

Absence Makes the Trust in Causal Models Grow Stronger

Samantha Kleinberg (samantha.kleinberg@stevens.edu)

Department of Computer Science, 1 Castle Point on Hudson
Hoboken, NJ 07030 USA

Eren Alay (ealay@stevens.edu)

Department of Computer Science, 1 Castle Point on Hudson
Hoboken, NJ 07030 USA

Jessecar K. Marsh (jessecar.marsh@lehigh.edu)

Department of Psychology, 17 Memorial Drive East
Bethlehem, PA 18018 USA

Abstract

People prefer complex explanations for complex phenomena, but make better choices when given only the information required. Thus there is a tension between the information people want, and the information they are able to use effectively. However, little is known about how the specific types of information included in causal models influences how people perceive them. We examine how omitting information influences how people reason about causal models, varying whether commonly known or unexpected information is removed (Experiment 1) or which parts of a causal path are omitted (Experiment 2). We find that omitting causal information participants expect to see lowers ratings of trust and other factors, while omitting less commonly known information improves ratings. However, causal paths can be simplified without harming perceptions of diagrams.

Keywords: causal models; complexity; simplicity

Introduction

Causal models, which depict causal relationships among a set of variables, can be powerful cognitive aids. Often these diagrams are represented as a set of nodes with edges connecting causes to effects. They can succinctly capture highly complex social and biological systems, such as the causes of depression or why people change jobs. Yet the types of comprehensive models created by domain experts or computational methods are often too complex for individual decision-makers to reason about and use successfully (Chan, 2001). We hear that “knowledge is power,” and are instructed to become informed decision makers to take control of our health, finances, or civic government. But do we really need to know every detail of how a bill becomes a law to make informed decisions on how to vote?

Recently, we showed that when making decisions about everyday scenarios, there is such a thing as too much information (Kleinberg & Marsh, 2021). When participants were given complex causal models they made the same choices as when they received no information at all, while simple causal models that include only the causal paths relevant to the answer led to significantly better choices. This work suggests that not including all causal relationships in the diagram presented, what we will call omitting information, can improve decisions by making it clearer what parts of a causal model to focus on. While this suggests a path for using causal models

to aid decisions, it raises new questions about the difference between the information people desire and the information that will help them. Namely, if we give people the simplified information that will help them make better choices, will they use it?

In everyday domains such as making decisions about jobs, health, or whether to bring an umbrella, we have existing knowledge and interpret new information in light of this knowledge. Zheng, Marsh, Nickerson, and Kleinberg (2020) found in fact that familiarity with a domain can impede use of causal models, even when individuals are able to use them successfully in novel scenarios. When a problem is about novel entities like blickets or numbered nodes in a Bayesian network, information mainly comes from the problem set-up rather than our prior experience, knowledge, or preferences. However, in real-world domains when we are given a simplified model or simplified guidance by an expert or government agency (e.g., get 150 minutes a week of exercise), information we expect to see may be missing. We may ask why the guidelines do not mention the type of exercise, recall that running far makes us tired for the rest of the day, or wonder if dietary changes are more impactful than physical activity.

It is an open question as to how the omission of expected information may influence trust in a model or a user’s willingness to use it. For example, an individual who has diabetes may be skeptical of a model of blood glucose management that does not include insulin. In general, information may be omitted due to the information not being relevant to the decision or choices made by the model’s creator. Further, a reasoner may think information is missing because of their own incorrect beliefs about what should be included in a model. Prior work has mainly examined preferences regarding simplicity and complexity of causal explanations, as opposed to causal models. Individuals often show a preference for simple explanations (Lombrozo, 2007; Liquin & Lombrozo, 2022), with simplicity being closely linked to the number of root nodes (Pacer & Lombrozo, 2017). In this view, a model with many root nodes is considered to have more “unexplained” causes, and is thus more complex. Recently, work on real-world domains has shown a preference for explanations that match the perceived complexity of the system (Lim & Oppen-

heimer, 2020), which may explain findings showing a preference for complex explanations in some tasks (Zemla, Sloman, Bechlivanidis, & Lagnado, 2017; Johnson, Valenti, & Keil, 2017). Other work showed that drawing causal models themselves led people to prefer simpler legal explanations, while people still preferred complex explanations when not required to model the evidence (Liefgreen & Lagnado, 2021). Taken together, these works suggest that in complex scenarios people may find a complex model more satisfying, but this does not yet address the question of whether people will trust or be willing to use a simplified model.

Evidence from studies of trust in automated systems suggest that expectations and existing knowledge do influence trust (Hoff & Bashir, 2015). Further, more detailed explanations of clinical decision support systems led to increased trust in the systems compared to less detailed explanations (Bussone, Stumpf, & O'Sullivan, 2015). This may provide further support for the complexity matching hypothesis of Lim and Oppenheimer (2020), and suggests that people may trust the simple models that improved decisions in (Kleinberg & Marsh, 2021) less than they trust complex models. However, there are many ways a model can be simplified and we do not yet know how the specific information content that is included or omitted may influence trust in and use of models.

This problem is also related to belief revision, in that in everyday situations we are updating our mental causal models in light of new evidence (Fernbach & Sloman, 2009). Recent work on how people update their causal models found that later learning about intermediate links in causal chains reduced people's estimates of the strength of a causal relationship (Stephan, Tentori, Pighin, & Waldmann, 2021). This suggests that omitting information in a chain may have an impact on how people perceive the causal relationships presented, though this work focused on fictitious diseases and genes. It is an open question of how in real world cases where participants may have previously known about the omitted link such omissions will be viewed.

In this study we test how omitted information influences perceptions of causal models. Experiment 1 tests whether participant expectations influence how they perceive omissions, by leaving out information that is either commonly reported or not commonly reported as a cause. We test models across natural, biological, and social domains. In Experiment 2 we test whether the role of a node in a causal model influences how participants perceive its omission. Specifically, we manipulate whether the omitted factor is a root node or a direct cause at the end of the same causal chain. Across both experiments participants rated how much they believe the models, how compatible they are with their existing knowledge, whether the models are understandable, whether they would be useful for making a decision, and finally whether they are trustworthy. Finally, we examine participant intuitions behind why information is omitted. Together, these experiments provide new insight into how participants reason about simplified models and how to create simple models that people believe.

Experiment 1

Prior work has shown that causal information can aid decisions when it is pared down to solely the information needed to make a decision (Kleinberg & Marsh, 2021). However, omitting information that participants expect to see may reduce their trust in the model, and thus make them less likely to use it. While Lim and Oppenheimer (2020) showed that people do prefer complex explanations of topics they perceive to be complex, the specific content of causal models has not yet been examined. We now test whether the absence of information participants expect to see influences their perceptions of a model, and if these perceptions differ from when unexpected information is omitted.

Method

Participants We recruited 150 U.S. residents aged 18-64 (73 female, 72 male, 5 identifying in other ways) from Prolific. Participants were compensated \$4.50 based on an expected study duration of 30 minutes. All participants were included in analysis.

Materials In this experiment we aim to manipulate how expected the omitted information is. Thus, we selected a variety of topics on which we can expect participants to have causal beliefs. We pilot tested 22 topics spanning natural, biological, and social phenomena with 72 subjects. Subjects answered an open ended question on causality for each topic (e.g. "Why do volcanoes erupt?" and "Why do people develop food allergies?"). After coding responses, we selected 8 domains that varied in complexity (average number of causes mentioned) and consensus (i.e., whether most participants mentioned one particular cause, or the number of mentions was uniformly distributed among a number of causes).

For each topic we developed a causal diagram based on the scientific literature in the domain. Diagrams were kept to a similar size, and varied from 4 to 6 nodes, and 3 to 5 edges (causal relationships). An example diagram depicting causes of developing gray hair is shown in Figure 1. For each diagram we created two variations: one where the cause that was most frequently mentioned in pilot testing was omitted, and one where we omitted the cause that was mentioned least frequently. In the example shown these were aging (most frequent) and illness (least frequent).

Procedure After consenting to the study, participants were instructed in the meaning of causal diagrams and what nodes and edges indicate. Participants saw one of three versions that varied what topics they would see in what format. Given that we have 8 diagrams, each group saw 3 in each of two formats (e.g., complete, most frequent missing) and 2 in the other format (e.g., least frequent missing). In this way, the format manipulation was within-subjects, with the exact topic presented in each format varying across participants. The order of topics was randomized for each participant. During the first stage of the study participants saw a single diagram per page with a series of statements capturing different facets of

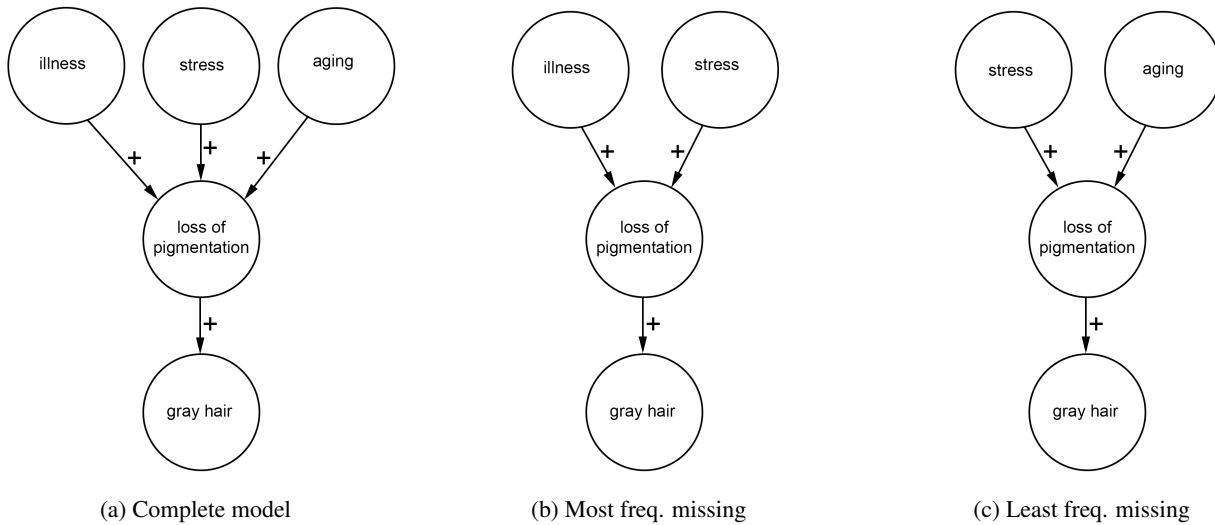


Figure 1: Diagrams used in Experiment 1. Diagram (a) is complete, while in (b) the most frequently mentioned cause (aging) is missing and in (c) the least frequently mentioned cause (illness) is omitted.

how a diagram may be perceived. The statements (and our labels for them, indicated in parentheses) were:

- How much do you believe the relationships shown in this diagram? (believability)
- How well does the information in this diagram fit with what you already know about the topic? (compatibility)
- How well could you explain the information in this diagram to a friend? (understandability)
- To what extent would this diagram help you make decisions about the topic? (utility)
- To what extent do you trust the information in this diagram? (trust)

Participants rated their agreement with each statement on a scale from 0 (not at all) to 6 (completely). After completing ratings for all 8 topics, we then elicited intuitions for the reasons behind missing information. We showed participants each incomplete diagram they had previously seen, now with the previously missing node and edge included and indicated in orange. Along with this diagram we asked “Some people also think [X] causes [Y]. Why do you think X was missing from the diagram you saw before? Please select the primary reason below.” The options provided were: the person creating the diagram didn’t know about it (missing knowledge), it’s a complex system and hard to know everything (too complex), experts disagree on causes (no consensus), it’s not important (unimportant), and other (with space for a free-text explanation). Finally, participants completed a brief demographic questionnaire.

Results

Influence of missing information expectations on perception We first analyzed the data to determine how different omissions influenced perceptions of diagrams. We created mean omission ratings by averaging across the topics presented as complete, missing most frequent, or missing least frequent for each participant. We did this separately for each of our 5 measures of interest (believability, compatibility, understandability, utility, and trust) We conducted a one-way ANOVA with omission type (complete, most, least) as a within-subjects variable for each of our 5 measures. We used Sidak-corrected follow-up tests to compare between the three groups. Across measures, we find that ratings are higher when diagrams omit the least frequent cause, followed by complete diagrams, with the diagrams missing the most frequent cause rated the least appealing for all variables, as shown in Figure 2. We now present analyses for each variable separately.

For believability, we found a main effect of omission type, $F(2, 298) = 15.8, p < .001, \eta_p^2 = .096$. While diagrams missing the least common cause were rated as more believable ($M = 4.72, SE = .084$) than complete diagrams ($M = 4.55, SE = .086$), this difference was not significant, $p = .130$. Least frequent diagrams were rated as more believable than diagrams missing the most frequent cause ($M = 4.23, SE = .096; p < .001$). Complete diagrams were also more believable than most frequent diagrams, $p = .003$.

We also found a significant main effect for compatibility, $F(2, 298) = 20.7, p < .001, \eta_p^2 = .122$. For these ratings, the diagram missing the least frequent cause ($M = 4.62, SE = .083$) was rated as more compatible than the complete ($M = 4.29, SE = .088$) or the most frequent cause diagram ($M = 4.01, SE = .100; ps < .001$). Complete diagrams were considered more compatible with people’s existing beliefs

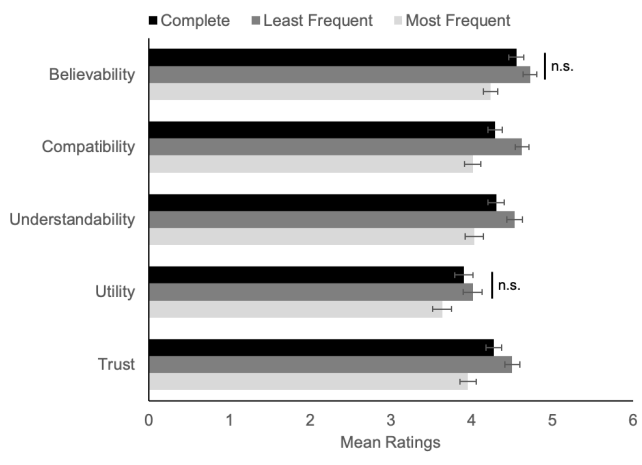


Figure 2: Experiment 1 mean perception ratings of diagrams by condition. Nonsignificant comparisons are marked. All other comparisons are significant at the $p = .05$ level

than most frequent diagrams, $p = .016$.

Findings for understandability were similar to compatibility. We found a significant main effect, $F(2, 298) = 14.2$, $p < .001$, $\eta_p^2 = .087$. The least frequent diagram ($M = 4.53$, $SE = .096$) was rated as more understandable than the complete ($M = 4.30$, $SE = .099$; $p = .038$) and the most frequent ($M = 4.03$, $SE = .113$; $p < .001$). Complete diagrams were more understandable than most frequent diagrams, $p = .015$.

Utility ratings showed a pattern similar to believability. We found a main effect of omission type, $F(2, 298) = 9.13$, $p < .001$, $\eta_p^2 = .058$. While diagrams missing the least common cause ($M = 4.01$, $SE = .116$) were rated as more useful than complete diagrams ($M = 3.90$, $SE = .116$), this difference was not significant, $p = .491$. Diagrams missing the most frequent cause ($M = 3.63$, $SE = .117$) were rated as having less utility than complete ($p = .016$) or least frequent diagrams, $p < .001$.

Finally, trust ratings again showed a preference for diagrams missing the least frequent cause. We found a significant main effect, $F(2, 298) = 18.4$, $p < .001$, $\eta_p^2 = .110$. The least frequent cause diagram ($M = 4.50$, $SE = .092$) was rated as more trusted than the complete ($M = 4.27$, $SE = .095$; $p = .021$) and the most frequent cause diagrams ($M = 3.95$, $SE = .101$; $p < .001$). Complete diagrams were trusted more than most frequent diagrams, $p = .004$.

Beliefs about reasons for omission We next explored the reasons people provided for why a cause was missing from a diagram. As a reminder, participants rated 2 to 3 diagrams for the most and for the least frequent missing cause diagrams. Participants could select the same or different answers for why a cause was omitted for each diagram. For each participant, we calculated the mean percentage of times they chose each of the five options within each of the omission types. As shown in Table 1, on average the most common response for both least and most frequent omission was that the system

Table 1: Reasons for missing information in Experiment 1. Numbers represent mean percentage of responses for each category.

Reason	Least freq	Most freq
Too complex	29.8%	29.3%
No consensus	29.6%	19.3%
Missing knowledge	18.3%	28.7%
Unimportant	11.0%	7.8%
Other	11.3%	15.8%

was too complex (least = 29.8%; most = 29.3%). In other words, people thought these were topics that were hard to know all possible causes for because of their complexity. For least frequent diagrams, a close second favorite response was that there was no expert consensus (29.6%). For most frequent diagrams, where a commonly reported cause was missing, the second most popular option was the designer missing knowledge (28.7%).

Discussion

Our results show that omissions can influence perception of diagrams in different ways depending on what information is omitted. When diagrams omit information participants likely expect to see (i.e., most frequently mentioned cause), the diagrams are rated lower in every category than complete diagrams. Yet when information that is unexpected is omitted, the resulting models are rated higher than both complete diagrams and most frequently missing ones. This suggests that participants are not only judging whether the complexity of a model matches the complexity of a topic (Lim & Oppenheimer, 2020), but also evaluating the information content.

In particular, participants thought the omission of expected information reflected more on the knowledge of the person creating the diagram, whereas the omission of less expected information was thought to suggest less consensus about the item's role. This has important implications for how choices about what to include in a diagram may influence the diagram's interpretation in ways their designers do not intend (e.g., leading to inferences about the designer's knowledge, or knowledge of the field). Critically, we find that omissions can both build trust and weaken it depending on the content that is omitted. Given prior work on how simplified models improve accuracy (Kleinberg & Marsh, 2021), this suggests a tradeoff between models people are willing to use and those they are able to use. Thus there is a need to find approaches to simplify models that do not diminish people's perceptions of their quality.

Experiment 2

Experiment 1 demonstrated that the content of omitted information influences how individuals perceive causal models. Models where a cause participants would expect to see was omitted were judged lower along all dimensions than complete models. On the other hand, removing information

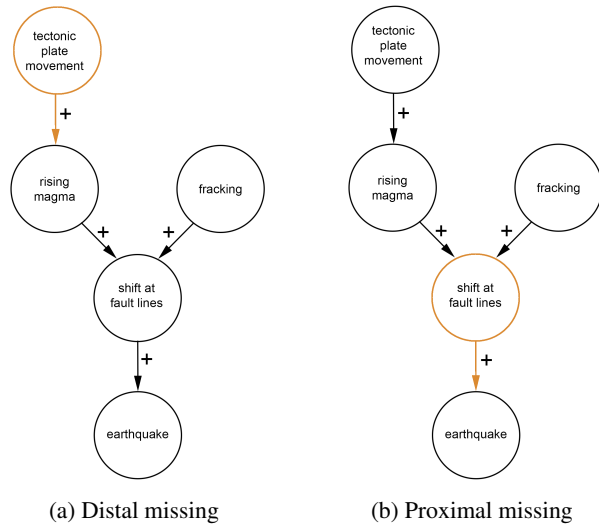


Figure 3: Diagrams used in Experiment 2. Both depict the highlighted diagrams shown to participants in the second half of the experiment, where we elicited their beliefs about reasons information is missing.

participants did not expect to see actually increased ratings. Importantly, in the previous experiment removing causes always removed an entire causal pathway from the overall diagram (for example the aging or illness pathway in Figure 1). Another way to simplify causal diagrams is to simplify the presented causal pathways, leaving out intermediate steps or capturing a multi-stage process with a single node. In Experiment 2 we now test whether leaving out nodes within the same causal pathway that are either closer or further away from the effect node changes people’s beliefs about the overall goodness of these diagrams. Prior work has shown that people often prefer to intervene on root nodes (Hagmayer & Sloman, 2009), which suggests that they may be more sensitive to the omission of such information. We now test this experimentally, varying the centrality of information omitted.

Method

Participants We recruited 150 U.S. residents aged 18-64 (39 female, 103 male, 5 identifying in other ways) from Prolific, with 149 completing. Participants were compensated \$3, as we found study actual duration to be around 20 minutes. All participants remain in analysis.

Materials In this experiment we aimed to manipulate the centrality of the cause omitted, leaving out a proximal or distal cause to determine whether centrality influences perceptions in the same way as salience. In Experiment 1 the diagrams used varied in structure, with two comprised solely of direct causes (therefore not possessing proximal versus distal causes). Thus to ensure all diagrams could be manipulated as intended we created a new set of diagrams for the same set of topics that all included longer chains of causes. Complete di-

Table 2: Reasons for missing information in Experiment 2. Numbers represent mean percentage of responses for each category.

Reason	Proximal	Distal
Missing knowledge	11.1%	12.7%
Too complex	21.1%	19.9%
No consensus	12.1%	9.3%
Unimportant	4.8%	9.7%
Implied presence	45.5%	42.6%
Other	5.4%	5.8%

agrams had 5 nodes, and 4 or 5 edges (causal relationships). For each topic we created two additional diagrams wherein we omitted either a distal cause (root node) or proximal cause (direct cause of the effect). For diagrams with two root nodes, we selected the more distal one (i.e., if one root has a chain with 3 links connecting it to the effect, and another root is connected to the effect by a 2 link chain, we omit the root beginning the 3 link chain). Figure 3 shows examples of the causes that were omitted for the diagram about earthquakes.

Procedure The procedure was the same as that of Experiment 1, with the only differences being in the diagrams used. All questions on diagram perception remained the same and we again randomly assigned participants to one of three conditions that varied which topic was paired with what type of diagram. When asking participants about reasons why information was missing we added one option based on a popular “other” response in Experiment 1: “it’s implied by/included in other parts of the diagram.” (implied presence).

Results and Discussion

We first examined how omitting a proximal versus distal cause influenced beliefs about the diagrams. As in Experiment 1, we created mean omission ratings for each of our 5 measures by averaging ratings across the topics that presented complete diagrams (complete), missing proximal causes (proximal), or missing distal causes (distal) for each participant. We conducted a one-way ANOVA with omission type (complete, proximal, distal) as a within-subjects variable for each of our five variables. Overall, participants’ ratings did not differ by condition. We found no significant main effect for believability ($p = .865$; complete $M = 4.25$, $SE = .092$; proximal $M = 4.29$, $SE = .097$; and distal $M = 4.24$, $SE = .086$), compatibility ($p = .577$; complete $M = 4.01$, $SE = .092$; proximal $M = 4.11$, $SE = .100$; and distal $M = 4.02$, $SE = .090$), understandability ($p = .212$; complete $M = 4.11$, $SE = .113$; proximal $M = 4.28$, $SE = .108$; and distal $M = 4.23$, $SE = .106$), utility ($p = .170$; complete $M = 3.57$, $SE = .109$; proximal $M = 3.74$, $SE = .113$; and distal $M = 3.59$, $SE = .105$), or trust ($p = .846$; complete $M = 3.95$, $SE = .096$; proximal $M = 3.94$, $SE = .105$; and distal $M = 3.90$, $SE = .093$).

We next turned to what reasons were selected for omis-

sions to understand why participants did not differentiate between missing diagram types. As in Experiment 1, we calculated the mean percentage of choices of each option for each omission type by participant, see Table 2. By a large margin, participants selected implied presence as the most popular reason for why a cause was missing for both proximal (45.5%) and distal (42.6%) causes. The popularity of this reason could help explain why participants did not differentiate their ratings based on the type of cause being omitted. Specifically, participants may have been assuming the presence of our missing causes even when they were not in the diagram. Taking the example in Figure 3, removing either a proximal or distal cause from the causal model of how an earthquake forms may not matter if people infer these parts of the causal mechanism regardless of whether they are directly presented. That is, even if a component is missing, people still feel the pathway is represented, and it appears not to make a significant difference which aspects of a specific path are included. Prior work found that mechanistic detail can be desirable for consumers choosing products with novel attributes (Fernbach, Sloman, Louis, & Shube, 2013), so it may be that familiarity is what allows such detail to be omitted without negative effects. Thus, in these familiar domains simplifying a causal model by removing components of a path, or potentially condensing the path (e.g., collapsing a process in one node summarizing it) may enable the creation of usable models that people also find trustworthy.

General Discussion

Across two experiments we found that omitting information from causal diagrams can both change and not change people's impression of those diagrams. When we removed root causes that were either commonly reported or uncommonly reported causes of an outcome, people were sensitive to this omission (Experiment 1). Specifically, people more favorably viewed simplified models that excluded an unexpected cause than complete models or simplified models that omitted an expected cause. However in that experiment the omitted information removed entire paths to an outcome. Omitting steps along a single causal pathway did not influence ratings (Experiment 2). Participants reported thinking that the omitted information was implied by what was shown, which may explain why the omissions did not change their perceptions of the diagrams.

Our work contributes new insights into the ongoing exploration of simplicity and complexity preferences. While there is evidence people prefer correspondingly complex explanations of complex phenomena (Lim & Oppenheimer, 2020) and that complex explanations of decision-support systems improved trust in the systems (Bussone et al., 2015), this work has primarily focused on comparing complex to simple causal models and has not pitted different simplifications against one another. To this end, we find that it is not simply the amount of information included that matters, but rather the type, and that people are more sensitive to the informa-

tion presented.

Our findings also have important implications for how we can provide better guidance in the form of causal models. Prior work has identified that simplified models lead to better choices (Kleinberg & Marsh, 2021), and we now shed light on how models can be simplified in ways that increase trust and belief in their evidence. First, we must realize that people do not always take a missing causal relationship to mean that said relationship is false. Officials creating guidance, like public health materials that explain how to slow the spread of a disease, may be tempted to include only the most accurate, relevant causes. While intuitively this seems like the correct approach, this may result in officials omitting causes that the lay public expects to be present. For example, if public health guidance on how you catch the common cold did not include "being out in cold air" people may not trust the information because that lay expected node is not present. This is not to say that public health guidance should support folk or incorrect information; rather, a better strategy may be to include the nodes but explicitly represent that they are not causally related. Just omitting such expected nodes may result in laypeople ignoring such guidance entirely. Second, models can be simplified so that they are not overwhelming by simplifying long causal chains. Korman and Khemlani (2020) found that people preferred a single complete model to multiple models, and we now find that removing nodes along a single path can reduce complexity without harming trust or other perceptions. As noted in philosophical theories of model completeness (Craver & Kaplan, 2020), the most complete model is not necessarily the one with the most detail, and adding more information does not always make a model better.

Acknowledgments

This work was supported in part by the NSF under award number 1907951 and the James S. McDonnell Foundation.

References

- Bussone, A., Stumpf, S., & O'Sullivan, D. (2015). The role of explanations on trust and reliance in clinical decision support systems. In *International Conference on Healthcare Informatics*.
- Chan, S. Y. (2001). The use of graphs as decision aids in relation to information overload and managerial decision quality. *Journal of information science*, 27(6), 417–425.
- Craver, C. F., & Kaplan, D. M. (2020). Are more details better? on the norms of completeness for mechanistic explanations. *The British Journal for the Philosophy of Science*, 71(1), 287–319.
- Fernbach, P. M., & Sloman, S. A. (2009). Causal learning with local computations. *Journal of experimental psychology: Learning, memory, and cognition*, 35(3), 678.
- Fernbach, P. M., Sloman, S. A., Louis, R. S., & Shube, J. N. (2013). Explanation fiends and foes: How mechanistic detail determines understanding and preference. *Journal of Consumer Research*, 39(5), 1115–1131.

- Hagmayer, Y., & Sloman, S. A. (2009). Decision makers conceive of their choices as interventions. *Journal of Experimental Psychology: General*, 138(1), 22–38.
- Hoff, K. A., & Bashir, M. (2015). Trust in automation: Integrating empirical evidence on factors that influence trust. *Human factors*, 57(3), 407–434.
- Johnson, S. G., Valenti, J., & Keil, F. C. (2017). Opponent uses of simplicity and complexity in causal explanation. In *Proceedings of the 39th annual meeting of the cognitive science society*.
- Kleinberg, S., & Marsh, J. K. (2021). It's complicated: Improving decisions on causally complex topics. In *Proceedings of the 43rd Annual Meeting of the Cognitive Science Society Society*.
- Korman, J., & Khemlani, S. (2020). Explanatory completeness. *Acta Psychologica*, 209, 103139.
- Liefgreen, A., & Lagnado, D. (2021). The role of causal models in evaluating simple and complex legal explanations. In *Proceedings of the 43rd annual meeting of the cognitive science society*.
- Lim, J. B., & Oppenheimer, D. M. (2020). Explanatory preferences for complexity matching. *PloS One*, 15(4), e0230929.
- Liquin, E. G., & Lombrozo, T. (2022). Motivated to learn: An account of explanatory satisfaction. *Cognitive Psychology*, 132, 101453.
- Lombrozo, T. (2007). Simplicity and probability in causal explanation. *Cognitive Psychology*, 55(3), 232–257.
- Pacer, M., & Lombrozo, T. (2017). Ockham's razor cuts to the root: Simplicity in causal explanation. *Journal of Experimental Psychology: General*, 146(12), 1761–1780.
- Stephan, S., Tentori, K., Pighin, S., & Waldmann, M. R. (2021). Interpolating causal mechanisms: The paradox of knowing more. *Journal of Experimental Psychology: General*, 150(8), 1500–1527.
- Zemla, J. C., Sloman, S., Bechlivanidis, C., & Lagnado, D. A. (2017). Evaluating everyday explanations. *Psychonomic Bulletin & Review*, 24(5), 1488–1500.
- Zheng, M., Marsh, J. K., Nickerson, J. V., & Kleinberg, S. (2020). How causal information affects decisions. *Cognitive Research: Principles and Implications*, 5(6), 1–24.