

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Moralization by Analogy: A Novel Perspective on Moral Change

Permalink

<https://escholarship.org/uc/item/8wr906s2>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Pedersen, Julie Maria Ejby

Wilks, Matti

Doumas, Leonidas A. A.

et al.

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Moralization by Analogy: A Novel Perspective on Moral Change

Julie M. E. Pedersen (j.pedersen@ed.ac.uk)¹, Matti Wilks¹, Leonidas A. A. Doumas¹ & Adam Moore¹

¹Department of Psychology, University of Edinburgh

Abstract

Moral feelings and beliefs are powerful drivers of behavior. While analogies are commonly invoked to induce moral sentiment, analogical reasoning has not been empirically tested as a mechanism for moralization. Across two studies (N=292, N=234), we investigated whether analogical reasoning facilitated moralization of neutral target actions by transferring moral significance from moralized source actions. Participants judged neutral target actions paired with either analogous immoral source actions or unrelated immoral actions (control). Target actions were judged as more immoral in the analogy condition compared to the control, but only when the source and target were considered comparable (i.e., mappable relations), and the source was judged as highly immoral. These factors interacted synergistically, consistent with analogical transfer: moral significance transferred to targets when relational mappings supported the inference. However, poorly-perceived analogies backfired, producing less moralization than controls. These findings provide novel evidence for analogical reasoning as a mechanism for moralization.

Keywords: Moralization; Moral change; Analogical reasoning

Introduction

Engaging moral feelings or beliefs can permit (Mooijman et al., 2018; Skitka et al., 2021) or prevent (Crimston et al., 2018; Hoffman et al., 2013) violence and harm towards others. People seeking to influence behavior or opinions, such as activists or politicians, regularly attempt to induce such moral sentiment by drawing analogies to already moralized topics. For example, with the drowning child scenario, Peter Singer (1972) argued that inaction towards death from famine and preventable diseases in less affluent countries is as morally condemnable as refusing to save a drowning child in front of you. The same strategy is commonly employed by animal activists (e.g., Isaacs, 2015; PETA, 2009) and in online discussions (e.g., Fiadotava et al., 2023). Despite many anecdotes supporting the use of analogy to induce or amplify moral sentiment, little work has investigated the role of analogical reasoning as a mechanism for moralization. Here, we provide initial evidence for analogical reasoning as a mechanism for moral change via the transfer of moral significance from an already moralized action to a normally innocuous action.

Analogical Reasoning as a Mechanism of Transfer

Analogical reasoning is a well-established process-level theoretical framework of inductive inference facilitated by the

transfer of properties from one (source) domain to another (target) domain (see Holyoak, 2012, for a review). This directional transfer leverages a well-understood source domain to draw inferences about a less well-understood target domain, facilitated by the mappable relations between them. By identifying such relational correspondences between the source and target domains (Gentner, 1983; Hummel & Holyoak, 1997), additional properties of the source domain can be transferred and, thus, inferred about the target domain (Halford et al., 1998; Holyoak et al., 1994; Hummel & Holyoak, 2003). For example, the physical laws governing water waves were applied to understand sound and light when there was some evidence that sound and light behaved like waves (Holyoak, 2012). Likewise, people may extrapolate from previous social experiences to new ones, such as understanding and navigating a love triangle (Hummel & Holyoak, 2003). Multiple studies demonstrate such mapping (e.g., Spellman & Holyoak, 1992) and transfer (e.g., Aho et al., 2022; Day & Goldstone, 2011; Gick & Holyoak, 1980). These findings are further supported by a large computational modeling literature with algorithmic and implementation-level accounts of analogical mapping (Falkenhainer et al., 1989; Forbus et al., 1995; Hummel & Holyoak, 1997) and transfer (Hummel & Holyoak, 2003), which reliably predict human task performance. Thus, there is strong empirical and rigorous theoretical backing supporting analogical mapping and transfer as a psychological process through which properties are extended to new domains.

Analogical inferences are also subject to multiple constraints. This includes existing knowledge and beliefs, both topically, but also in terms of the exact relational structure of individual representations (e.g., causal vs. non-causal relational representations; Gentner & Toupin, 1986; Holyoak & Thagard, 1989; Lassaline, 1996). These constraints affect both whether analogical reasoning is initiated, most prominently by influencing spontaneous generation of appropriate sources (e.g., Gick & Holyoak, 1980), and what is transferred when a source is provided, shaping which correspondences and properties are identified and inferred (Schunn & Dunbar, 1996; Spellman et al., 2001). Additionally, analogical reasoning requires manipulating complex representations and, therefore, depends on working memory and other cognitive capacities (e.g., Holyoak & Hummel, 2001). As such, while analogical transfer can occur between a source and a target, it is subject to individual circumstances, perceived coherence of the mapping, and other beliefs retrieved during the process.

Moralization as Analogical Transfer

We propose that analogical reasoning, specifically analogical transfer, could be a mechanism for moralization by facilitating the extension of moral sentiment from an already moralized behavior to an otherwise neutral behavior. Indeed, there is some evidence for this mechanism in the context of meat-eating. Mapping morally relevant properties to meat-eating both in specific (e.g., similar social/cognitive capacities between pigs and dogs; Horne et al., 2021) and abstract (e.g., harm avoidance; Feinberg et al., 2019) scenarios led to more moralized attitudes. In line with the broader moral psychology literature, these findings also demonstrate the importance of direct or implied harm for moral judgments (Haslam et al., 2020; Rhee et al., 2019; Schein & Gray, 2016, 2018). Thus, there is some evidence that drawing parallels from moralized to non-moral behaviors, particularly when such analogies involve harm, can induce moralization. Yet, no work has systematically tested the role of analogical reasoning between moral and neutral domains as a cognitive mechanism of such moralization.

The Present Research

Engaging moral feelings and beliefs is a powerful and widespread strategy to incite harmful or helpful behavior. While we observe analogies used to induce moral sentiment, analogical reasoning has not previously been tested as a mechanism for moral change. We posit that analogical reasoning can facilitate moralization as a form of moral belief or attitude change by mapping morally relevant properties to extend moral sentiment from a moralized action to a neutral action. Across two studies, we test if moral sentiment can be extended to normally innocuous actions (*target actions*, e.g., mowing the lawn) by highlighting similarities with typically moralized immoral actions (*source actions*, e.g., deforestation). We consider how this interacts with participants' processing of the analogies, including to what extent they found the moralized action immoral and how good the analogy was (i.e., whether the source and target action were considered comparable).

Study 1

In Study 1, we preregistered a between-participants manipulation to test if analogies between morally salient (source) actions and morally neutral (target) actions would produce moralization of the target. In line with our proposed mechanism that analogical transfer facilitates moralization, we hypothesized that target actions presented in an analogy would be considered more immoral than the same actions when presented with an unrelated source action (H1). This mechanism would also predict that transfer depends on the strength of the source representation, therefore, we hypothesized that this would be amplified by the perceived immorality of the source action (H2). Finally, reflecting the prediction that analogical transfer should depend on the degree of mapping between source and target, we hypothesized that this amplification by source immorality would be strengthened by greater

perceived comparability between the source and target actions (H3). Supplementary materials and study preregistration are available on the [OSF](#).

Methods

Participants. We recruited 300 UK-based adults who spoke fluent English via Prolific. Eight participants were excluded based on preregistered criteria (failing both attention checks), leaving a final sample of 292 participants (see Supplementary Materials for sample demographics). Two participants had some missing data (see Figure 1).

Design. We conducted a between-participants manipulation of analogies between morally salient and neutral actions. We operationalized this as two sets of vignettes, analogy and control vignettes (see Materials for more details).

Materials. We produced and pre-tested 12 vignettes (six control and six analogy vignettes; provided in the Supplementary Materials). Each vignette had the same two-paragraph structure: a source paragraph describing the harms of a morally salient action (e.g., cutting down the rainforest destroys habitats and threatens biodiversity), followed by a target paragraph describing the harms of a morally neutral action (e.g., mowing and maintaining a grass lawn destroys habitats and threatens biodiversity). In two separate pre-tests, we verified that all source actions were considered immoral and all target actions were considered morally neutral ($N = 87$; see Supplementary Materials) and that all vignette components were reasonably believable and clear ($N = 79$; see Supplementary Materials). The analogy vignettes contained six pairs of morally salient and neutral actions that formed and were presented as analogies with comparative language (e.g., cutting down the rain forest is like mowing and maintaining a grass lawn). Meanwhile, control vignettes shuffled morally salient and neutral actions and presented these without comparative language. We used the same shuffled combinations for all participants in the control condition to balance variability.

Procedure and Measures. The experiment was administered online using JsPsych (de Leeuw et al., 2023). Participants were randomly assigned to either the analogy or control condition. They first completed six moral judgment trials. On each trial, participants read a vignette associated with their assigned condition and judged the morality of the source and target action described (order of source and target judgments were randomized between-participants). Moral judgments were measured with a sliding scale from extremely immoral (-100) to extremely moral (100) in five-step increments, where the scale mid-point represented moral neutrality (neither moral nor immoral). After each moral judgment, participants were also asked to explain their judgment in open text. Upon completing all moral judgments, participants were asked to judge the source-target comparability of the six vignettes on a 6-point Likert scale (0-5; see Figures 1C and 3C). During all moral and comparability judgments, the vignette

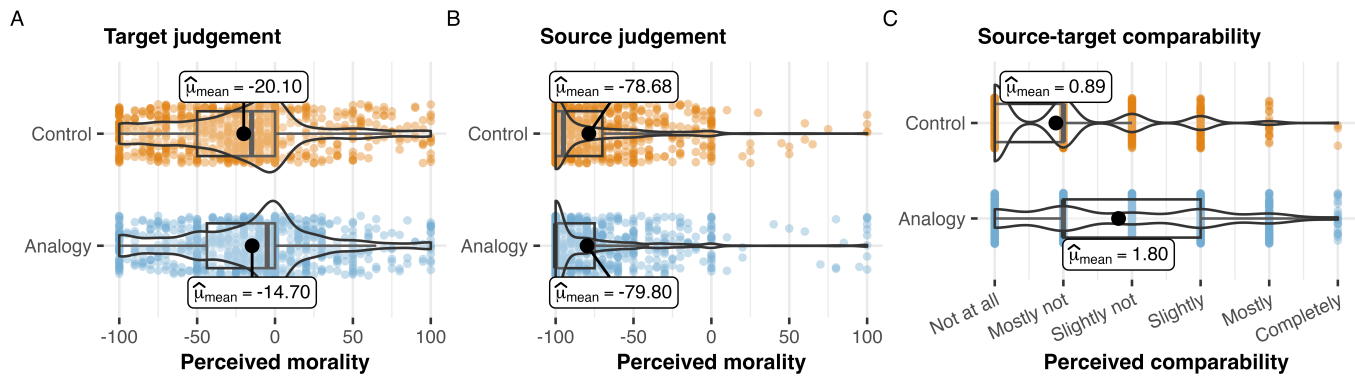


Figure 1: Distributions and means of marginal responses: perceived (A) target immorality, (B) source immorality, and (C) source-target comparability.

remained visible. Scale endpoints for moral and comparability judgments and the presentation order of the vignettes were randomized between participants. Counterbalancing of source target rating order and scale end points did not affect results (see Supplementary Materials).

Data Analysis. We fitted a preregistered Bayesian multi-level linear regression model predicting perceived target immorality with condition (analogy/control; reference level: control), perceived source action immorality, perceived source-target comparability, and all two- and three-way interactions between these. All variables were standardized for model estimation. We fitted random intercepts for participants and targets, reducing the preregistered random effect structure to stabilize model estimation. We considered effects credible if their 95% highest density intervals (HDIs) excluded zero. We dropped two non-credible coefficients (perceived source immorality and its interaction with source-target comparability) from the fixed effects per preregistered criteria. All effects were estimated with preregistered priors and all effects remained credible with weakly informative priors¹.

Results

Across conditions, source actions were consistently judged very immoral (see Figure 1B). In contrast, target moral judgments were more variable and, on average, less immoral in the analogy condition than the control condition (see Figure 1A). Source-target pairs in the analogy condition were considered more comparable on average than those in the control condition, though mean comparability was low in both conditions (see Figure 1C). Notably, in the analogy condition, bimodality in the comparability judgments indicated that participants were somewhat divided in their acceptance of the analogies.

The results of the statistical model (Bayesian $R^2 = 0.36$, 95% HDI[0.32, 0.39]) indicated a complex relationship between these variables (see Figure 2). At mean levels of comparability (notably low: $M = 1.35$) and source immorality, targets in the

analogy condition were credibly judged as less immoral than in the control condition ($\beta = 0.3$, 95% HDI[0.19, 0.42]; opposite direction of H1), while control targets were not judged credibly different from moral neutrality ($\beta = -0.11$, 95% HDI[-0.40, 0.16]). In both conditions, higher perceived source-target comparability was associated with judging the target action as more immoral ($\beta = -0.22$, 95% HDI[-0.29, -0.17]). This effect effectively doubled in the analogy condition due to a credible interaction ($\beta = -0.19$, 95% HDI[-0.24, -0.13]). These effects should be interpreted noting the distribution of comparability judgments; 78 % of control source-target pairs were considered "not at all" or "mostly not" comparable, while this was only the case for 51 % of the comparability judgments in the analogy condition (see Figure 1C). For source immorality, there was no credible effect on target immorality in the control condition (dropped from final model estimation), but higher source immorality was credibly associated with targets being judged as more immoral in the analogy condition ($\beta = 0.14$, 95% HDI[0.10, 0.19]; in line with H2). Finally, in the analogy condition, we found a synergistic relationship between source immorality and source-target comparability such that target actions were judged as even more immoral when comparability was high and source immorality was high ($\beta = 0.11$, 95% HDI[0.08, 0.14]; supporting H3). There was no credible interaction between perceived source immorality and source-target comparability in the control condition (dropped from final model estimation). Means and random effects for each target are provided in the Supplementary Material.

Interim Discussion

Study 1 demonstrated that neutral actions can become moralized through analogy with immoral actions, but only when these actions were perceived to be comparable and the source was considered highly immoral. This pattern aligns with predictions of an analogical reasoning account where the source analogue and relational mapping facilitate analogical transfer. Interestingly, when source and target were not comparable (i.e., mapping did not occur), targets in the control condition were judged as more immoral than in the analogy condition,

¹Weakly informative prior: $\beta_s \sim N(0, 1)$. Our preregistered priors reflected the existing literature (see Introduction).

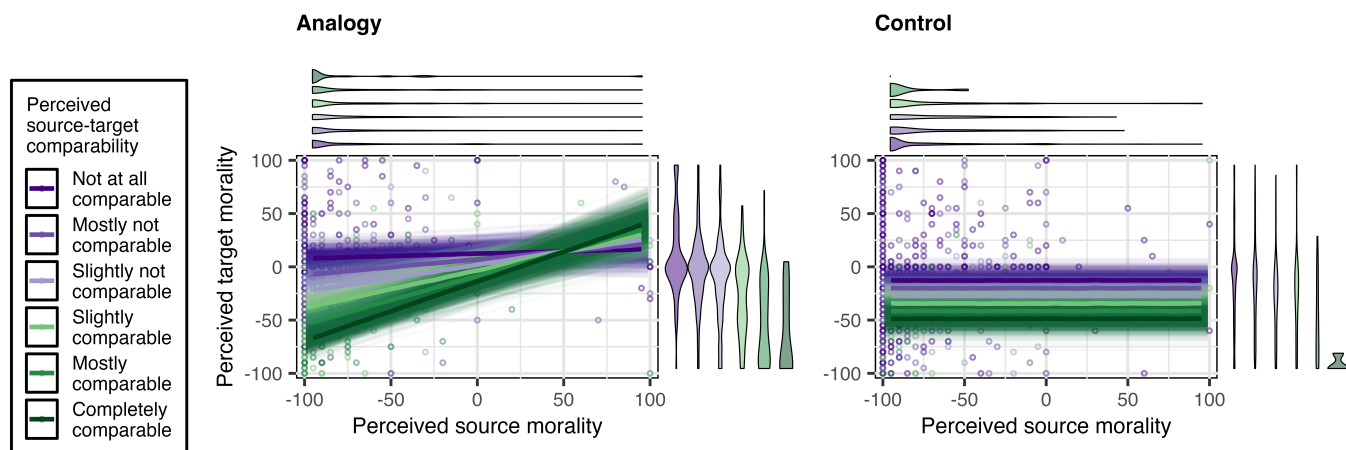


Figure 2: Predicted target moralization (thick line) with posterior draws (thin lines) by source immorality (x-axis), condition (facet), and source-target comparability (color). Points and marginal violin plots indicate raw data distributions of moral judgments for sources (top) and targets (right). Standardized model predictions were re-scaled to the raw scales of the data.

suggesting that poorly perceived analogies may backfire. We discuss this further in the General Discussion.

Given the importance of perceived source-target comparability to an analogical reasoning account, it is notable that the average comparability judgment in the analogy condition was relatively low and not much higher than the control condition average. Analogy vignettes were designed to have mappable elements while control condition vignettes were shuffled source-target pairs. We expected this to produce a larger difference in perceived comparability than we observed. This could be a limitation of our design where participants are presented with analogies, which may not align with their existing beliefs (Holyoak & Powell, 2016; Horne et al., 2015), leading to low perceived comparability through interference with relational mapping (Hummel & Holyoak, 1997). Alternatively, it may be an artifact of the between-participants design. Participants in the analogy condition only saw vignettes with mappable elements and may calibrate their comparability judgments to this, and likewise for participants in the control condition (e.g., Tversky, 1977). Study 2 addresses the latter possibility.

Study 2

In Study 2, we aimed to replicate the findings of Study 1 in a within-participants design and in a different population. The within-participants design ensured that participants all saw a combination of analogous and non-analogous source and target pairs. Based on the findings of Study 1 and our proposed mechanism of moralization by analogical transfer, we preregistered the following hypotheses: Higher source-target comparability will increase target moralization in the control condition (H1a) but more so in the analogy condition (H1b). In the analogy condition, the higher the immorality of the source, the more moralized the target will be (H2), amplified further by higher source-target comparability (H3). Supplementary materials and study preregistration are available on the [OSF](#).

Methods

Participants. We recruited 250 undergraduate students. Sixteen participants were excluded based on preregistered criteria (failing both attention checks or being underage), leaving a final sample of 234 participants with no missing data. Sample demographics are provided in the Supplementary Materials.

Design. We conducted a within-participants manipulation of analogies between morally salient and neutral actions. We operationalised this as two sets of vignettes, analogy and control vignettes (see Materials).

Materials. We used the same vignettes as Study 1. We created 15 different vignette sets, each containing two analogy vignettes (identical to Study 1) and two control vignettes constructed by recombining one component from each of the four unused analogy vignettes. This prevented participants from identifying control vignettes as shuffled analogies.

Procedure and Measures. The experiment was administered online using JsPsych (de Leeuw et al., 2023) according to the same procedure as Study 1. Participants were randomly assigned to one of the 15 vignette sets (counter-balanced). They then completed four moral judgment trials followed by four comparison trials using the Study 1 measures. Counterbalancing of source target rating order and scale end points did not affect results (see Supplementary Materials).

Data Analysis. We preregistered and fitted a multilevel Bayesian regression model with the same fixed effect structure, variable standardization, and inference criteria (95% HDI excluding zero) as Study 1. We fitted an independent random intercept and condition slope by participant as well as random intercepts by target and source nested within target, reducing the preregistered random effect structure to stabilize

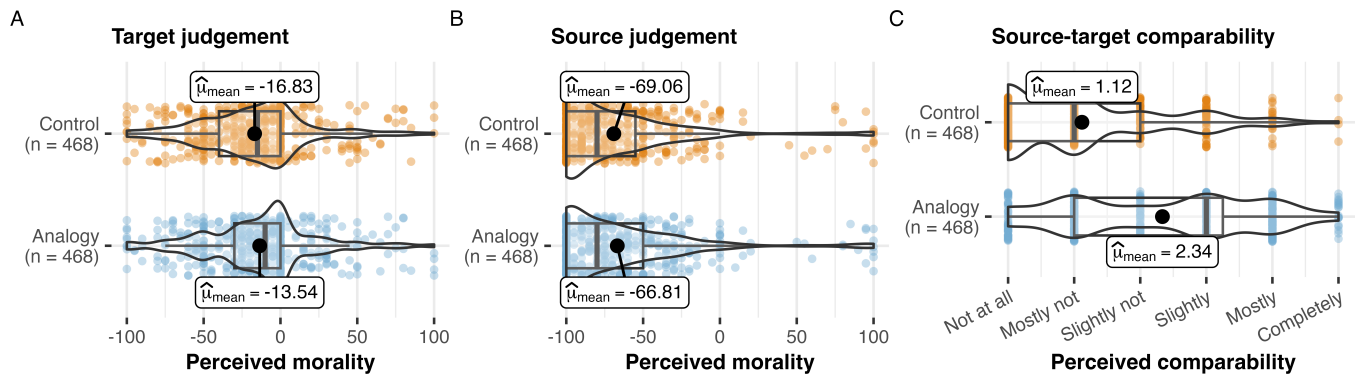


Figure 3: Distributions and means of marginal responses: perceived (A) target immorality, (B) source immorality, and (C) source-target comparability.

model estimation. Study 1 results were used as priors², with weakly informative priors for parameters that were removed in Study 1. As in Study 1, we removed non-credible parameters per preregistered criteria. All effects were estimated with preregistered priors; retained effects remained credible with weakly informative (WI) priors unless otherwise noted.

Results

Study 2 replicated key findings from Study 1 (see Figure 3A-C). However, both target and source actions were judged as more morally moderate on average than in Study 1 (see Figure 3A-B). In line with expectation, sources and targets in the analogy condition were considered more comparable than in Study 1 (see Figure 3C).

Overall, the statistical model (Bayesian $R^2 = 0.37$, 95% HDI[0.32, 0.43]) replicated the pattern of observed effects in Study 1 (compare Figures 2 and 4). However, in line with the dampened moral judgments of sources and targets observed in the raw data, effects were weaker. Replicating Study 1, at mean levels of comparability (notably low: $M = 1.72$) and source immorality, targets in the analogy condition were credibly judged as less immoral than in the control condition (i.e., more moral; $\beta = 0.31$, 95% HDI[0.23, 0.40]), while control targets were not judged credibly different from moral neutrality ($\beta = -0.19$, 95% HDI[-0.41, 0.03]). In the control condition, target actions were judged as credibly more immoral, the more comparable the source and target were ($\beta = -0.22$, 95% HDI[-0.25, -0.18]; in line with H1a). Moreover, like in Study 1, this was credibly more so in the analogy condition due to an interaction ($\beta = -0.16$, 95% HDI[-0.25, -0.11]; supporting H1b). However, this effect was not reproducible with weakly informative priors ($\beta_{WI} = -0.04$, 95% HDI[-0.17, 0.01]). These effects should be considered along with the distribution of comparability judgments; in the control condition, 72 % of source-target pairs were judged as "not at all" or "mostly not" comparable, while only 37 % of comparability judgments in the analogy condition were in that

range. In the analogy condition, higher source immorality was credibly associated with higher target immorality ($\beta = 0.13$, 95% HDI[0.09, 0.16], in line with H2), although this was not reproduced with weakly informative priors ($\beta_{WI} = 0.09$, 95% HDI[0, 0.17]). We found no effect of perceived immorality of the source action in the control condition (dropped from model estimation). Finally, replicating Study 1, in the analogy condition, higher source-target comparability credibly amplified the positive relationship between source and target immorality, leading to more immoral target judgments for more immoral source judgment and higher perceived source-target comparability ($\beta = 0.10$, 95% HDI[0.07, 0.14]; supporting H3). There was no credible interaction between source immorality and source-target comparability in the control condition (dropped from final model estimation). Means and random effects for each target are provided in the Supplementary Material.

General Discussion

Anecdotally, analogies are often used to induce moral sentiment. However, to our knowledge, this paper represents the first empirical test of analogy as a mechanism for moral reasoning. We find that analogies drawing parallels between the harms of immoral (source) actions and neutral (target) actions led to moralization more so than unmatched pairings, only if the source action is considered immoral and comparable to the target action. We also observed a potential backfire effect when these conditions were not met.

The observed pattern of effects strongly aligns with our proposed account of moralization via analogical transfer. That is, moralization should occur through (a) mapping relational correspondences between source and target (Gentner, 1983; Holyoak, 2012), and (b) transferring moral significance based on those mappings (Hummel & Holyoak, 2003). In our studies, source-target comparability reflected mapping quality while source immorality captured transferable moral content. In line with an analogical reasoning account, these factors interacted to facilitate moralization of otherwise morally neutral actions when presented as analogies, but not when presented in a non-analogy context. While comparability and source

²Normally distributed with the β estimate and associated SE from Study 1 as the mean and SD.

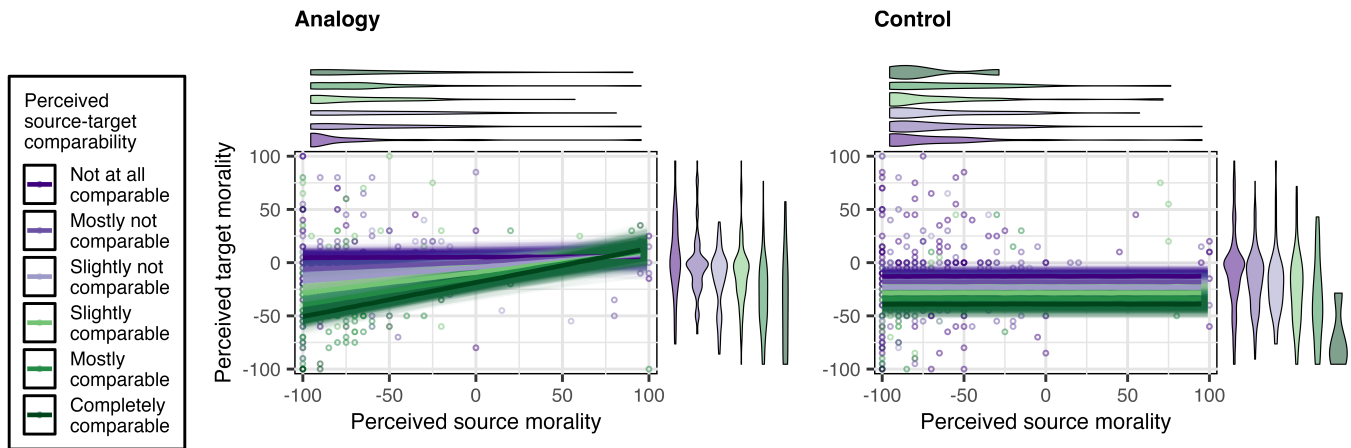


Figure 4: Predicted target moralization (thick line) with posterior draws (thin lines) by source immorality (x-axis), condition (facet), and source-target comparability (color). Points and marginal violin plots indicate raw data distributions of moral judgments for sources (top) and targets (right). Standardized model predictions were re-scaled to the raw scales of the data.

immorality also showed some independent effects on target moralization for analogies, these were less robust (see Study 2). Our results, thus, support that analogical transfer may facilitate moralization. This fits with existing empirical research on moralization (Feinberg et al., 2019; Horne et al., 2021) and further emphasizes analogical transfer as a theoretical framework to guide further study and understanding of moralization. Although these findings are not conclusive, they provide initial empirical support for analogical transfer as a mechanism for moralization of neutral actions and opens avenues for further research.

We also observed that targets in the analogy condition were judged as less immoral than in the control condition at low levels of source-target comparability. This suggests that a poorly received analogy may backfire, producing less moralization than without using a comparison. While the mechanism underlying this effect is unclear, one possibility is that when analogies conflict with participants' existing beliefs, they may provoke resistance due to incoherence (Holyoak & Powell, 2016; Horne et al., 2015) or interfere with the mapping process (Hummel & Holyoak, 1997). These explanations are in line with the sensitivity of the initiation and content of analogical inference to both external (e.g., Spellman et al., 2001) and internal constraints (e.g., Holyoak & Hummel, 2001). This backfire effect suggests that moralization by analogy is conditional, perhaps with the similar constraints to analogical inference.

Limitations and Future Directions

While our findings support analogical reasoning as a mechanism for moralization, this initial work has limitations. First, it is unclear exactly how moralization by analogy should behave according to theories of analogical reasoning. Analogical reasoning is constrained by individual-level factors such as working memory (e.g., Holyoak & Hummel, 2001) as well as multiple sub-processes that can affect the outcome

of an analogical inference, including source retrieval (e.g., Gick & Holyoak, 1980), relational mapping (e.g., Spellman & Holyoak, 1992), and analogical transfer (e.g., Spellman et al., 2001). Computational simulations of analogical inference (e.g., Hummel & Holyoak, 2003) using our materials could determine if our findings align with the theoretical and process-level assumptions of analogical reasoning. Second, because we provided the relational mappings, it is unclear whether participants actually engaged in analogical reasoning rather than more passive processing. As moralization has not been studied in the context of analogical reasoning, and spontaneous analogical transfer is difficult to manipulate (see Gick & Holyoak, 1980), we used pre-constructed analogies (i.e., vignettes with sources and targets) to test its potential as an explanatory mechanism. Future research should test designs where participants generate mappings themselves (e.g., Spellman & Holyoak, 1992) in the moral domain to further investigate the role of analogical reasoning mechanisms in moralization.

Conclusion

Moral feelings and beliefs are powerful drivers of harmful and helpful behavior. Analogies are often invoked to induce or amplify moral sentiment towards actions by aligning them with already moralized actions. Yet, the role of analogical reasoning has not been empirically tested in the context of moralization. In two studies, we found evidence for analogical transfer as a mechanism of such moral change. We showed that normally neutral actions can be moralized when presented as analogies with already moralized source actions. However, analogies that were perceived as poor comparisons hindered rather than facilitated moralization. Moreover, we observed a pattern of effects consistent with facilitators of analogical reasoning. Together, these findings support analogical reasoning as a potential *mechanism* for moralization of previously neutral actions.

Acknowledgments

AI was used to support editing the vignettes, coding the experiment, and formatting the figures in R. This work was supported in part by the Economic and Social Research Council through the Scottish Graduate School of Social Sciences [ESRC Grant number: ES/P000681/1] and the University of Edinburgh.

References

- Aho, K., Roads, B. D., & Love, B. C. (2022). System alignment supports cross-domain learning and zero-shot generalisation. *Cognition*, 227, 105200. <https://doi.org/10.1016/j.cognition.2022.105200>
- Crimston, C. R., Hornsey, M. J., Bain, P. G., & Bastian, B. (2018). Toward a Psychology of Moral Expansiveness. *Current Directions in Psychological Science*, 27(1), 14–19. <https://doi.org/10.1177/0963721417730888>
- Day, S. B., & Goldstone, R. L. (2011). Analogical transfer from a simulated physical system. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 37(3), 551–567. <https://doi.org/10.1037/a0022333>
- de Leeuw, J. R., Gilbert, R. A., & Luchterhandt, B. (2023). jsPsych: Enabling an Open-Source Collaborative Ecosystem of Behavioral Experiments. *Journal of Open Source Software*, 8(85), 5351. <https://doi.org/10.21105/joss.05351>
- Falkenhainer, B., Forbus, K. D., & Gentner, D. (1989). The structure-mapping engine: Algorithm and examples. *Artificial Intelligence*, 41(1), 1–63. [https://doi.org/10.1016/0004-3702\(89\)90077-5](https://doi.org/10.1016/0004-3702(89)90077-5)
- Feinberg, M., Kovacheff, C., Teper, R., & Inbar, Y. (2019). Understanding the process of moralization: How eating meat becomes a moral issue. *Journal of Personality and Social Psychology*, 117(1), 50–72. <https://doi.org/10.1037/pspa0000149>
- Fiadotava, A., Astapova, A., Hendershott, R., McKinnon, M., & Jürgens, A.-S. (2023). Injecting fun? Humour, conspiracy theory and (anti)vaccination discourse in popular media. *Public Understanding of Science*, 32(5), 622–640. <https://doi.org/10.1177/09636625221147019>
- Forbus, K. D., Gentner, D., & Law, K. (1995). MAC/FAC: A Model of Similarity-Based Retrieval. *Cognitive Science*, 19(2), 141–205. https://doi.org/10.1207/s15516709cog1902_1
- Gentner, D. (1983). Structure-mapping: A theoretical framework for analogy. *Cognitive Science*, 7(2), 155–170. [https://doi.org/10.1016/S0364-0213\(83\)80009-3](https://doi.org/10.1016/S0364-0213(83)80009-3)
- Gentner, D., & Toupin, C. (1986). Systematicity and surface similarity in the development of analogy. *Cognitive Science*, 10(3), 277–300. [https://doi.org/10.1016/S0364-0213\(86\)80019-2](https://doi.org/10.1016/S0364-0213(86)80019-2)
- Gick, M. L., & Holyoak, K. J. (1980). Analogical problem solving. *Cognitive Psychology*, 12(3), 306–355. [https://doi.org/10.1016/0010-0285\(80\)90013-4](https://doi.org/10.1016/0010-0285(80)90013-4)
- Halford, G. S., Wilson, W. H., & Phillips, S. (1998). Processing capacity defined by relational complexity: Implications for comparative, developmental, and cognitive psychology. *Behavioral and Brain Sciences*, 21(6), 803–831. <https://doi.org/10.1017/S0140525X98001769>
- Haslam, N., Dakin, B. C., Fabiano, F., McGrath, M. J., Rhee, J., Vylomova, E., Weaving, M., & Wheeler, M. A. (2020). Harm inflation: Making sense of concept creep. *European Review of Social Psychology*, 31(1), 254–286. <https://doi.org/10.1080/10463283.2020.1796080>
- Hoffman, S. R., Stallings, S. F., Bessinger, R. C., & Brooks, G. T. (2013). Differences between health and ethical vegetarians. Strength of conviction, nutrition knowledge, dietary restriction, and duration of adherence. *Appetite*, 65, 139–144. <https://doi.org/10.1016/j.appet.2013.02.009>
- Holyoak, K. J. (2012, March). Analogy and Relational Reasoning. In K. J. Holyoak & R. G. Morrison (Eds.), *The Oxford Handbook of Thinking and Reasoning*. Oxford University Press. Retrieved September 17, 2021, from <http://oxfordhandbooks.com/view/10.1093/oxfordhb/9780199734689.001.0001/oxfordhb-9780199734689-e-13>
- Holyoak, K. J., & Hummel, J. E. (2001). Toward an Understanding of Analogy within a Biological Symbol System. <https://doi.org/10.7551/mitpress/1251.003.0008>
- Holyoak, K. J., & Powell, D. (2016). Deontological coherence: A framework for commonsense moral reasoning. *Psychological Bulletin*, 142(11), 1179–1203. <https://doi.org/10.1037/bul0000075>
- Holyoak, K. J., & Thagard, P. (1989). Analogical Mapping by Constraint Satisfaction. *Cognitive Science*, 13(3), 295–355. https://doi.org/10.1207/s15516709cog1303_1
- Holyoak, K. J., Novick, L. R., & Melz, E. (1994). Advances in connectionist and neural computation theory. Vol. 2: Analogical connections. In K. J. Holyoak & J. A. Barnden (Eds.), *Analogical connections* (pp. 113–180). Ablex.
- Horne, Z., Powell, D., & Hummel, J. (2015). A Single Counterexample Leads to Moral Belief Revision [eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1111/cogs.12223>]. *Cognitive Science*, 39(8), 1950–1964. <https://doi.org/10.1111/cogs.12223>
- Horne, Z., Rottman, J., & Lawrence, C. (2021). Can coherence-based interventions change dogged moral beliefs about meat-eating? *Journal of Experimental Social Psychology*, 96, 104160. <https://doi.org/10.1016/j.jesp.2021.104160>
- Hummel, J. E., & Holyoak, K. J. (1997). Distributed representations of structure: A theory of analogical access and mapping. *Psychological Review*, 104(3), 427–466. <https://doi.org/10.1037/0033-295X.104.3.427>
- Hummel, J. E., & Holyoak, K. J. (2003). A symbolic-connectionist theory of relational inference and generalization. *Psychological Review*, 110(2), 220–264. <https://doi.org/10.1037/0033-295X.110.2.220>
- Isaacs, A. (2015, October). Q&A: Animal Rights Activist and Holocaust Survivor Alex Hershaft. Retrieved February 2, 2026, from <https://momentmag.com/qa-animal-rights-activist-and-holocaust-survivor-alex-hershaft/>

- Lassaline, M. E. (1996). Structural alignment in induction and similarity. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 22(3), 754–770. <https://doi.org/10.1037/0278-7393.22.3.754>
- Mooijman, M., Hoover, J., Lin, Y., Ji, H., & Dehghani, M. (2018). Moralization in social networks and the emergence of violence during protests. *Nature Human Behaviour*, 2(6), 389–396. <https://doi.org/10.1038/s41562-018-0353-0>
- PETA. (2009, April). If Your Dog Tasted Like Pork, Would You Eat Her? Retrieved February 2, 2026, from <https://www.peta.org/features/dog-pig/>
- Rhee, J. J., Schein, C., & Bastian, B. (2019). The what, how, and why of moralization: A review of current definitions, methods, and evidence in moralization research [eprint: <https://compass.onlinelibrary.wiley.com/doi/pdf/10.1111/spc3.12511>]. *Social and Personality Psychology Compass*, 13(12), Article e12511. <https://doi.org/10.1111/spc3.12511>
- Schein, C., & Gray, K. (2016). Moralization and Harmification: The Dyadic Loop Explains How the Innocuous Becomes Harmful and Wrong. *Psychological Inquiry*, 27(1), 62–65. <https://doi.org/10.1080/1047840X.2016.1111121>
- Schein, C., & Gray, K. (2018). The Theory of Dyadic Morality: Reinventing Moral Judgment by Redefining Harm. *Personality and Social Psychology Review*, 22(1), 32–70. <https://doi.org/10.1177/1088868317698288>
- Schunn, C. D., & Dunbar, K. (1996). Priming, analogy, and awareness in complex reasoning. *Memory & Cognition*, 24(3), 271–284. <https://doi.org/10.3758/BF03213292>
- Singer, P. (1972). Famine, Affluence, and Morality. *Philosophy and Public Affairs*, 1(3), 229–243. <http://www.jstor.org/stable/2265052>
- Skitka, L. J., Hanson, B. E., Morgan, G. S., & Wisneski, D. C. (2021). The Psychology of Moral Conviction. *Annual Review of Psychology*, 72, 347–366. <https://doi.org/10.1146/annurev-psych-063020-030612>
- Spellman, B. A., & Holyoak, K. J. (1992). If Saddam is Hitler then who is George Bush? Analogical mapping between systems of social roles. *Journal of Personality and Social Psychology*, 62(6), 913–933. <https://doi.org/10.1037/0022-3514.62.6.913>
- Spellman, B. A., Holyoak, K. J., & Morrison, R. G. (2001). Analogical priming via semantic relations. *Memory & Cognition*, 29(3), 383–393. <https://doi.org/10.3758/BF03196389>
- Tversky, A. (1977). Features of similarity [Place: US Publisher: American Psychological Association]. *Psychological Review*, 84(4), 327–352. <https://doi.org/10.1037/0033-295X.84.4.327>