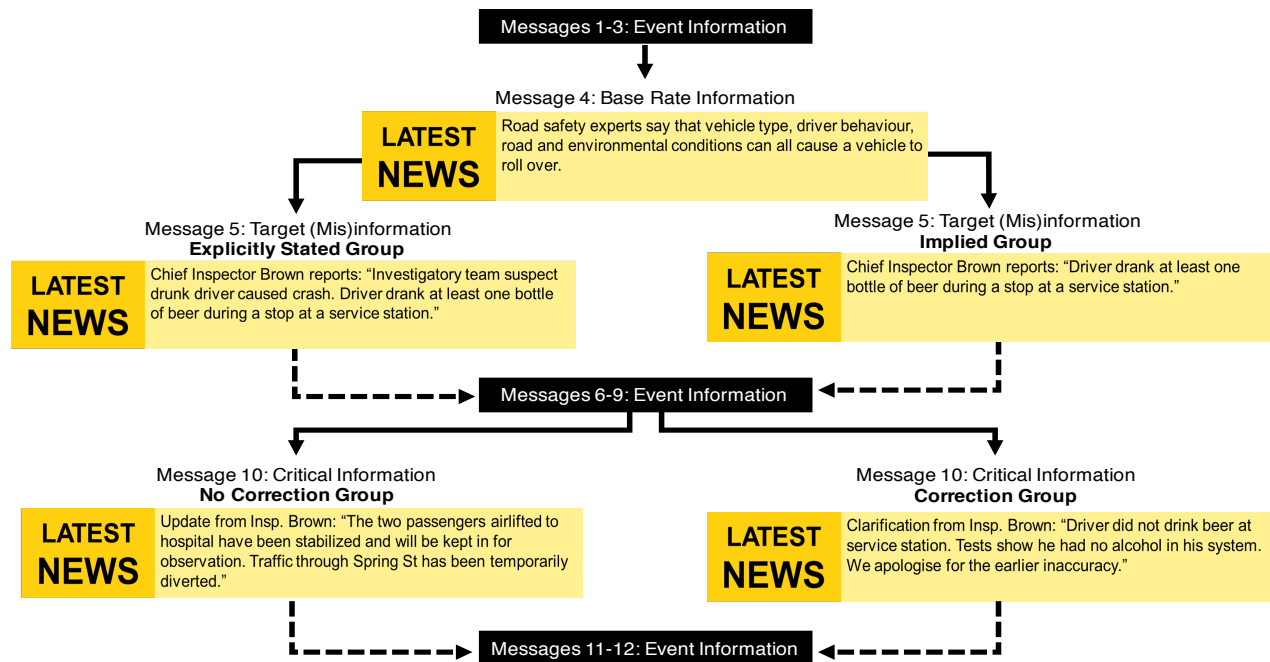


Figure 1 Schematic diagram of the version of *continued influence effect* task used in the present study. Messages relate to the ‘van crash’ narrative. Messages 1-3 provide general information about the event beginning with the van accident being reported. Message 4 provides (base rate) information about common causes of vehicular accidents. Target (mis)information is presented at Message 5 and is then corrected for correction group at Message 10. The ‘warehouse fire’ narrative followed the same structure.

Stimuli & Design

The stimuli were generated in Qualtrics (Qualtrics, Provo, UT). Participants read one of 8 versions of a fictional news report that either described a fire at a warehouse or a van crash, each consisting of 12 discrete messages. The ‘fire’ report was a modified version of stimuli used in previous research (e.g., (Ecker et al., 2010; Johnson & Seifert, 1994). The ‘crash’ report was a new story. The two event reports were constructed to be structurally similar but superficially distinct.

Fig. 1 illustrates how message content was varied across experimental conditions, as well as the message presentation format. The effect of correction information (No Correction, Correction), Event Report (Fire, Crash), and Misinformation (Explicitly Stated, Implied) on reference to target (mis)information was assessed between groups; participants were randomly assigned to one of the 8 experimental groups.

The messages were presented in as ‘latest news’ in a rolling news format that originated from the same fictional news outlet; each message was no longer than 280 characters. Messages in the same position within the sequence were matched for length across event narratives. Each message appeared one at a time for a minimum of 5 seconds each; there was no maximum reading time. Participants clicked a button to proceed to the next message; they were unable to return and view previous messages.

The misinformation in the fire report (Message 5) either implied (Fire Chief Lucas issues statement: “Cans of oil paint and pressurized gas cylinders were present in storeroom before fire”) or explicitly stated (Fire Chief Lucas issues statement: “Investigation team suspect fire caused by carelessly stored flammable liquids. Cans of oil paint and pressurized gas cylinders were present in storeroom before fire”). Message 10 varied depending on condition and either corrected the target (mis)information (e.g., Correction from Chief Lucas: No flammable items actually in storeroom. No paint or gas had ever been present in the warehouse. We apologize for our earlier error), or remained uncorrected (e.g., Update from Chief Lucas: The warehouse employees taken to hospital were treated for smoke inhalation and have been released. Temporary accommodation is available for evacuated residents). Additional information in the report gave further details about the event that could not be interpreted as evidence in support of the cause implied by misinformation (e.g., Three warehouse workers working overtime have been taken to hospital to be treated for smoke inhalation). The only exception was Message 8 which could be interpreted as an effect of the cause implied or explicitly stated by the misinformation (e.g., thick, oily smoke + sheets of flames hinder firefighter’s efforts, intense heat has made the fire difficult to bring under control). The crash report followed the same format and the content of the messages can be seen in Fig. 1.

After reading all of the messages in the report, participants completed a questionnaire consisting of 7

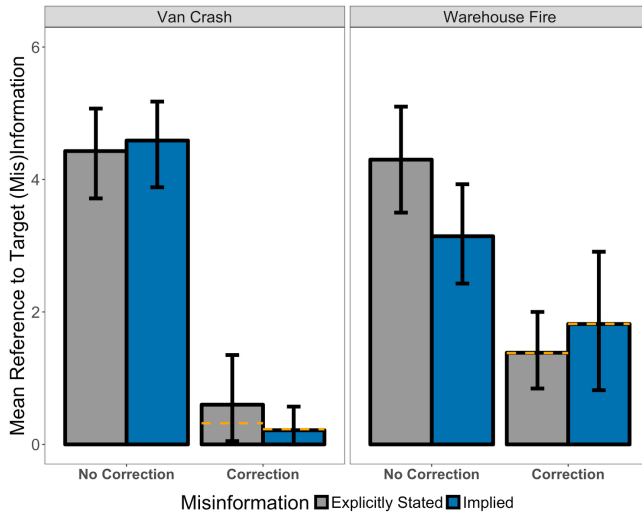


Figure 2 Mean reference to target (mis)information as a function of event report, correction information, and misinformation. Error bars represent 95% confidence interval of the mean. Bracket represent significant comparisons between correction information conditions. Dashed lines represent means after excluding participants who did not recall the correction.

inference questions, 7 fact recall questions and 2 correction recall questions (16 questions in total). The inference and fact recall question blocks were intermixed and presented in a random order; except the question relating to the cause of the event, which always came last. Inference questions probed participants' causal understanding of the news report (e.g., "Is there any evidence of careless management in relation to this fire?"), and included a question querying what participants thought the most likely cause of the fire was. Fact recall questions enquired about the event details that were consistent across misinformation and correction information conditions (e.g., "Which hospital were the workers taken to?"). Two further questions assessed recall and understanding of the correction message – this question was of course only relevant to participants in conditions featuring a correction. Participants typed a response to each of 16 questions in a text box, were required to use a minimum of 25 characters, and encouraged to answer using full sentences.

Results

Coding of Responses

The main dependent variable extracted from responses to inference questions was 'reference to target (mis)information'. References that explicitly stated, or strongly implied, that the target (mis)information caused, or contributed, to the outcome of the event were scored a 1 or were otherwise scored as 0. Table 1 shows examples of responses that was coded as a reference to target (mis)information and an example of a response that was not

coded as such. References to flammable materials which did not specifically mention storage (e.g., "It could have been avoided by keeping flammable objects or items in place") were not treated as references to misinformation because there was no specific mention of gas, paint, liquids, substances, chemicals, or the fact they were (allegedly) kept in the storeroom. Similarly, references to driver behavior that did not mention intoxication or drunkenness were not counted as references to misinformation (e.g., "by having him be more alert drinking coffee").

The maximum individual score for inference questions was 7. Responses to factual questions were scored for accuracy; correct or partially correct responses were scored 1 and incorrect responses were scored 0. The maximum factual score was 7. Correction recall scores were computed using the same criteria; the maximum individual correction recall score was 2.

Inter-Coder Reliability Responses were coded by a trained coder. A second, independent judge then coded 10% of responses from each narrative. Inter-rater agreement was 0.95 and Cohen's $K = 0.89 \pm 0.03$, indicating a high level of agreement between coders, both of which are higher than the benchmark values of 0.7 and 0.6 (Krippendorff, 2012; Landis, & Koch, 1977), and there was no systematic bias between raters, $\chi^2 = 1.92, p = .17$.

Table 1 Example of response coding criteria

Event Narrative	Inference Question	Example of Response Scored 1	Example of Response Scored 0
Van Accident	How could this accident have been avoided?	If the driver had not been drinking.	He was in a court battle with his ex-wife.
Warehouse Fire	How could the fire at the warehouse have been avoided?	They should store flammable substances separately.	Pay more attention to the signs and smells of a fire. Not overworking workers.

Inferences

Reference to target (mis)information on inference questions were subjected to a three-way factorial ANOVA. Fig. 2 shows mean references to target (mis)information across conditions. The ANOVA yielded a significant third order interaction, $F(1, 107) = 4.15, p = .04, \eta_p^2 = .04$ [.00, .13]. Simple second order interaction effects revealed an interaction between correction information and misinformation directness for the fire report, $F(1, 45) = 4.29, p = .04, \eta_p^2 = 0.09$ [.00, .26], but not for the crash report, $F(1, 62) = 0.32, p = .58, \eta_p^2 = .01$ [.00, .09]. This was followed up with simple main effects analysis which showed that, for the fire report, there was a significant effect of correction information for explicitly stated misinformation, $F(1, 22) = 31.33, p < .001, \eta_p^2 = .59$ [.27,

.73], but not implied misinformation, $F(1, 23) = 3.80, p = .06, \eta_p^2 = .14$ [.00, .39].

This means that implied misinformation was more resistant to correction than explicitly stated misinformation, but only for the people who read the fire report. There were also a greater number of references to uncorrected (mis)information when the cause was explicitly stated than when it was implied. In contrast, a correction to both implied and explicitly stated misinformation led to a robust reduction in reliance on misinformation for the crash narrative. This suggests that story content does interact with misinformation type (i.e., whether misinformation explicitly states or implies a likely cause of the outcome). Moreover, these results also suggest that some misinformation is more easily updated or that corrections are more effective for some misinformation than for others. However, given that there were as few as 11 participants in one cell, and the fact that the effect size for the third order interaction was small, caution should be exercised when generalizing from these results².

Fact Recall

Manipulations of event report, misinformation type, or correction information were not expected to affect recall. Mean fact recall scores ranged from 4.15 to 6.18 (out of 7). Contrary to expectations, however, the ANOVA revealed a significant three-way interaction, $F(1, 107) = 8.42, p = .005, \eta_p^2 = .07$ [.01, .18]. Further examination of the interaction, again, showed a significant interaction between misinformation and correction information, $F(1, 45) = 6.29, p = .02, \eta_p^2 = .12$ [.00, .30]. When misinformation was explicitly stated the group featuring a correction recalled significantly fewer facts than the uncorrected group, $F(1, 22) = 8.31, p = .009, \eta_p^2 = .27$ [.02, .51]. The reason for the difference between these groups is not entirely clear but could also be the reason that a correction was more effective in the explicit condition in the warehouse fire narrative.

Correction Recall

Correction recall scores were analysed in to confirm whether ability to recall accurately differed between event narratives or misinformation conditions. The no correction groups were not included in the analysis as their responses were not informative. There was no interaction between misinformation and event report for correction recall scores, $F(1, 52) = 0.09, p = .76, \eta_p^2 = .00$ [.00, .08]; nor were there main effects of misinformation, $F(1, 52) = 3.44, p = .07, \eta_p^2 = .06$ [.00, .21], or event reports, $F(1, 52) = 0.10, p = .75, \eta_p^2 = .00$ [.00, .08]. Mean correction recall scores ranged between 1.62 – 1.89 (out of 2), thus indicating good overall recall of the correction message.

Discussion

Misinformation presented in news stories often influences beliefs and reasoning about described events even when refuted. The present work extends on previous findings showing that misinformation that implies a likely cause of an adverse outcome is more resistant to correction than misinformation that explicitly states a likely cause. This is thought to occur because people's mental models of events are more elaborate when they have to infer causal relations between pieces of information in the story than when causal relations are explicitly provided.

One limitation of prior work that has compared implied and explicitly stated misinformation is that findings were only obtained for a single news story (e.g., Rich & Zaragoza, 2017). The present study tested the generalizability of this finding with two different stories describing events, presented as rolling news reports. One of the event reports has been commonly used in previous continued influence studies (e.g., Johnson & Seifert, 1994) whereas the other was a newly developed story. The event report used in the present work differed to those used in previous studies along two main dimensions. First, the event reports were structured in such that they did not push people to make misinformation related inferences. More specifically, most of the additional event information did was neutral with respect to the likely cause implied or explicitly stated by the initial misinformation (i.e., the information provided not evidence in support of the misinformation). Second, information presented immediately before misinformation gave information about other common causes of the adverse outcome (e.g., common causes of crashes are road conditions or common causes of industrial fires are electrical problems). This piece of information was included to encourage participants to think of other likely causes of the outcome described in the report.

The results of the present study do indeed suggest that differences between implied and explicit misinformation do interact with the features of the described event. In the crash scenario corrections to both implied and explicitly stated misinformation led to robust reductions in reference to misinformation. Reference to misinformation was close to zero for both implied and explicitly stated misinformation conditions. This finding is clearly not due to the fact that participants refer to this information less overall, as participants referred more to the uncorrected misinformation in the crash than the fire scenarios.

Conversely, in the warehouse fire narrative implied misinformation was more resistant to correction than explicitly stated misinformation, consistent with Rich and Zaragoza's (2016) findings. The present study also found that participants referred to explicitly stated (mis)information more than implied misinformation when it was uncorrected.

The results of the present work provide evidence that story content not only interacts with explicitly stated and implied misinformation, but also with a corrections' effectiveness. One reason that a correction was substantially

² The difference between implied and explicitly stated misinformation was only observed when analysing data from the restricted sample of participants who did not fail the instructional attention check questions. This could be due to the fact that participants who failed the attention checks were also less likely to properly encode the information included in the story and would therefore be unaffected by the manipulations.

more effective at reducing reference to misinformation in crash than the fire event report could be related to the type of the correction in each scenario. The crash correction gave clear evidence that the driver's drunkenness was not the cause of the crash (i.e., that a test showed the driver had no alcohol in system), thereby eliminating the possibility that the driver's drunkenness caused the crash. In contrast, the correction in the fire report corrected the earlier statement (i.e., that flammable substances were in the storeroom before the fire) without providing any evidence of this, merely indicating the fire did not occur in the way that was initially stated or implied.

One possible reason that the continued influence effect was not replicated for the novel scenario is that it is less ambiguous than scenarios commonly used in continued influence studies. Although the correction used in Rich and Zaragoza's study appeared to give clear evidence in contradiction of the misinformation (i.e., the son was out of town and therefore could not have committed the crime), participants could still have drawn an inference about the son's involvement in the theft. For example, he could have given someone the keys and told them where the jewelry was kept. In contrast, in the crash scenario once participants know the driver had no alcohol in his system they are forced to update their inference and have to rely on other possible causes of the crash, or indeed, conclude that there is insufficient evidence to infer a clear cause.

It is not entirely clear why explicitly stated misinformation was referred to more when it was uncorrected in the fire report but referred to less after a correction. One possibility is that the information about oil paint and gas cylinders was made more relevant by stating that police suspected that this was the cause of the fire and therefore was more available to participants when answering questions. In contrast, when the information was corrected the explicit statement of the cause made the discrepancy between the correction and misinformation more apparent and therefore participants were better able to revise their model (see Ecker, Hogan, & Lewandowsky, 2017 for a similar point).

Further investigation is necessary to draw firm conclusions about the continued influence effect did not appear for the crash story. Given the relatively low sample size in some conditions following exclusion of participants who answered instruction check questions incorrectly, the overall power of the study may have been low. Despite these limitations, the present study's results highlight the importance of testing the boundary conditions of the continued influence effect. Furthermore, these results constrain the generalizability of continued influence findings, more generally, and show that the effect is not guaranteed to emerge under all conditions.

References

Cook, J., Ecker, U., & Lewandowsky, S. (2015). Misinformation and how to correct it. *Emerging Trends in the Social and Behavioral Sciences: An Interdisciplinary, Searchable, and Linkable Resource*.

- Davies, M. F. (1997). Belief persistence after evidential discrediting: The impact of generated versus provided explanations on the likelihood of discredited outcomes. *Journal of Experimental Social Psychology*, 33(6), 561–578.
- Ecker, U. K. H., Lewandowsky, S., & Apai, J. (2011). Terrorists brought down the plane!—No, actually it was a technical fault: processing corrections of emotive information. *Quarterly Journal of Experimental Psychology* (2006), 64(2), 283–310. <https://doi.org/10.1080/17470218.2010.497927>
- Ecker, U. K. H., Lewandowsky, S., Swire, B., & Chang, D. (2011). Correcting false information in memory: manipulating the strength of misinformation encoding and its retraction. *Psychonomic Bulletin & Review*, 18(3), 570–578. <https://doi.org/10.3758/s13423-011-0065-1>
- Ecker, U. K. H., Lewandowsky, S., & Tang, D. T. W. (2010). Explicit warnings reduce but do not eliminate the continued influence of misinformation. *Memory & Cognition*, 38(8), 1087–1100. <https://doi.org/10.3758/mc.38.8.1087>
- Ecker, U. K., Hogan, J. L., & Lewandowsky, S. (2017). Reminders and Repetition of Misinformation: Helping or Hindering Its Retraction? *Journal of Applied Research in Memory and Cognition*.
- Ecker, U. K., Lewandowsky, S., Fenton, O., & Martin, K. (2014). Do people keep believing because they want to? Preexisting attitudes and the continued influence of misinformation. *Memory & Cognition*, 42(2), 292–304.
- Horne, Z., Powell, D., Hummel, J. E., & Holyoak, K. J. (2015). Countering antivaccination attitudes. *Proceedings of the National Academy of Sciences*, 112(33), 10321–10324.
- Johnson, H. M., & Seifert, C. M. (1994). Sources of the Continued Influence Effect: When Misinformation in Memory Affects Later Inferences. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 20(6), 1420–1436. <https://doi.org/10.1037/0278-7393.20.6.1420>
- Krippendorff, K. (2012). *Content analysis: An introduction to its methodology*. SAGE Publications.
- Landis, J. R., & Koch, G. G. (1977). The measurement of observer agreement for categorical data. *Biometrics*, 159–174.
- Lewandowsky, S., Ecker, U. K. H., Seifert, C. M., Schwarz, N., & Cook, J. (2012). Misinformation and Its Correction: Continued Influence and Successful Debiasing. *Psychological Science in the Public Interest*, 13(3), 106–131. <https://doi.org/10.1177/1529100612451018>
- Lewandowsky, S., Ecker, U. K., & Cook, J. (2017). Beyond Misinformation: Understanding and Coping with the “Post-Truth” Era. *Journal of Applied Research in Memory and Cognition*, 6(4), 353–369.
- Mercier, H., & Sperber, D. (2017). *The enigma of reason*. Harvard University Press.
- Myers, J. L., Shinjo, M., & Duffy, S. A. (1987). Degree of causal relatedness and memory. *Journal of Memory and Language*, 26(4), 453–465.
- Pennington, N., & Hastie, R. (1988). Explanation-based decision making: Effects of memory structure on judgment. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 14(3), 521.
- Rich, P. R., & Zaragoza, M. S. (2016). The continued influence of implied and explicitly stated misinformation in news reports. *Journal of Experimental Psychology: Learning, Memory & Cognition*, 42(1), 62–74. <http://dx.doi.org/10.1037/xlm0000155>
- Starbird, K., Maddock, J., Orand, M., Achterman, P., & Mason, R. M. (2014). Rumors, false flags, and digital vigilantes: Misinformation on twitter after the 2013 boston marathon bombing. *ICConference 2014 Proceedings*.