

Implicit questions shape information preferences

Sehrang Joo (sehrang.joo@yale.edu)

Department of Psychology, Yale University
New Haven, CT, 06520 USA

Sami R. Yousif (sami.yousif@yale.edu)

Department of Psychology, Yale University
New Haven, CT, 06520 USA

Frank C. Keil (frank.keil@yale.edu)

Department of Psychology, Yale University
New Haven, CT, 06520 USA

Abstract

We ask questions about everything from why clocks tick to why the sky is blue. Although people sometimes prefer teleological explanations over mechanistic explanations in response to ‘why’ questions, why questions are ambiguous—referring either to a ‘how’ question *or* a ‘for what purpose’ question. In this paper, we examine the relation between these *implicit* questions and explanation preferences. First, we asked whether people have specific expectations regarding ‘why’ questions: How do they interpret these ambiguous cases and does this vary across domains? Indeed, people have strong, domain-specific expectations that mirror well-documented explanation preferences. People also have preferences about which specific question they would prefer to have answered. In other words, ‘why’ questions are ambiguous but not treated as such — and this has consequences for downstream explanation preferences. We explore these consequences in light of both the philosophical and psychological literature on explanation.

Keywords: why questions; explanation; teleology; mechanism

Introduction

We can—and do—ask ‘why’ questions about all sorts of things that we encounter, both in our everyday lives and as scientists. Yet these questions are intrinsically ambiguous. Take, for example, a question as simple as “Why do trees have leaves?” The question may refer to either (a) *how* the leaves grew on the tree (a mechanistic explanation), or (b) the *purpose* or *function* behind trees having leaves (a teleological explanation). Despite this intrinsic ambiguity, questions like these may not seem ambiguous when we encounter them: You may automatically adopt some assumption about which question the agent was likely seeking or have a preference for which question you yourself would rather have answered. Here, we explore whether (and how) these expectations and preferences shape the explanation preferences that follow.

Teleological and Mechanistic Explanation

Both adults and children sometimes prefer teleological over mechanistic explanations. For instance, adults are found to prefer teleological explanations such as “The mononykus has a long tail so that it can keep its balance while it runs” over mechanistic explanations such as “The mononykus has a long tail because its feathers were big and stuck out from behind

its body” (Kelemen, 1999). Children, too, show a teleology bias (e.g., Schachner et al., 2017; Banerjee & Bloom, 2014; Kelemen, 1999) and endorse teleological explanations even in cases where such explanations seem *wrong* to most adults. For example, in response to the question, “Why are the rocks pointy?”, children prefer the teleological explanation, “The rocks are pointy so that animals won’t sit on them and smash them.”

These teleological biases are interpreted as part of a broad developmental theory: Children’s “promiscuous” teleology preferences are taken as evidence that teleological reasoning is an early-developing cognitive default—and possibly as evidence that children have intuitions about God(s) that lead them to view the world as intelligently designed (Kelemen, 2004; but see also ojaalehto et al., 2013 for an alternate proposal). Moreover, ‘promiscuous teleology’ suggests that a tendency to think teleologically is not only an early-developing bias but also a persistent one. Even adults (and, in some cases, even trained physicists; Kelemen et al. 2013) fall back on teleology when under time pressure or cognitive load (Kelemen & Rosset, 2009; Kelemen et al. 2013). As a result of an observed preference for teleological explanations, teleology is framed as a persistent but scientifically unwarranted bias that may “have subtle enduring effects on our species’ intellectual progress” by impeding our ability to advance scientific, mechanistic understanding (Kelemen et al., 2013, p. 1081).

Teleology: An information preference?

Despite a large body of work regarding the flaws of teleological explanation, children’s (and adults’) teleological biases need not be interpreted this way. Another, perhaps simpler, possibility is that their preferences may not be about explanation at all; instead, they may reflect the appeal of the kind of *information* presented in teleological explanations. Consider, for instance, a child encountering a household item, like a microwave, for the first time. She is far more likely to benefit from knowing the purpose of a microwave than from knowing its inner workings. In general, children especially may need to know not only how the world *does* work, but also how it *will* work.

Viewing the teleological preference as one for information rather than for explanation *per se* may help reconcile findings of ‘promiscuous teleology’ with cases where children seek relevant explanatory and mechanistic information (e.g., Greif et al., 2006). At once, children seem to understand (and seek) the kind of information that adults classify as scientifically warranted and yet also have a very poor understanding of what constitutes a valid explanation. If, however, children are simply most interested in first learning about teleological information, then perhaps there is an intelligent basis for children’s endorsement of teleological explanations—and, by extension, of adults’ teleological biases, as well. These preferences might simply be a reflection of what we desire to know most, and therefore of a bias for *information* rather than for *explanation*.

In contrast to views dismissing teleology as a “scientifically unwarranted” bias (Kelemen & Rosset, 2009; Kelemen et al., 2013), we ask whether adults might remain intelligent information seekers *even as* they are drawn to teleology. This perspective shift would flip our understanding of a teleological bias from ‘incorrect default’ to ‘rational tendency’ and help reconcile such biases with rational scientific inquiry. Further, recharacterizing adults’ teleological biases in this way would also suggest a way to recharacterize children’s ‘promiscuous’ teleology.

Current Study

The current study is organized with the following general questions in mind: (1) Are there distinct implicit questions (e.g., ‘how’ and ‘purpose’ questions) within ambiguous ‘why’ questions?; (2) Are adults sensitive to those implicit questions when reasoning about ‘why’ questions and their answers?; and (3) Might these questions shape adults’ preferences, prior to any explanations?

We begin in the simplest way possible, by directly asking people what specific question is implied by various ‘why’ questions (Experiment 1). From there, we use a novel ‘jeopardy paradigm’ and examine what questions are implied by teleological and mechanistic explanations (Experiment 2). We then use these implicit ‘how’ and ‘purpose’ questions to investigate people’s information preferences (Experiment 3). We submit that not only *are* there implicit questions within ‘why’ questions, but that teasing apart this ambiguity also reveals that people’s teleological preferences may exist long before they encounter an explanation.

Experiment 1: What did they want to know?

‘Why’ questions are ambiguous—but do people assume that they imply specific questions in particular contexts? Here, we showed participants an ambiguous ‘why’ question and asked them what the agent asking the question really wanted to know: a mechanistic ‘how’ question or a teleological ‘purpose’ question.

Method

Participants One hundred adult participants completed a survey online through Amazon Mechanical Turk. The sample

size was chosen on the basis of independent pilot data and was preregistered. All participants lived in the United States.

Stimuli Each survey item consisted of a ‘why’ question (e.g., “Why does the mononykus have such a long tail?”), a ‘how’ (e.g., “How did the mononykus’ tail become long?”) and a ‘purpose’ (e.g., “What is the purpose of the mononykus’ long tail?”). Total materials consisted of twelve such sets of questions, four each in the domain of animals, non-living natural kinds (NLNK), and artifacts. These stimuli were adapted from Kelemen (1999), with a few notable modifications, including: (1) ‘How’ and ‘purpose’ questions (rather than teleological and mechanistic explanations) were used in this experiment; (2) One additional domain (artifacts) was added, to include the three domains most commonly investigated in the context of teleological explanations (e.g., Lombrozo & Gwynne, 2014; Lombrozo & Carey, 2006); (3) Stimuli were presented in a fully randomized order (rather than in pairs). See Table 1 under “‘Why’ Question” and “Questions” for example stimuli.

Procedure All participants saw all twelve items (in a different random order for each participant). In each case, participants were simply asked what the agent asking the question “really want[ed] to know”. They then chose between a ‘how’ question (e.g., “How did the mononykus’ tail become long?”) and a ‘purpose’ question (e.g., “What is the purpose of the mononykus’ long tail?”). The questions themselves were also presented in a random order within each question for each participant. No other information was collected. Data, materials, and preregistration information for this experiment and all following can be found on the Open Science Framework (OSF) [here](#).

Results and Discussion

The results of Experiment 1 are shown in Figure 1. When presented with ‘why’ questions about animals and artifacts, a significant proportion of participants thought the agent was really asking a ‘purpose’ question (animals: $p=.86$; artifacts: $p=.95$), $p<.001$ (binomial tests). In contrast, a significant proportion of participants chose the ‘how’ question for non-living natural kinds, ($p=.87$), $p<.001$.

In following with the work on which we most closely based our stimuli (Kelemen, 1999), we next also analyzed participants’ expectations by treating participants’ responses across items in the same domain as an average. Participants’ responses were scored as a 1 if they chose the ‘purpose’ question and a 0 if they chose the ‘how’ question. We then added these scores within each domain for each participant. We again found that participants thought that ‘why’ questions implied ‘purpose’ questions when asked about animals ($M=3.44$, $SD=1.02$, $t(99)=14.14$, $p<.001$, $d=1.41$), and artifacts ($M=3.79$, $SD=.64$, $t(99)=27.96$, $p<.001$, $d=2.80$). For non-living natural kinds, participants instead thought that ‘why’ questions implied ‘how’ questions ($M=.53$, $SD=.89$, $t(99)=16.47$, $p<.001$, $d=1.65$). Importantly, these expectations mirror people’s explanation preferences: People

Table 1: One full example from each domain is shown. Full stimuli are available at our OSF page [here](#).

'Why' Question	Explanations	Questions
[Animal] Why does the mononykus have such a long tail?	The mononykus has a long tail because its feathers were big and stuck out from behind its body. The mononykus has a long tail so that it can keep its balance when it runs.	How did the mononykus' tail become long? What is the purpose of the mononykus' long tail?
[Non-living natural kind] Why are the rocks so pointy?	The rocks are pointy because little bits of stuff piled up on top of one another over a long time. The rocks are pointy so that animals won't sit on them and smash them.	How did the rocks become pointy? What is the purpose of the rocks being pointy?
[Artifact] Why is the baking tool full of holes?	The baking tool is full of holes because it was cut with something sharp. The baking tool is full of holes so that it can hook onto small parts on a hot oven rack.	How did the baking tool's holes form? What is the purpose of the baking tool's holes?

thought that 'why' questions implied 'purpose' questions in domains where they prefer teleological explanations and 'how' questions in domains where they prefer mechanistic explanations. (For the purpose of comparing preferences, we ran an additional experiment, not reported here, that investigated people's base explanation preferences. These results are also shown in Figure 1.)

Finally, we also explicitly compared participants' expectations between domains by using a repeated-measures ANOVA. There was a significant main effect of domain, $F(2,99)=434.21$, $p<.001$. Post-hoc tests demonstrated that participants' were more likely to assume the 'why' question was implying a 'purpose' question for animals than for non-living natural kinds, $t(99)=20.67$, $p<.001$, Bonferroni corrected, $d=2.07$. Similarly, participants were also more likely to assume the 'why' question was implying a 'purpose' question for artifacts than for non-living natural kinds, $t(99)=25.87$, $p<.001$, Bonferroni corrected, $d=2.59$. While participants were also more likely to expect the 'purpose'

questions for animals than for artifacts, these effects were much smaller, $t(99)=3.78$, $p<.001$, Bonferroni corrected, $d=.38$. People not only assume that a 'why' question is actually seeking some more specific information, but also that the implied question can be broken down in a domain-specific way.

In other words, while the question that is being asked is the identical general 'why' question, the question that is pragmatically implied can be broken down into more specific 'how' and 'purpose' questions. Thus, when first encountering a 'why' question, people may already have an expectation about the more specific kind of information being requested—and, as a result, perhaps also about the kind of answer that would adequately address that question. This expectation raises a key question: How much of people's explanation preferences are truly about *explanation*, and how much might actually be driven by the questions or the kind of information promised by the question.

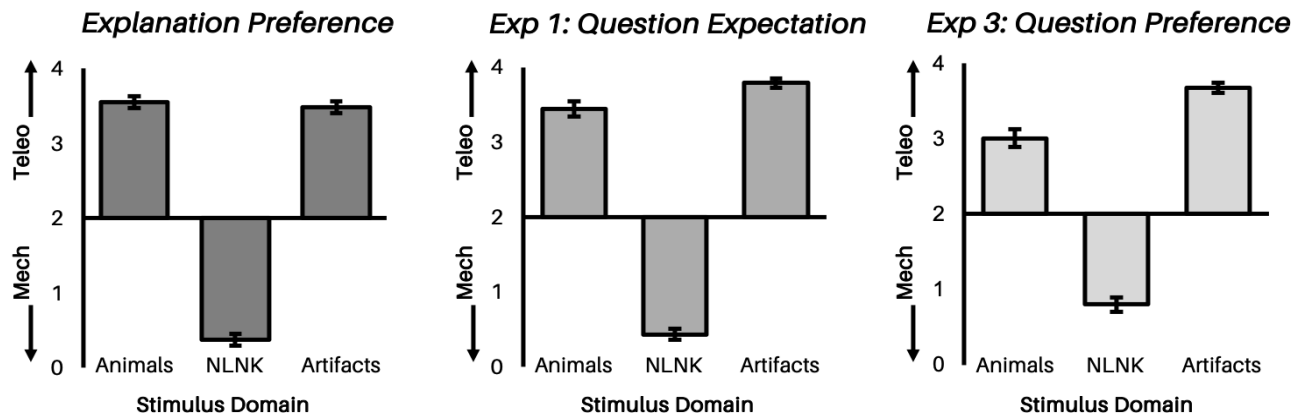


Figure 1: Results demonstrating a teleological preference for animals and artifacts. From left to right: (1) an unreported experiment (see OSF page [here](#)) that established people's explanation preferences; (2) Experiment 1; (3) Experiment 3. The x-axis represents chance performance, and domain of the stimulus is represented along the axis. Participants' scored responses are represented along the y-axis. In all cases, participants' preferred the teleological choice for animals and artifacts but the mechanistic choice for NLNK (non-living natural kinds). Error bars represent +/- 1 SE.

Experiment 2: ‘Jeopardy’

Do people also think that teleological and mechanistic explanations *result* from different questions? Here, we had participants complete a novel ‘jeopardy’ task. Instead of asking participants to choose between two different explanations, we provided them with the explanations and asked them what *question* likely led to that answer.

Method

Participants One hundred new adult participants completed a survey online through Amazon Mechanical Turk (data from 7 additional participants were collected but excluded for failing the training item; data from 5 additional participants were collected but excluded for failing to follow directions or produce interpretable questions; see the Procedure section). This sample size was chosen to be identical to those in the previous experiment. This experiment was also preregistered.

Stimuli Total materials consisted of twenty-four explanations; twelve were mechanistic explanations and twelve were teleological explanations. These explanations were adapted from Kelemen (1999) and corresponded to the questions used in Experiment 1. See Table 1 under “Explanations” for example stimuli.

Procedure Participants were assigned to one of two between-subjects conditions. Participants in each condition saw twelve explanations, each about a different item. One group saw two explanations of each kind (mechanistic and teleological) in each domain (animals, non-living natural kinds, artifacts). The other group saw the exact opposite set of explanations, flipping which items were shown with mechanistic vs. teleological explanations. The goal of these between-subjects conditions was simply for participants to only see one explanation of any given item, and there were no other differences between conditions. Items were presented in a different random order for each participant.

Participants first saw a training item that introduced the ‘jeopardy’ task. They were told that instead of being shown questions and answering them, they would be shown *answers* and be asked to generate the questions that led to those answers. As an example, they were given the answer “The book is over there, on top of the shelf but underneath the scarf” and asked to give a question that might have led to this answer. Participants who did not provide a question along the lines of “Where is the book?” were excluded and replaced for failing to understand the task (see the Participants section).

In each of the test items, participants were shown an image of the item and were told that somebody had provided an answer about that item (e.g., “Someone answered, ‘The mononykus has a long tail because its feathers were big and stuck out from behind its body.’”). They were asked, “What question might have led to this answer?” and were explicitly prompted not to use the word “why”. No other information was collected.

Participants who either (1) systematically generated non-interpretable questions or (2) ignored the direction not to use

the word “why” were excluded and replaced (see the Participants section). Questions were then coded into ‘purpose,’ ‘how,’ and ‘other’ categories. A question was coded as a ‘purpose’ question if it contained the words “purpose,” “for,” “use,” “function”, or “advantage.” A question was coded as a ‘how’ question if it contained the words “how”, “cause”, “made”, and “led.” All criteria for coding questions were preregistered and are available on our OSF page [here](#).

Results and Discussion

The results of Experiment 2 are shown in Figure 2. People thought that teleological explanations resulted from predominately ‘purpose’ questions and that mechanistic explanations resulted from predominately ‘how’ questions.

Participants’ questions were not evenly distributed across categories, $X^2(5, N=923)=166.37, p<.001$. The same was true when analyzing questions within each domain; whether thinking about answers regarding animals, ($X^2(5, N=317)=35.26, p<.001$), non-living natural kinds, ($X^2(5, N=286)=121.87, p<.001$), or artifacts, ($X^2(5, N=320)=116.58, p<.001$) participants generated different kinds of questions for different kinds of explanations.

In keeping with prior work, we again also analyzed participants’ likelihood of generating ‘how’ vs. ‘purpose’ questions (for both mechanistic and teleological explanations) by treating participants’ questions across items in the same domain as an average. ‘How’ questions were scored as a -1, ‘purpose’ questions were scored as a 1, and ‘other’ questions were scored as a 0. We then added these scores for each participant across all teleological explanations and, separately, across all mechanistic explanations. When presented with mechanistic explanations, participants were more likely to generate ‘how’ questions, ($M=-1.74, SD=2.35, t(99)=7.40, p<.001, d=.74$). Conversely, when presented with teleological explanations, participants were more likely to generate ‘purpose’ questions, ($M=1.42, SD=2.15, t(99)=6.60, p<.001, d=.66$). The difference between participants’ questions in these two different explanation types was significant, $t(99)=10.13, p<.001, d=1.01$. This pattern was independently true within each domain, all $p<.001$, all $d>.5$.

These results suggest that people interpret teleological explanations as answers to ‘purpose’ questions and mechanistic explanations as answers to ‘how’ questions. When encountering explanations, people therefore have strong assumptions about what kinds of questions those explanations could plausibly be answering.

Taken together, Experiments 1 and 2 suggest that reasoning about what constitutes a better explanation might consist of reasoning about (1) what sort of question is being asked and (2) what kind of answer satisfies that question. In cases where adults endorse teleological explanations, for instance, they both think that a ‘why’ question is *really* seeking information about something’s purpose and that a teleological explanation would answer that ‘purpose’ question. These explanations might therefore not be better

explanations—in which case they would be taken as somehow more robustly explanatory—but simply ‘better’ in virtue of being able to answer the relevant question.

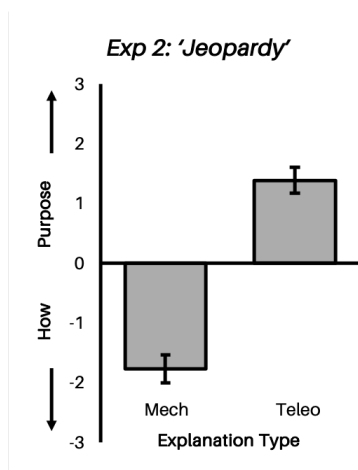


Figure 2. The x-axis represents chance performance, and explanation type is represented along the axis. Participants’ coded questions are represented along the y-axis, such that bars above the x-axis represent a tendency to generate ‘purpose’ questions for that explanation type. Error bars represent ± 1 SE.

Experiment 3: Explanation vs. Information

Given that people both think ‘why’ questions imply more specific questions *and* that teleological and mechanistic explanations differentially address these more specific questions, could people simply have a preference for teleological *information* (i.e., information about something’s purpose)? If people have a preference driven by these implicit questions, then their expectations about the fit between a given ‘why’ question and a possible answer may constitute an information preference that operates prior to considering explanations. Here, we presented participants with a ‘how’ question and a ‘purpose’ question and simply asked them, “Which of these questions would you rather have answered?”

Method

All elements of the experimental design were identical to those of Experiment 1, except as stated below. One hundred new participants completed the survey online through Amazon Mechanical Turk. This sample size was chosen to be identical to that in the previous experiment. This experiment was also preregistered.

Participants were now asked to choose between a ‘how’ and a ‘purpose’ question, on the basis of which they would rather have answered. No ‘why’ questions were given; for each item, participants saw a picture of the item and two possible questions (i.e., a ‘how’ and a ‘purpose’ question) to choose from. See Table 1 for example stimuli.

Results and Discussion

The results to Experiment 3 are shown in Figure 1. When presented with animals and artifacts, a significant proportion of participants preferred ‘purpose’ questions, (animals: $\hat{p}=.75$; artifacts: $\hat{p}=.92$), $p<.001$ (binomial tests). In contrast, a significant proportion of participants preferred ‘how’ questions for non-living natural kinds, ($\hat{p}=.80$), $p<.001$ (binomial test).

We also analyzed participants’ expectations by treating participants’ responses across items in the same domain as an average. Participants’ responses were scored as a 1 if they chose the ‘purpose’ question and a 0 if they chose the ‘how’ question. We then added these scores within each domain for each participant. We again found that participants significantly preferred ‘purpose’ questions when asked about animals, ($M=3.01$, $SD=1.12$, $t(99)=8.99$, $p<.001$, $d=.90$) and artifacts, ($M=3.68$, $SD=.65$, $t(99)=25.87$, $p<.001$, $d=2.59$). For non-living natural kinds, participants significantly preferred mechanistic explanations, ($M=.80$, $SD=.99$, $t(99)=12.19$, $p<.001$, $d=1.22$). Finally, we also explicitly compared participants’ expectations between domains by using a repeated-measures ANOVA. There was a significant main effect of domain, $F(2,99)=245.62$, $p<.001$. Post-hoc tests demonstrated that participants were more likely to choose the ‘purpose’ question for animals than for non-living natural kinds, $t(99)=13.32$, $p<.001$, Bonferroni corrected, $d=1.33$. Similarly, participants were more likely to choose the ‘purpose’ question for artifacts than for non-living natural kinds, $t(99)=21.21$, $p<.001$, Bonferroni corrected, $d=2.12$. While participants were also more likely to choose the ‘purpose’ question for artifacts than for animals, these effects were much smaller, $t(99)=6.87$, $p<.001$, Bonferroni corrected, $d=.69$.

Importantly, these patterns of results mirror adults’ explanatory preferences (see Figure 1). In other words, in the cases where adults exhibit a teleological explanation preference, they *also* prefer ‘purpose’ questions. This distinction is subtle but critical; it suggests that adults’ preferences for teleology may exist long before they receive or consider an explanation. When people prefer teleological explanations, their explanatory preference may be the downstream consequence of a preference for certain questions (i.e., ‘purpose’ questions)—or by a preference for the kind of *information* that it would take to answer that question.

General Discussion

The experiments reported here demonstrate that despite their intrinsic ambiguity, ‘why’ questions are often not interpreted as such: Implicit questions within a ‘why’ question influence evaluations of both the question itself and the explanations that it prompts. In Experiment 1, people assumed ‘why’ questions imply more specific ‘how’ and ‘purpose’ questions in different contexts. In Experiment 2, through a ‘jeopardy’ paradigm, people also thought that *explanations themselves* imply more specific questions, with teleological explanations

implying ‘purpose’ questions and mechanistic explanations implying ‘how’ questions. In Experiment 3, people preferred the kinds of questions that mirror their explanatory preferences. These experiments collectively demonstrate that people seem to have (1) an expectation about what question (‘how’ vs. ‘purpose’) is implied when an agent asks a ‘why’ question, (2) an expectation about what kind of explanation (mechanistic vs. teleological) would actually answer the relevant implicit question, and (3) a preference for some questions (implied or otherwise) over others in the first place.

These results suggest that there may be an alternative way to understand people’s preference for a given explanation. Despite the ambiguity of the ‘why’ questions they regularly ask (and encounter), people have strong expectations about the kinds of questions that they *implicitly* seek (and answers they receive). If their explanatory preferences reflect a downstream consequence of these expectations, then people may exhibit an *information* preference—one which tracks teleology preferences without presuming that people find teleological explanations to be robustly explanatory. Future research should therefore address the possibility that people’s preferences may not be driven primarily by their understanding of causal explanation.

Such a view may help reconcile cases where even adults are ‘promiscuously’ teleological (e.g., Kelemen & Rosset, 2009; Kelemen et al. 2013) with cases where adults are sensitive to the causal relevance of various explanations (e.g., Lombrozo & Carey, 2006; see also Liquin & Lombrozo, 2008). Adults may be drawn to teleological questions and answers particularly under speeded conditions, but override these information preferences when explicitly reasoning about the merits of teleology as an explanation. If people have a preference for teleology as a kind of *information*, then endorsement of teleological explanations may be a result of this more general preference and not of a mistaken tendency to view something’s purpose as genuinely causal. On this view, while teleological explanations are not always causal explanations, they may be independently alluring simply because they answer the most relevant question at hand.

This perspective need not be at odds with prior research (e.g., Kelemen et al., 2013; Rose & Schaffer, 2015) that correlates people’s endorsement of teleological explanations with the degree to which they seem to think of nature as the product of an intelligent designer. The critical suggestion of an information preference would be that endorsement of teleology and beliefs about intelligent design can be teased apart; while beliefs in intelligent design may increase endorsement of teleology, endorsement of teleology need not imply anything about whether one believes in intelligent design. Instead, individuals’ teleological biases may result from a combination of factors.

Importantly, for teleological *information* to be valuable, it need only increase understanding across some interesting or useful dimension. One promising possibility is that teleology may increase understanding by offering information on how the explanandum fits into a larger worldview. In contrast to mechanistic explanations (which generally provide

information about how an explanandum’s parts fit together to construct it), teleological explanations may be best understood as offering information about how the explanandum itself fits into some wider picture. For instance, learning that an animal’s tail is good for keeping its balance helps place the tail in the context of the animal as a whole. Similarly, learning about the functions of everyday objects like microwaves or unfamiliar tools may also help contextualize them in the world more generally.

Re-characterizing the teleology preference as centered around information would also be significant in reframing the larger theoretical proposals made about teleological biases. We argue that adults exhibit an information preference: Might children be similarly motivated? If children are like adults, then their endorsement of explanations such as “The rocks are pointy so that the animals won’t sit on them and smash them” need not result from thinking the rocks were designed that way by a creator. Rather, the information provided in the teleological explanation may simply be better-suited to helping them understand a largely unfamiliar world, perhaps by establishing a broad connection between things-which-are-pointy and things-which-are-not-sat-on.

Take the very simple example of encountering an unfamiliar and complicated-looking machine. Learning that the machine is for making coffee allows you to characterize it in terms of more familiar objects (e.g., a simple coffee pot), predict what it will do (e.g., produce a cup of coffee), and interact meaningfully with it. Such information is highly useful in contextualizing the unfamiliar. Teleological explanations are also found to be more generalizable than other kinds of explanations, such that learning about *this* coffee machine’s function is highly generalizable knowledge to coffee machines at large (Lombrozo & Gwynne 2014; see however Lockhart et al., 2019 and Chuey et al. 2020 on how mechanistic *expertise* seems more generalizable). Such research suggests a mechanism by which teleological information is clearly useful. Importantly, such information might be particularly useful to a child learning to navigate a largely unfamiliar world. Children in particular may like teleological explanations simply because they address the questions they care most about, and teleology itself may be a useful heuristic for learning about the world.

In short, we demonstrate that while our simplest, most fundamental, and yet most interesting question, “Why?”, is ambiguous—it is not interpreted as such. This simple fact shapes the way that we understand not only such ‘why’ questions, but also the answers that follow them. Just as we expect a certain *kind* of answer when asking someone who looks like they might be crying, “How are you feeling?” (e.g., not an answer like “I’m cold”), ambiguous ‘why’ questions are often asked with a particular sort of answer in mind. These expectations have the potential to reshape our understanding of people’s explanation and information preferences alike, suggesting that seeking to understand the world around us must begin by understanding the questions that we ask about it.

Acknowledgements

For helpful comments and conversation we thank Joshua Knobe, Paul Franks, Nicole Betz, Richard Ahl, Brynn Sherman, Katie Coyne, and all members of the Yale Cognition and Development Lab. This project was supported by a National Science Foundation Graduate Research Fellowship awarded to S. R. Yousif.

natural phenomena. *Journal of Experimental Child Psychology*, 157.

References

- Banerjee, K., & Bloom, P. (2014). "Everything Happens for a Reason": Children's Beliefs About Purpose in Life Events. *Child Development*, 86.
- Chuey, A., Lockhart, K., Sheskin, M., & Keil, F.C. (2020). Children and adults selectively generalize mechanistic knowledge. *Cognition*, 199, 104231.
- Greif, M. L., Nelson, D. G. K., Keil, F. C., & Gutierrez, F. (2006). What Do Children Want to Know About Animals and Artifacts? *Psychological Science*, 17.
- Kelemen, D. (1999). Why are rocks pointy? Children's preference for teleological explanations of the natural world. *Developmental Psychology*, 35.
- Kelemen, D. (2004). Are Children "Intuitive Theists"? Reasoning About Purpose and Design in Nature. *Psychological Science*, 15.
- Kelemen, D., & DiYanni, C. (2005). Intuitions About Origins: Purpose and Intelligent Design in Children's Reasoning About Nature. *Journal of Cognition and Development*, 6.
- Kelemen, D., & Rosset, E. (2009). The Human Function Compunction: Teleological Explanation in Adults. *Cognition*, 111.
- Kelemen, D., Rottman, J., & Seston, R. (2013). Professional Physical Scientists Display Tenacious Teleological Tendencies: Purpose-Based Reasoning as a Cognitive Default. *Journal of Experimental Psychology: General*, 142.
- Liquin, E. G., & Lombrozo, T. (2018). Structure-function fit underlies the evaluation of teleological explanations. *Cognitive Psychology*, 107.
- Lombrozo, T., & Carey, S. (2006). Functional explanations and the function of explanation. *Cognition*, 99.
- Lombrozo, T., & Gwynne, N. Z. (2014). Explanation and inference: mechanistic and functional explanations guide property generalization. *Frontiers in Human Neuroscience*, 8, 700.
- Lockhart, K. L., Chuey, A., Kerr, S., & Keil, F. C. (2019). The privileged status of knowing mechanistic information: An early epistemic bias. *Child Development*, 90.
- ojalehto, b., Waxman, S.R., and Douglas L.M. (2013). Teleological Reasoning About Nature: Intentional Design or Relational Perspectives? *Trends in Cognitive Sciences*, 17.
- Rose, D., & Schaffer, J. (2015). Folk Mereology is Teleological. *Noûs*, 51.
- Schachner, A., Zhu, L., Li, J., & Kelemen, D. (2017). Is the bias for function-based explanations culturally universal? Children from China endorse teleological explanations of