

Affective And Cognitive Underpinnings of Moral Condemnation
When News of Transgressions Goes Viral

Online Supplement

Contents

Equations for Computing Indirect Effects in Experiments 1 and 2 (p. 2)

Experiment S1 (p. 4)

Experiment S2 (p. 7)

Details of Multilevel Mediation Analysis in Experiment 3 (p. 13)

The Present Research vs. Effron (2022): Effect-Size Comparison (p. 15)

References (p. 17)

Equations for Computing Indirect Effects in Experiments 1 and 2

The main text reports the results of a multilevel analysis testing three indirect effects of repetition on unethicity judgments: one indirect effect through affective intensity, one through infamy inferences, and one through unusualness inferences (see Figure 3 in the main text). This analysis requires estimating four equations:

$$\text{Equation 1: } \text{affect}_{ij} = \beta_0 + \beta_1 \text{repetition}_{ij} + \beta_2 \text{transgression1}_j + \dots + \beta_{14} \text{transgression13}_j + \beta_{15} \text{version1}_j + \beta_{16} \text{version2}_j + u_i + e_{ij}$$

$$\text{Equation 2: } \text{infamy}_{ij} = \beta_0 + \beta_1 \text{repetition}_{ij} + \beta_2 \text{transgression1}_j + \dots + \beta_{14} \text{transgression13}_j + \beta_{15} \text{version1}_j + \beta_{16} \text{version2}_j + u_i + e_{ij}$$

$$\text{Equation 3: } \text{unusual}_{ij} = \beta_0 + \beta_1 \text{repetition}_{ij} + \beta_2 \text{transgression1}_j + \dots + \beta_{14} \text{transgression13}_j + \beta_{15} \text{version1}_j + \beta_{16} \text{version2}_j + u_i + e_{ij}$$

$$\text{Equation 4: } \text{unethicity}_{ij} = \beta_0 + \beta_1 \text{repetition}_{ij} + \beta_2 \text{affect}_{ij} + \beta_3 \text{infamy}_{ij} + \beta_4 \text{unusualness}_{ij} + \beta_5 \text{transgression1}_j + \dots + \beta_{17} \text{transgression13}_j + \beta_{18} \text{version1}_j + \beta_{19} \text{version2}_j + u_i + e_{ij}$$

Equations 1–4 predict, respectively, affect, infamy, unusualness, and unethicity ratings from the i th participant for the j th transgression, and in each equation:

- β_0 is the fixed intercept
- transgression1–transgression13 are dummy codes for the 14 specific transgressions (transgression14 is omitted because it is the reference group)
- version1 and version2 are dummy codes for the 3 versions of the headline participants rated in the judgment phase (version3 is omitted because it is the reference group).
- u_i is the random-intercept variance for participants
- e_{ij} is the residual variance.

In Equations 1–3, β_1 is the path from repetition to, respectively, affect, infamy, and unusualness (i.e., the a -paths in the mediation model). In Equation 4, β_2 , β_3 , and β_4 represent,

respectively, the paths from affect, infamy and unusualness to unethicity judgments (i.e., the b -paths in the mediation model)

In Stata 17, we estimated these questions using the *gsem* command and computed each of the three indirect effects by multiplying the relevant a and b paths using the *nlcom* command in Stata 17 (StataCorp, 2021). For example, the indirect effect from repetition to affect to unethicity judgments is equal to the product of β_1 in Equation 1 and β_2 in Equation 4.

Experiment S1

Experiment 2 found that repetition significantly reduced moral condemnation overall when each repeated exposure came from a different news source. Experiment S1 tested whether this effect would replicate when we omitted measures of affect, infamy inferences, and unusualness perceptions, which might have called participants' attention to the theorized processes.

Method

We pre-registered Experiment S1 at https://aspredicted.org/V8C_2HC.

Participants

We posted slots on Prolific for 1,000 UK nationals who had not completed any of our lab's prior repetition studies. An *a priori* power simulation, conducted using the *simr* function in R (Green & MacLeod, 2015), suggested that this number of participants, with 14 repeated measures each, would provide 94% power to detect a raw mean difference of -1.50 between repeatedly-seen and new headlines, given variance estimates obtained from pilot data.

At the time we ran this experiment, Prolific's users skewed female, so we requested a gender-balanced sample. Participants could not begin the study if they did not consent, failed a reading comprehension question, or were on a mobile device. $N = 1,015$ people accessed the study; after applying the pre-registered exclusion criteria (i.e., duplicate IP address or Prolific ID) and dropping people who could not be analyzed because they did not provide any responses, the final sample size was $N = 1,003$ people (496 women, 489 men, and 6 nonbinary individuals; M age = 41 years, $SD = 15$) who provided 13,987 responses.

Procedure

Participants followed the same procedure described in Experiments 1–2. This time, however, we only measured moral judgments (i.e., how unethical each wrongdoing was), not affect or unusualness.

Results and Discussion

The results again showed a significant moral repetition effect: Participants rated repeatedly-seen transgressions ($M = 68.65$, $SE = 4.33$) as less unethical than new transgressions ($M = 70.52$, $SE = 4.33$), $b = -1.87$, $z = 4.36$, $p < .001$, $d_z = -.11$ in a mixed regression model with fixed effects for which of the three headlines about a given transgression participants rated, plus random intercepts for participants and item (see Table S1 for complete results and robustness check with random slopes). Because our theoretical model does not predict whether affective desensitization or infamy inferences will be the stronger process, we did not pre-register directional predictions for this analysis.

Most participants (84%) indicated correctly that all the headlines were from different news sources, and the results remained robust when only analyzed responses from these participants. (Fewer than 1% of people thought all headlines were from the same news source; the remainder thought that some headlines had been from the same sources and some had been from different sources).

Table S1*Mixed Regression Results from Experiment S1*

Fixed Effects	Pre-registered Model					Robustness Check				
	<i>b</i>	<i>SE</i>	<i>z</i>	<i>p</i>	95% <i>CI</i>	<i>b</i>	<i>SE</i>	<i>z</i>	<i>p</i>	95% <i>CI</i>
Repeated Headline	-1.87	0.43	-4.36	0.00	[-2.71, -1.03]	-1.87	0.71	-2.65	0.01	[-3.26, -0.49]
Version 2	-2.20	1.26	-1.75	0.08	[-4.68, 0.27]	-2.24	1.26	-1.78	0.08	[-2.54, 2.42]
Version 3	0.03	1.26	0.02	0.98	[-2.45, 2.50]	-0.06	1.26	-0.05	0.96	[-2.54, 2.42]
Intercept	71.24	4.39	16.22	0.00	[62.63, 79.85]	71.28	4.42	16.13	0.00	[62.62, 79.94]
Random Effects	Variance		<i>SE</i>		95% <i>CI</i>	Variance		<i>SE</i>		95% <i>CI</i>
Intercepts										
Participants	220.09		11.95		[197.87, 244.80]	218.25		12.04		[195.89, 243.18]
Headlines	258.38		97.91		[122.95, 543.00]	261.64		99.30		[124.35, 550.50]
Slopes										
Participants						13.45		7.83		[4.30, 42.12]
Headlines						4.24		2.64		[1.26, 14.35]
Residual	643.80		8.00		[628.32, 659.67]	638.88		8.21		[623.00, 655.17]

Note. The robustness check adds random slopes. Repeated headline was dummy coded (1 = repeatedly seen, 0 = new).

Experiment S2

In Experiment S1, as in Experiments 1 and 2, each repeated exposure to a transgression conveyed that a different news outlet had published a headline about it. Experiment S2 tested whether Experiment S1's results would replicate when each repeated exposure to a transgression involved reading an identical headline ostensibly shared online by a different social media user.

Method

Experiment S2 was pre-registered at https://aspredicted.org/PXR_PXX

Participants

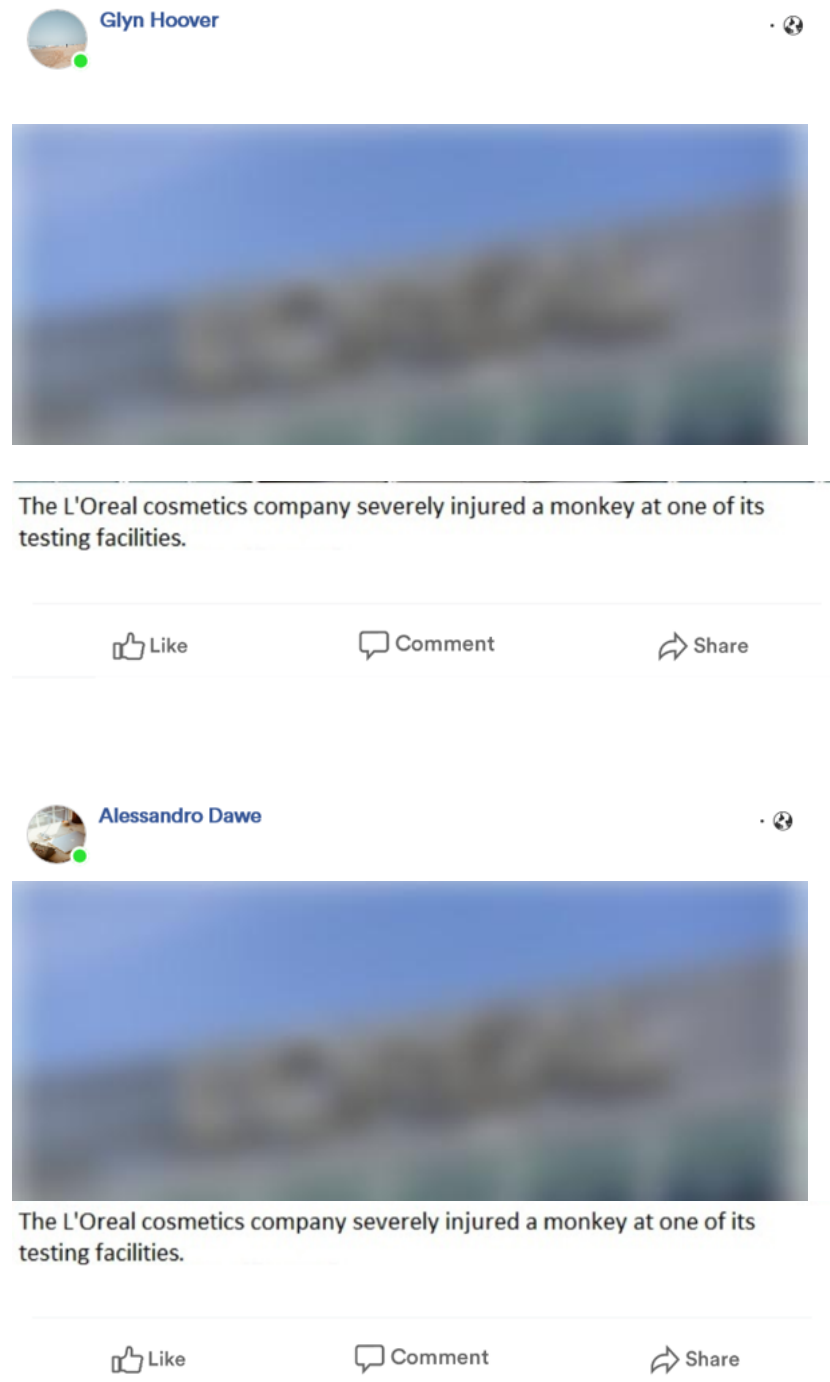
We requested 800 UK-based participants from Prolific who had not completed other repetition studies from our lab, and who passed a reading comprehension check. More than 800 began the study, and after dropping responses submitted from duplicate IP addresses or duplicate participant IDs, as pre-registered, we were left with 807 participants who provided 11,298 responses to the dependent measure. We did not conduct a power analysis but based this sample size on benchmarks from prior research (Effron, 2022).

Materials

The stimuli were 14 headlines about the corporate wrongdoings from Experiments 1–3. We created 3 Facebook-style posts for each headline and an accompanying image. Each post was identical, except had ostensibly been shared by a different social-media user, whose name and avatar were indicated above the headline (see Figure S1 for an example). The posts omitted information about the news source.

Figure S1

Example Headlines, Ostensibly from Different Social Media Users, in Experiment 1



Note. Participants saw an unblurred image of an office building with a sign reading “L’Oreal.” The images above were blurred for copyright reasons.

Procedure

The procedure closely followed Experiments 1, 2, and S1, except this time, each repeated exposure to a headline about a particular transgression reflected that a different social-media user had posted the headline.

At the beginning of the experiment, participants read:

In this study, you'll rate some news headlines shared by different social media users. The name of each user will be displayed above each headline. **You will only see each social media user once in this study.** However, it is possible for different social media users to share the same headline.

Then they completed the familiarization and distraction phases described in Experiment 1.

We next introduced the judgment phase by telling participants, “Now you’ll rate headlines **shared by a new set of social media users**. These users are different than the ones you saw before, even though some of the headlines they shared may be the same” (emphasis in original). Then, as in Experiments 1–S1, participants saw and rated each of the fourteen headlines in our bank, in randomized orders, on the dependent measure: “How unethical is the behavior described in the headline?” Recall that they had been exposed to a randomly selected half of the 14 headlines they had seen in the familiarization phase (*repeatedly-seen headlines*), and the other half they were seeing for the first time (*new headlines*). The avatars and names associated with each headline in the judgment phase were different than the two associated with the headlines in the exposure phase. We counterbalanced which avatars and names appeared in the judgment phase versus the familiarization phase.

We also manipulated whether participants imagined that their moral judgments would be public or private. As this manipulation is not central to our theoretical model and did not significantly moderate the results, we provide details after the *Results and Discussion* section.

Results and Discussion

As in Experiments 2 and S1, participants rated transgressions as less unethical when they had seen them before ($M = 72.58$, $SE = 4.59$) than when they had not ($M = 75.67$, $SE = 4.56$), $b = -3.09$, $SE(b) = .43$, $z = 7.25$, $p < .001$, $d_z = -.15$ in a mixed regression model with random intercepts for participants and headlines, plus fixed effects for the ancillary manipulation (see Table S1 for complete results and robustness check).

Almost everyone (97%) noticed that some of the headlines had been repeated, and of these participants, 91% indicated correctly that a different social-media user had shared the headlines each time. The results remained robust when we excluded participants who incorrectly thought that the same social-media users had shared multiple headlines.

Ancillary Manipulation and Results

We hypothesized that the moral repetition effect would be attenuated or eliminated when participants imagined that their moral judgments would be made public. As we detail below, this hypothesis was not supported.

Public/private manipulation. At the beginning of the judgment phase, participants randomly assigned to a *public* condition were told:

Suppose you posted something on one of your social-media accounts (whichever one you use most) for anyone to see. Please write the initials of five people who would be likely to see your post (for example, HS). These people can be friends, coworkers, family, acquaintances, “followers,” or anyone you like.

Then they read:

Imagine that your ratings of the headlines will be posted on social media for anyone to see. The people who see your ratings could be friends, coworkers, family, or even strangers. They might include the five people you just listed.

Before each rating, they were reminded that they should answer the questions as if their responses would be made public on social media.

Participants in the *private* condition were instead told:

Suppose you were planning to travel to five U.S. states next year on vacation (whichever states you like); cost is not an issue. Please write the abbreviations of the five states you would most like to visit (for example, NY). These states can be places you've been before, places you've never visited, or anywhere you like.

Then they were read:

Remember that your responses to the following questions will be kept completely confidential.

They were reminded of this confidentiality before providing each unethicality rating.

We asked all participants at the end of the study whether they complied with the manipulation.

Ancillary results. Our main result – that participants rated repeatedly-seen transgressions as less unethical than new transgressions – held regardless of whether participants imagined that their responses would be made public, $b = -3.92$, $SE(b) = .60$, $z = 6.53$, $p < .001$, or would remain private, $b = -2.25$, $SE(b) = .61$, $z = 3.71$, $p < .001$, and we found no support for our hypothesis that the moral repetition effect would be larger in the public than in the private condition. The interaction term was not significant, $b = -1.66$, $SE(b) = .85$, $p = .051$, and it was in the opposite direction than hypothesized. This pattern would be interesting to follow up on in subsequent research, but is not the focus of the theoretical model we test in the present investigation.

Table S2*Mixed Regression Results from Experiment S2*

Fixed Effects	Pre-registered Model					Robustness Check				
	<i>b</i>	<i>SE</i>	<i>z</i>	<i>p</i>	95% <i>CI</i>	<i>b</i>	<i>SE</i>	<i>z</i>	<i>p</i>	95% <i>CI</i>
Repeated Headline	-3.09	0.43	-7.25	0.00	[-3.93, -2.26]	-3.09	0.67	-4.65	0.00	[-4.40, -1.79]
Public	-0.09	1.12	-0.08	0.94	[-2.28, 2.10]	0.01	1.11	0.01	0.99	[-2.17, 2.19]
Intercept	75.72	4.62	16.37	0.00	[66.65, 84.78]	75.67	4.67	16.21	0.00	[66.52, 84.82]
Random Effects	<i>Variance</i>		<i>SE</i>		95% <i>CI</i>	<i>Variance</i>		<i>SE</i>		95% <i>CI</i>
Intercepts										
Participants	214.26		12.51		[191.09, 240.24]	210.03		12.57		[186.79, 236.16]
Headlines	289.90		109.81		[137.98, 609.09]	295.72		112.18		[140.60, 621.97]
Slopes										
Participants						15.95		7.34		[6.47, 39.31]
Headlines						3.41		2.31		[0.90, 12.85]
Residual	513.45		7.09		[499.73, 527.54]	508.38		7.25		[494.38, 522.78]

Note. The robustness check adds random slopes. Repeated headline was dummy coded (1 = repeatedly seen, 0 = new).

Details of Multilevel Mediation Analysis in Experiment 3

The main text reports the results of a multilevel analysis testing two indirect effects of repetition on unethicity judgments: one indirect effect through affective intensity, and one through unusualness inferences (see Figure 6 in the main text). This analysis requires three equations:

$$\text{Equation 5: } \text{affect}_{ij} = \beta_0 + \beta_1 \text{repetition}_{ij} + \beta_2 \text{transgression1}_j + \dots + \beta_{14} \text{transgression13}_j + u_i + e_{ij}$$

$$\text{Equation 6: } \text{unusual}_{ij} = \beta_0 + \beta_1 \text{repetition}_{ij} + \beta_2 \text{transgression1}_j + \dots + \beta_{14} \text{transgression13}_j + u_i + e_{ij}$$

$$\text{Equation 7: } \text{unethicity}_{ij} = \beta_0 + \beta_1 \text{repetition}_{ij} + \beta_2 \text{affect}_{ij} + \beta_3 \text{unusual}_{ij} + \beta_4 \text{transgression1}_j + \dots + \beta_{16} \text{transgression13}_j + u_i + e_{ij}$$

Equations 5–7 predict, respectively, affect, infamy, unusualness, and unethicity ratings from the i th participant for the j th transgression, and in each equation:

- β_0 is the fixed intercept
- transgression1–transgression13 are dummy codes for the 14 specific transgressions (transgression14 is omitted because it is the reference group)
- u_i is the random-intercept variance for participants
- e_{ij} is the residual variance.

In Equations 5 and 6, β_1 is the path from repetition to, respectively, affect and unusualness (i.e., the a -paths in the mediation model). In Equation 7 β_2 and β_3 represent, respectively, the paths from affect and unusualness to unethicity judgments (i.e., the b -paths in the mediation model)

We conducted this analysis in Stata 17, estimating the three equations using the *gsem* command and multiplying the relevant a and b paths using the *nlcom* command (StataCorp,

2021). For example, the indirect effect from repetition to affect to unethicality judgments is equal to the product of β_1 in Equation 5 and β_2 in Equation 7.

The Present Research vs. Effron (2022):

Effect-Size Comparison

Our theoretical model predicts that the moral repetition effect will be weaker in contexts that strengthen the infamy-inferences process relative to the affective-desensitization process. Comparing the repetition effect sizes observed in Experiment 1–S2 vs. 3 provides one test of this prediction. As an additional test, we examined the effect size observed in prior research.

In Experiment 1–S2’s paradigm, repetition sent a *modest* signal of infamy, conveying that at least 3 social-media users or news outlets had shared information about a transgression. As noted, repetition reduced moral condemnation. By contrast, in the paradigms used in prior research (e.g., Effron, 2022), repetition sent *no* signal of infamy, conveying *nothing* about the number of individuals or outlets who had shared such information. Our model thus predicts a smaller average effect size in Experiments 1–S2 than in this prior research.

To test this prediction, we first computed the average repetition effect size across two comparable prior studies: Effron (2022) Experiments 1a and 1b (total $N = 1,261$). Those experiments offer the most relevant benchmark because their stimuli included the same 14 wrongdoings as in the present research, also formatted as headlines (but without information about who shared or published those headlines).¹ A mixed regression model with random intercepts for participant and headline, plus fixed effects for study, showed that repetition significantly reduced moral condemnation, $b = -3.11$, $SE = .30$, $z = 10.48$, $p < .001$.

¹ The data for the experiments in Effron (2022) are available at that paper’s OSF site, https://osf.io/6xpq8/?view_only=eb17d8e867aa4fe7b5007baa17086859, and we also reposted them on the present paper’s OSF site. Only the “judge sharing” condition in those studies is comparable to the present paradigms; in another condition, which we do not analyze here, participants judged the unethicity of sharing content on social media rather than committing the wrongdoings described in the headlines. The only other methodological difference between our experiments and this prior work was that Effron (2022) informed participants in advance that the news was fake.

As our model predicts, the effect size we observed in Experiments 1–4 was significantly smaller than this benchmark from prior research, $z = 1.964$, $p < .05$, computed as follows:

$$z = \frac{(b_1 - b_2)}{\sqrt{SE(b_1)^2 + SE(b_2)^2}}$$

These were the results when we excluded responses to two wrongdoings in Effron’s studies that were not in the present experiments’ stimulus set. The results supported the same conclusion, however, when we included those responses, $z = 2.44$, $p = .0149$.

Thus, consistent with our model, the moral repetition effect was smaller in paradigms where repetition signaled that the wrongdoings were infamy than in paradigms where repetition sent no such signal.

References

- Effron, D. A. (2022). The moral repetition effect: Bad deeds seem less unethical when repeatedly encountered. *Journal of Experimental Psychology: General*, 151(10), 2562–2585.
<https://doi.org/10.1037/xge0001214>
- Green, P., & MacLeod, C. J. (2015). SIMR: An R package for power analysis of generalized linear mixed models by simulation. *Methods in Ecology and Evolution*, 7(4), 493–498.
<https://doi.org/https://doi.org/10.1111/2041-210X.12504>