

Supplemental Information for

How are Competitive Framing Environments Transformed by Person-to-Person

Communication?

An Integrated Social Transmission, Content Analysis, and Eye Movement Monitoring Approach

Study 1

The original news stories had an average of 69.2 words for the context section. The context section of the reproductions decreased across the waves ($B = -5.97$, $SE = 1.13$, $p < .001$)¹ and contained an average of 34.08 words ($SD = 12.15$) in wave 1, 26.16 words ($SD = 7.50$) in wave 2, and 22.14 words ($SD = 6.30$) in wave 3.

We also estimated a mixed-effects regression model in which we treated the number of words in the context section that people were exposed to as a fixed effect and our primary independent variable.² Our dependent variable was the extent to which a frame a participant was exposed to appeared in his or her memory-based reproduction of the news story (0 = did not appear, 1 = did appear). A negative and significant effect of word count ($B = -.02$, $SE = .007$, $p < .01$) suggests that as the number of words increased for the context section, people were less likely to remember the frame. We estimated another model in which we included wave as a covariate to account for other message features that may have changed across the waves. The coefficient for the word count for the context section does not reach conventional levels of statistical significance with the inclusion of this variable ($B = -.01$, $SE = .01$, $p = .39$).

Investigating an Alternative Explanation

As mentioned in the main paper, as part of the battery of surveys participants completed after the reproduction task, participants were shown all the original versions of the news stories and frames (even participants in waves 2 and 3 were shown the original frames). Participants were then asked to rate the effectiveness of the argument conveyed in the frame (1 = Definitely not effective, 6 = Definitely effective). To determine whether frames that were ideologically incongruent with the participants' general ideological beliefs were rated as less effective than an

¹ This was estimated using a regression model with wave as an independent variable. Word count of the context reproduction was used a dependent variable.

² Results in the supplementary information are based on mixed-effects regression models.

ideologically congruent frames, we estimated a mixed-effects regression model in which ideological-congruency was used an independent variable (participant ideology congruent with frame's ideological association = 1, participant ideology incongruent with frame's ideological association = 0). We then used participant's effectiveness ratings of the frames as a dependent variable. A positive and significant coefficient ($B = 1.24$, $SE = .11$, $p < .001$) suggests that ideologically-congruent frames were rated as more effective than ideologically-incongruent frames. This suggests that participant's individual ratings of a frame's reflect participants' general pre-existing attitudes.

We estimated an interaction model in which we modeled the norming ratings of a frame's associated ideology (higher values indicate more conservative ideology), the participant's self-reported ideology (higher values indicate more conservative ideology), and the interaction between the two as independent variables. Our dependent variable was whether a frame a participant was exposed to appeared in his or her reproduction of the news story (0 = did not appear, 1 = did appear). A positive coefficient for the interaction would suggest that as ideological congruency between the frames and participants increased, the ability to remember frames increased – an outcome consistent with the selective transmission account. In contrast, a negative coefficient for the interaction would suggest that as ideological congruency between the frames and participants increased, the ability to remember frames decreased – an outcome *inconsistent* with the selective transmission account. We found a negative and significant interaction ($B = -0.19$, $SE = .06$, $p < .001$), a result inconsistent with the selective transmission account.

Study 2

In study 2, we also examined whether attention to other parts of a message (the context

section) decreases people's ability to remember frames. A typical news story contains several relevant parts (contextual information about a political event, information about frames). Given the limitation that people cannot visually attend to all pieces of information at the same time, people need to make strategic decisions about what pieces of information they should direct their attention to. This competition for visual attention can have important consequences. Focusing on other parts of a message that do not contain frames (contextual descriptions of the event) can decrease attention to the frames. Given the strong link between attention and memory (Loftus, 1972; Neuschatz, Lampinen, Preston, Hawkins, & Toggia, 2002; Pertzov, Avidan, & Zohary, 2009), this decrease in visual attention to the frames may partly explain why people may not remember them. Specifically, we propose the following hypothesis:

H₃: As visual attention to contextual information unconnected to frames increases (and attention to frames decreases), people will be less likely to remember frames.

This hypothesis is important to test because framing studies typically assume that individuals pay attention to frames. This foundational assumption is often not directly tested. One of our contributions to the scholarly literature on framing is to show how eye movement monitoring technology can be used to test this assumption.

Results

One goal for study 2 was to test Hypothesis 3: As visual attention to contextual information unconnected to frames increases (and attention to frames decreases), people will be less likely to remember frames. We defined regions of interest around the sentences that provided background information about the event described in the news story (context region of interest) and regions of interest around the sentences that contained the frames (frames region of interest). We then calculated the proportion of time individuals directed their gaze to the context region of interest by dividing the amount of time individuals directed their gaze to the context

region of interest by the total amount of time they directed their gaze at both the context and frame regions of interest combined.

We estimated a mixed-effects model in which we treated the proportion of time individuals directed their gaze at the context section as a fixed effect and our primary independent variable. Our dependent variable was the total number of frames that appeared in their memory-based reproduction of the news story (zero, one, or two). A negative and significant effect of gaze directed at the context section ($B = -1.95$, $SE = 0.39$, $p < .001$) suggests that as individuals directed greater attention to the context section (at the expense of attention to the frames), people were less likely to remember frames – a finding consistent with Hypothesis 3.

Discussion

In study 2, we further explored the mechanisms underlying why people may forget frames. We reasoned that individuals may direct their visual attention to other equally important parts of a message (context information) at the cost of directing less attention to frames. A decrease in visual attention to frames could lead to a weaker or non-existent memory representation that results in people forgetting them. Indeed, we found that greater attention to the other parts of the message (and decreased attention to the section containing the frames) decreased people's ability to remember the frames.

References

- Loftus, G. R. (1972). Eye fixations and recognition memory for pictures. *Cognitive Psychology*, 3(4), 525-551.
- Neuschatz, J. S., Lampinen, J. M., Preston, E. L., Hawkins, E. R., & Toggia, M. P. (2002). The effect of memory schemata on memory and the phenomenological experience of naturalistic situations. *Applied Cognitive Psychology*, 16(6), 687-708.
- Pertsov, Y., Avidan, G., & Zohary, E. (2009). Accumulation of visual information across multiple fixations. *Journal of Vision*, 9(10), 1-12.