

Explaining to Others Prompts Children to Favor Inductively Rich Properties

Caren M. Walker¹, Tania Lombrozo¹, Cristine H. Legare² & Alison Gopnik¹

(caren.walker@berkeley.edu, lombrozo@berkeley.edu, legare@psy.utexas.edu, gopnik@berkeley.edu)

¹Department of Psychology, University of California, Berkeley, 3210 Tolman Hall, Berkeley, CA 94720 USA

²Department of Psychology, University of Texas at Austin, 1 University Station #A8000, Austin, TX 78712 USA

Abstract

Three experiments test the hypothesis that engaging in explanation prompts children to favor inductively rich properties when generalizing to novel cases. In Experiment 1, preschoolers prompted to explain during a causal learning task were more likely to override a tendency to generalize according to perceptual similarity and instead extend an internal feature to an object that shared a causal property. In Experiment 2, we replicated this effect of explanation in a case of label extension. Experiment 3 demonstrated that explanation improves memory for internal features and labels, but impairs memory for superficial features. We conclude that explaining can influence learning by prompting children to favor inductively rich properties over surface similarity.

Keywords: Explanation; causal learning; category labels; non-obvious properties; inductive inference

Introduction

The world has a complex structure, and the challenge of causal learning is to discover the nature of this structure to facilitate prediction and action. This is not a trivial task; it is sometimes impossible to predict how an object will behave based on its appearance. In fact, perceptually similar objects can be endowed with very different causal properties: Poison hemlock may *look* identical to wild carrot, but it is certainly not good to eat. Learning to override perceptual features in favor of non-obvious but inductively rich properties is thus an important achievement.

Previous research has examined the role of obvious (perceptual) properties versus non-obvious (internal or abstract) properties in children's inferences. Young children can use both perceptual and non-perceptual properties in categorizing objects (e.g., Gelman & Markman, 1987; Gopnik & Sobel, 2000), but adults often group objects according to common internal properties, labels, and causal affordances (regardless of perceptual similarity) in cases where young children tend to group objects based on salient perceptual similarity.

To illustrate this shift, consider the findings from Nazzi and Gopnik (2000). Children observed four objects placed on a toy, one at a time. Two of these objects were shown to be causally effective – they made the toy play music – and two were inert. One of the causal objects was then held up and labeled (e.g., “*This is a Tib.*”), and children were asked to give the experimenter the other “*Tib.*” In no-conflict trials, perceptual and causal properties were always correlated. However, in conflict trials, the same perceptual properties appeared across causal and inert objects. All children were more likely to choose the causal object in the no-conflict trials than in the conflict trials, but analyses of

conflict trials revealed a developmental shift: when generalizing the novel label, 3.5-year-olds relied on perceptual cues over causal cues, while 4.5-year-olds relied on causal cues over perceptual cues.

Subsequent work has demonstrated a comparable shift in generalizing *internal parts* (as opposed to a category label). Sobel et al. (2007) used a similar procedure to demonstrate that 4-year-olds, but not 3-year-olds, are more likely to infer that objects have shared internal parts when they share causal properties than when they share external appearance.

These examples – and many others (see Keil, 1989; Gelman, 2003) – demonstrate that by 5 years of age, children begin to favor inductively rich but subtle cues, such as category membership and internal parts, over perceptual similarity when generalizing from known to unknown cases. But how is this transition achieved? Here we explore the hypothesis that the process of generating *explanations* is an important mechanism in scaffolding this transition.

Explanation and Causal Learning

Accounts of explanation from both philosophy and psychology suggest an important relationship between explanation and causal learning: By explaining past observations we uncover information likely to support future judgments and interventions (e.g., Lombrozo, 2012; Walker, Williams, Lombrozo, & Gopnik, 2012).

Consistent with this idea, research with adults finds that prompts to explain can improve learning (e.g., Chi et al., 1994) and promote the discovery and application of broad generalizations underlying what is being explained (e.g., Williams & Lombrozo, 2010). Prior research also suggests that even young children's explanations have characteristics that make them well suited to highlighting inductively rich properties: they often invoke broad generalizations (Walker et al., 2012) and go beyond appearances (Legare, 2012).

For example, Walker et al. (2012) found that prompting preschool-aged children to explain causal events made them more likely to favor broad patterns in generalizing causal properties to novel objects. In the first of these studies, children were presented with evidence that was consistent with two candidate causes (e.g., “green objects make the toy go” versus “yellow objects make the toy go”), where one accounted for more observations. Children who were prompted to explain were more likely than controls to generalize according to the candidate cause that accounted for more of the data. In the second study, the cause that accounted for more of the data was contrasted with an alternative cause that was more consistent with children's

prior knowledge (e.g., “large blocks make the toy go”). In this case, those who explained were *less* likely to generalize according to the cause that accounted for more of the evidence, and instead privileged their prior knowledge.

Additionally, young children’s explanations often appeal to non-perceptual properties, including unobservable causes (Legare, Wellman & Gelman, 2009) and labels (Legare, Gelman, & Wellman, 2010), and studies find that prompting children to explain can lead them to favor causal over perceptual learning (Legare & Lombrozo, under review).

We therefore predict that by encouraging learners to consider broad generalizations, explaining can encourage learners to go beyond appearances to favor non-obvious but inductively rich properties as a basis for generalization.

In the following experiments, we use a method similar to Nazzi and Gopnik (2000) and Sobel et al. (2007) to examine whether generating explanations prompts children to favor generalizing internal parts (Experiment 1) and labels (Experiment 2) on the basis of causal over perceptual similarity. In Experiment 3, we examine whether the effects of explanation derive from a special relationship between explanation and inductively rich properties, or from a global boost in performance, as might be expected if explaining simply increased attention. Together, these experiments shed light on the role of explanation in the construction of generalizations that support causal inference.

Experiment 1

Experiment 1 examined whether explanation influenced preschoolers’ extension of a hidden, internal property to other objects that shared either perceptual or causal properties. Children observed four sets of three objects individually placed on a toy that played music when “activated” (see Gopnik & Sobel, 2000). Each set contained three objects: one that activated the toy (*target object*), one that was perceptually identical to the *target object*, but failed to activate the toy (*perceptual match*), and one that was perceptually dissimilar to the *target object*, but successfully activated the toy (*causal match*). After each outcome was observed, children were asked to either explain (*explain* condition) or report (*control* condition) that outcome. Next, children received additional information about the target object: an internal part was revealed. Children were asked which one of the two other objects in the set (i.e., the *perceptual match* or *causal match*) shared the internal property with the *target object*. This method pit highly salient perceptual similarity against shared causal properties; children could base their generalizations on either one, but not both. We hypothesized that children who were asked to explain each outcome would be more likely than children in the control condition to consider the property with the greatest inductive richness and therefore select the *causal match* over the *perceptual match*.

Method

Participants A total of 108 children were included in Experiment 1, with 36 3-year-olds ($M = 40.9$ months; $SD =$

3.7, range: 35.8 – 47.7), 36 4-year-olds ($M = 53.3$ months; $SD = 3.6$, range: 48.5 – 59.8), and 36 5-year-olds ($M = 64.4$ months; $SD = 3.0$, range: 60.1 – 70.4). Eighteen children in each age group were randomly assigned to each of the two conditions (*explain* and *control*).

Materials The toy was similar to the “blicket detectors” used in past research on causal reasoning (Sobel & Gopnik, 2000), and consisted of a 10” x 6” x 4” opaque cardboard box containing a wireless doorbell that was not visible to the participant. When an object “activated” the toy, the doorbell played a melody. The toy was in fact surreptitiously activated by a remote control.

Twelve wooden blocks of various shapes and colors were used (see Fig. 1). A hole was drilled into the center of each block. Eight blocks contained a large red plastic map pin glued inside the hole; the remaining four blocks were empty. All of the holes were covered with a dowel cap, which covered the opening to conceal what was inside. Each of the four sets of blocks was composed of three individual blocks. Two were identical in color and shape. One of these blocks (the *target object*) contained a map pin. The other of these blocks (the *perceptual match*) did not. The third block (the *causal match*) was perceptually dissimilar to the other two, and, like the *target object*, it contained a map pin.

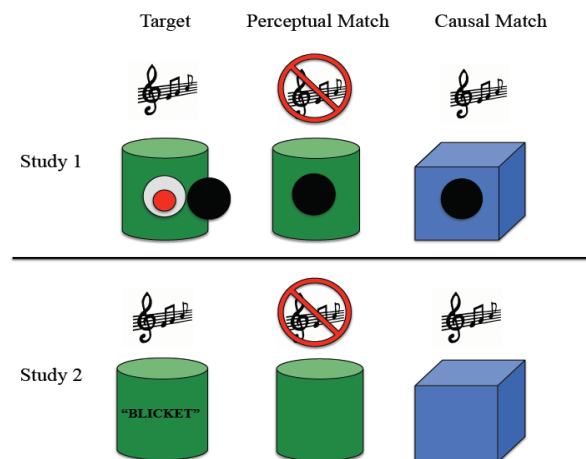


Figure 1: Sample set of objects used in Experiment 1 (top) and Experiment 2 (bottom). Each row corresponds to a single set of items. There were a total of four sets.

Procedure Children participated in a brief warm-up game with the experimenter. Following this warm-up, the toy was placed on the table. The child was told, “This is my toy. Some things make my toy play music and some things do not make my toy play music.” Then the first set of three blocks was brought out and placed in a row on the table. The order of presentation of the three blocks was randomized. One at a time, the experimenter placed a block on the toy. Two of the three blocks in each set (the *target object* and the *causal match*) caused the toy to activate and

play music. The *perceptual match* did not. After children observed each outcome, they were asked for a verbal response. In the *explain* condition, children were asked to explain the outcome: “Why did/didn’t this block make my toy play music?” In the *control* condition, children were asked to describe the outcome (with a yes/no response): “What happened to my toy when I put this block on it? Did it play music?” After all three responses had been recorded, the experimenter demonstrated each of the three blocks on the toy a second time to facilitate recall.

Next the experimenter pointed to the set of objects and said, “Look! They have little doors. Let’s open one up.” The experimenter selected the *target object* and removed the cap to reveal the red map pin that had been hidden inside. The experimenter said, “Look! It has a little red thing inside of it. Can you point to the other one that also has something inside?” Children were then encouraged to point to one of the two remaining objects (i.e., the *perceptual match* or the *causal match*) to indicate which contained the same inside part, and this selection was recorded. Children could either select the block that was perceptually identical to the target or the object that shared the causal property, but not both.

Following their selection, children were not provided feedback, nor were they allowed to explore the blocks. Instead, all blocks were removed from view, and the next set was produced. This procedure was repeated for the three remaining sets. Each child participated in a total of four sets. Children were given a score of “1” for selecting the *causal match* and a “0” for selecting the *perceptual match*; children thus received 0-4 points across the four sets.

Results and Discussion

Data were analyzed with a 2 (condition) x 3 (age group) ANOVA, which revealed main effects of condition, $F(1, 102) = 50.70, p < .001$, and age, $F(2, 102) = 7.34, p < .01$ (see Fig. 2), with no significant interaction. Overall, children who were asked to explain ($M = 2.98, SD = 1.23$) were more likely than controls ($M = 1.61, SD = 1.58$) to generalize the internal part of the *target object* to the *causal match* rather than the *perceptual match*. Pairwise comparisons revealed no difference in performance between 3- and 4-year-olds, $p = .86$. However, 3- and 4-year-olds each selected the *causal match* significantly less often than 5-year-olds, $p < .01$.

One-sample t -tests were performed to assess whether explaining prompted children to override a preference for perceptual similarity and select the causal match. The 3-year-olds and 4-year-olds in the control condition selected the *perceptual match* significantly more often than chance, $t(17) = -3.69, p < .01$, and $t(17) = -2.53, p < .05$, respectively. Those in the explain condition selected the *causal match* significantly more often than chance, $t(17) = 3.01, p < .01$, and $t(17) = 2.48, p < .05$, respectively. Five-year-olds in the control condition performed no differently from chance ($M = 2.61, SD = 1.72$), $t(17) = 1.51, p = .15$, and 5-year-olds in the explain condition selected the *causal*

match significantly more often than expected by chance, $t(17) = 4.57, p < .001$.

These data suggest that in the absence of an explanation prompt, children relied primarily on information about the target object’s salient perceptual features to predict whether a novel object would share an internal property. However, when children of the same age were asked to generate an explanation for the evidence that they observed, they instead privileged the target object’s causal efficacy in making inferences about internal properties.

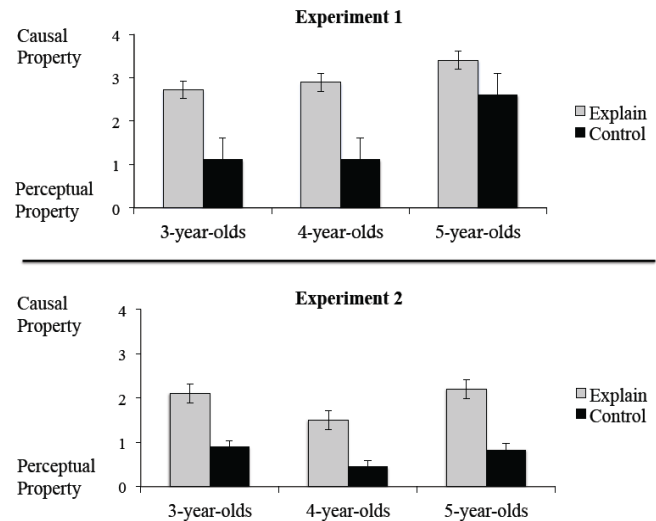


Figure 2: Average responses in explain and control conditions for Experiment 1 (top) and Experiment 2 (bottom). Higher numbers indicate a larger number of trials (of 4) on which an internal part (Experiment 1) or a label (Experiment 2) was generalized in line with a shared causal property over perceptual similarity.

Qualitative Data Explanations for the first object most often appealed to appearance (38%), with a minority (5%) appealing to internal properties. Explanations for the second set of objects, which occurred after observing the first set, appealed to appearance (33%) and internal properties (32%) equally often. By the final set, explanations most often appealed to internal parts (38%).¹

Experiment 2

Experiment 2 examined whether the influence of explanation on children’s inferences was restricted to consideration of internal parts, or whether these effects generalize to other inductively rich properties as well. A similar method was used to examine children’s generalization of a novel label from a target object to either a perceptually-matched or a causally-matched object.

¹ In the interest of space, we do not report the full qualitative analyses of children’s explanations in this paper. Instead, we provide a brief summary of these data Experiments 1-2.

Method

Participants A total of 108 children were included in Experiment 2, with 36 3-year-olds ($M = 42.1$ months; $SD = 3.8$, range: 35.9 – 48.0), 36 4-year-olds ($M = 54.0$ months; $SD = 3.0$, range: 48.4 – 59.9), and 36 5-year-olds ($M = 65.0$ months; $SD = 3.8$, range: 60.6 – 70.9). Eighteen children in each age group were randomly assigned to each of the two conditions (*explain* and *control*).

Materials The same toy was used in Experiment 2. Twelve wooden blocks of various shapes and colors were also used. There were a total of four sets of objects, each containing three blocks. As in Experiment 1, two of these blocks (the *target object* and the *perceptual match*) were perceptually identical (same color and shape) and one of these blocks (the *causal match*) was distinct (see Fig. 1).

Procedure The procedure for Experiment 2 was identical to Experiment 1, with one major exception: Instead of revealing a hidden internal property, the experimenter held up the *target object* and labeled it, saying, “See this one? This one is a blicket! Can you point to the other one that is also a blicket?” Again, children received a total of four sets of objects, and could receive 0-4 points.

Results and Discussion

Data were analyzed with a 2 (condition) $\times$ 3 (age group) ANOVA, which revealed a main effect of condition, $F(1, 102) = 13.51$, $p < .001$ (see Fig. 2), and no other significant effects. Overall, children who were asked to explain ($M = 1.91$, $SD = 1.83$) were more likely than controls ($M = .72$, $SD = 1.47$) to generalize the label to the *causal match*.

We next considered the data against chance responding. One-sample t -tests revealed that 3-, 4-, and 5-year-olds in the control condition selected the *perceptual match* significantly more often than chance, $t(17) = -2.93$, $p < .01$, $t(17) = -3.69$, $p < .01$, and $t(17) = -3.10$, $p < .01$, respectively. In the explanation condition, the average of children’s selections did not differ significantly from chance, $t(17) = .12$, $p = .90$, $t(17) = -1.26$, $p = .23$, and $t(17) = .375$, $p = .712$, respectively. However, the data for this condition were distributed bimodally, with approximately half the children providing a majority of causal choices and half perceptual choices. The percentage of children selecting the *causal match* on 3 or 4 trials was 50% for 3-year-olds, 44% for 4-year-olds, and 56% for 5-year-olds. The *distribution* of children’s selections did differ significantly from that expected by chance in all age groups, $\chi^2(4) = 84.26$, $p < .001$, $\chi^2(4) = 66.49$, $p < .001$, and $\chi^2(4) = 83.97$, $p < .001$, respectively.

Like the younger children in Experiment 1, children in the control condition relied primarily on information about a target object’s salient perceptual features to predict whether a novel object would share a category label. However, when children of all ages were asked to generate an explanation for the evidence that they observed, they

considered the target object’s causal efficacy significantly more often in making inferences about shared labels.

Qualitative Data Appearance explanations were most common overall (28%); however, there was an increase in the proportion of explanations appealing to kind or explicitly mentioning a label, with 7% in the first set and 19% in the final set.

Experiment 3

The findings from Experiments 1 and 2 confirm our prediction that explanation encourages children to favor inductively rich properties (i.e., causality) as a basis for generalization. However, the findings are also consistent with an alternative explanation: that prompts to explain increased children’s overall attention or engagement, resulting in “better” performance. Experiment 3 tests this alternative.

In Experiment 3, children were asked to explain or report causal outcomes after observing four unique objects, two of which activated the toy. Because we did not observe relevant age differences in Experiments 1-2, Experiment 3 was restricted to 4-year olds. After each object was placed on the toy, three properties were revealed: an internal part, a label, and a sticker (added to the object). The internal parts and the labels correlated with the toy’s activation (i.e., all and only objects that activated the toy had a particular inside part and label) while the sticker did not. Children then completed a memory task.

The purpose of Experiment 3 was to assess whether the effects of explanation observed in Experiments 1-2 could be due to a global and indiscriminate boost in attention. Based on our interpretation of Experiments 1 and 2, we predicted that a prompt to explain would improve memory for inside parts and labels, but not for the sticker, which was neither correlated with other properties nor plausibly inductively potent. If the effects of explanation can instead be attributed to a global increase in attention or engagement, one would predict improved memory for all features.

Method

Participants A total of 36 4-year-olds were included in Study 3 ($M = 53.8$ months; $SD = 3.7$ months; range = 47.9 – 59.7). Eighteen children were randomly assigned to one of two conditions (*explain* and *control*).

Materials Experiment 3 used the same toy as in the previous experiments. Four unique wooden blocks (distinct colors and shapes) were also used (see Table 1). As in Experiment 1, all blocks had a hole drilled into the center. Two of the blocks had a red, round plastic map pin glued inside and two of the blocks had a white, square eraser glued inside the hole. Four stickers were used during the experiment (two heart stickers and two star stickers). Several small cards were constructed as memory aids during the test phase of the experiment. Half of the cards had an image of a black music note (placed in front of the objects that children believed activated the toy), and half of the

cards had an image of a black music note crossed out with a red “X” (placed in front of the objects that children believed did not activate the toy). Four additional cards were constructed: one with a red circle, one with a white square, one with a heart sticker, and one with a star sticker. These cards were used to facilitate the forced choice test.

Table 1. List of properties for objects used in Experiment 3.

	Object 1	Object 2	Object 3	Object 4
Causal	Yes	No	Yes	No
Internal	Red	White	Red	White
Label	Toma	Fep	Toma	Fep
Sticker	Heart	Heart	Star	Star

Procedure As in Experiments 1 and 2, the experimenter introduced the toy. The experimenter then produced a single block and placed it on the toy. The child observed as the block did or did not cause the toy to play music. As before, children in the *explain* condition were asked to explain the outcome for each of the blocks and children in the *control* condition were asked to report the outcome (with a “yes/no” response). After the response was recorded, the experimenter repeated the demonstration a second time.

The experimenter provided three additional pieces of information about the object: the type of internal part was revealed (“Look! It has a little door on it! Let’s open it up. Look, there is a [red]/[white] thing inside.”), a label was provided (“See this one? This one here? This one is a [Fep]/[Toma]!”), and a sticker was placed on the bottom (“Now I am going to put a sticker on it! I am going to put a [heart]/[star] sticker on the bottom, see?”). The experimenter repeated each property twice, and then the block was removed from view. The entire procedure was repeated for the three remaining blocks, one at a time. All children observed the causal property first. The order of the remaining three properties was counterbalanced.

Next, the experimenter placed all four objects on the table in front of the child in random order, and told the child that they would now play a “memory game.” To assess recall for the causal property of each object, the experimenter produced two cards – one with a music note, and one with a crossed out music note. The experimenter asked the child to point to the card that indicated whether the block did or did not play music. The child responded once for each of the four objects. Depending upon the child’s response, the experimenter would then place an additional card (with a music note or a crossed-out music note) in front of the object, which would remain throughout.

To assess recall for the internal part, the experimenter produced two cards – one with a red circle and one with a white square. The experimenter asked the child to point to the card that indicated the type of thing inside the block. The child responded once for each of the four objects. To assess recall for the label, the experiment said, “Some of these blocks were called ‘Tomas’ and some of these blocks

were called ‘Feps’. What was this one called, a ‘Toma’ or a ‘Fep’?” The child responded once for each object. The order of presentation was counterbalanced across trials.

Finally, to assess recall for the type of sticker added to the block, the experimenter produced two cards – one with a heart sticker and one with a star sticker. The experimenter asked the child to point to the card that indicated the type of sticker added to the bottom of the block. The child responded once for each of the four objects.

Memory for internal parts, labels, and stickers was solicited in the same order as the corresponding properties were presented to that child in the demonstration phase of the experiment. For each property, children were given a score of “1” for accurate recall and a “0” for inaccurate recall. Because there were a total of four objects, each child could receive between 0 and 4 points for each property.

Results and Discussion

Memory for the objects’ causal properties was analyzed with a one-way ANOVA, which revealed that children in the *explain* condition were significantly more accurate ($M = 3.93$) than controls ($M = 3.39$), $F(1, 34) = 8.42, p < .01$.

Next, a repeated measures ANOVA with the other object properties (internal part, label, sticker) as the repeated measure and condition (*explain*, *control*) as the between subjects variable revealed a main effect of object properties, $F(2, 68) = 6.96, p < .01$, and the predicted interaction between object properties and condition, $F(2, 68) = 8.30, p < .01$ (see Fig. 3). Children who were prompted to explain were significantly more accurate than controls in reporting the labels, $F(1, 34) = 9.34, p < .01$, but *less* accurate than controls in recalling the sticker type, $F(1, 34) = 5.16, p < .05$.

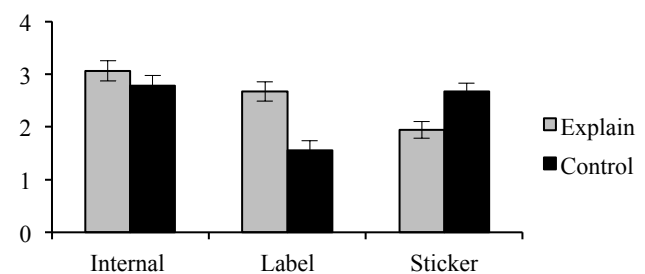


Figure 3: Average memory score (out of 4 trials) for each property assessed in Experiment 3. Error bars correspond to one SEM in each direction.

These data provide evidence against the possibility that engaging in explanation simply improves overall attention to the task. Instead, children who explained were more likely to recall the properties that were inductively rich, while ignoring a superficial, perceptual property that did not correlate with other features.

General Discussion

In each of these experiments, prompting young children to explain made them more likely to privilege inductively rich, non-obvious properties over salient surface similarity in

making novel inferences. Children in the control condition, who were not prompted to explain, based their judgments on perceptual similarity. These effects were not restricted to a particular property or domain, as comparable effects were observed across two quite different properties: internal parts (Experiment 1) and novel labels (Experiment 2).

Although explanation led to fewer perceptual responses in Experiment 1 than in 2, this difference parallels the disparity in children's baseline performance in the control condition (see Fig. 2). In other words, children were more willing to privilege internal parts over appearances than labels over appearances, in line with previous research (Gopnik & Sobel, 2000; Sobel et al., 2007). For our purposes, the critical finding is that explanation decreased perceptual responding in both cases.

Finally, the results of Experiment 3 provide additional support by demonstrating improved memory for a correlated cluster of inductively rich properties in children prompted to explain. Importantly, Experiment 3 also provides evidence that effects of explanation are selective: Children who explained had impaired memory for the superficial property. This provides evidence against the idea that explanation produces a general benefit for learning by globally and indiscriminately increasing engagement or motivation.

Why might explaining lead children to favor inductively rich properties? Wellman and Liu (2007) suggest that explaining makes an occurrence sensible by reference to a larger framework: The explainer appeals to the interplay between evidence and current theories to construe the phenomenon as an instance of a larger, coherent system. In line with this idea, recent computational approaches to cognitive development have proposed that the formation of generalizations at multiple levels of abstraction enables learners to learn quickly and generalize effectively to novel cases (Kemp, Perfors & Tenenbaum, 2007). By prompting children to favor inductively rich regularities, explanation could play a role in pushing children beyond immediate observations to consider higher-order generalizations that support abstract knowledge.

Similarly, we have argued that engaging in explanation constrains the learner to consider an event as an instance of a broad generalization (see also Lombrozo, 2012). Recall that previous research found that explaining magnified effects of prior knowledge in the service of broad generalization (Walker et al., 2012). In the current study, the belief that internal parts and labels are inductively rich is itself an important instance of higher-order prior knowledge. We propose that explaining contributes to the formation of causal theories by constraining learners to consider properties that are most likely to generalize to novel cases. In the experiments presented here, this constraint improved children's ability to override highly salient perceptual cues.

Acknowledgments

Thanks to Rosie Aboody, Dustin Bainto, Ela Banerjee, Brynna Ledford, Donna Ni, Gabriel Poon, UC Berkeley Early Childhood Centers, and Lawrence Hall of Science.

Research supported by the McDonnell Foundation and NSF grant DRL-1056712 to T. Lombrozo.

References

- Chi, M. T. H., DeLeeuw, N., Chiu, M.H., & LaVancher, C. (1994). Eliciting self-explanations improves understanding. *Cog Sci*, 18, 439-477.
- Gelman, S. & Markman, E. (1987). Young children's induction from natural kinds: The role of categories and appearances. *Child Dev*, 58: 1532-1541.
- Gelman, S. (2003). *The essential child: Origins of essentialism in everyday thought*. Oxford: Oxford Press.
- Gopnik, A. & Sobel, D. (2000). Detecting blickets: How young children use information about novel causal powers in categorization and induction, *Child Dev*, 71(5): 1205-1222.
- Keil, F.C. (1989). *Concepts, kinds, and cognitive development*. Cambridge, MA: MIT Press.
- Kemp, C., Perfors, A., & Tenenbaum, J.B. (2007). Learning overhypotheses with hierarchical Bayesian models. *Dev Sci*, 10, 307-321.
- Legare, C.H. (2012). Exploring explanation: Explaining inconsistent information guides hypothesis-testing behavior in young children. *Child Dev*, 83, 173-185.
- Legare, C.H., Gelman, S.A., & Wellman, H.M. (2010). Inconsistency with prior knowledge triggers children's causal explanatory reasoning. *Child Dev*, 81, 929-944.
- Legare & Lombrozo (under review). Unique and selective benefits of explanation and exploration on learning in early childhood.
- Legare, C.H., Wellman, H.M., & Gelman, S.A. (2009). Evidence for an explanation advantage in naïve biological reasoning. *Cog Psych*, 58, 177-194.
- Lombrozo, T. (2010). Explanation and abductive inference. In K.J. Holyoak and R.G. Morrison (Eds.), *Oxford Handbook of Thinking and Reasoning* (pp. 260-276), Oxford, UK: Oxford Univ. Press.
- Nazzi, T. & Gopnik, A. (2000). A shift in children's use of perceptual and causal cues to categorization. *Dev Sci*, 3(4): 389-396.
- Sobel, D.M. Yoachim, C.M. Gopnik, A. Meltzoff, A.N. & Blumenthal, E.J. (2007). The blicket within: Preschoolers' inferences about insides and causes. *J. of Cog Dev*, 8(2): 159-182.
- Walker, C.M., Williams, J.J., Lombrozo, T. & Gopnik, A. (2012). Explaining influences children's reliance on evidence and prior knowledge in causal induction. In N. Miyake, D. Peebles, & R.P. Cooper (Eds.), *Proceedings of the 34th Ann Conf of the Cog Sci Soc*, pp. 1114-1119. Austin, TX: Cog Sci Soc.
- Wellman, H.M. & Liu, D. (2007). Causal reasoning as informed by the early development of explanations. In *Causal Learning: Psychology, Philosophy, & Computation* (eds.) L. Schulz & A. Gopnik, pp. 261-279.
- Williams, J.J. & Lombrozo, T. (2010). The role of explanation in discovery and generalization: Evidence from category learning. *Cog Sci*, 34: 776-806.