

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Intuitive archeology: Detecting social transmissino in the design of artifacts

Permalink

<https://escholarship.org/uc/item/7p40g7d2>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 40(0)

Authors

Schachner, Adena

Brady, Timothy F

Oro, Kiani

et al.

Publication Date

2018

Intuitive archeology: Detecting social transmission in the design of artifacts

Adena Schachner (schachner@ucsd.edu), Timothy F. Brady (timbrady@ucsd.edu),
Kiani Oro (koro001@fiu.edu), Michelle Lee (m.lee@ucsd.edu)

University of California, San Diego, Department of Psychology
9500 Gilman Drive M/C 0109, San Diego, CA 92093-0109 USA

Abstract

Human-made objects (artifacts) often provide rich social information about the people who created them. We explore how people reason about others from the objects they create, characterizing inferences about when social transmission of ideas (copying) has occurred. We test whether judgments are driven by perceptual heuristics, or structured explanation-based reasoning. We develop a Bayesian model of explanation-based inference from artifacts and a simpler model of perceptual heuristics, and ask which better predicts people's judgments. Our artifact-building task involved two characters who built toy train tracks. Participants viewed pairs of tracks, and judged whether copying had occurred. Our explanation-based model accurately predicted on a trial-by-trial basis when participants inferred copying; the perceptual heuristics model was significantly less accurate. Efficient design 'explained away' similarity, making similarity weaker evidence of copying for efficient tracks. Overall, data show that like intuitive archeologists, people make rich explanation-based inferences about others from the objects they create.

Keywords: social cognition; Bayesian inference; explanation; social transmission; imitation; artifact; design

Introduction

We live in environments saturated by human-made objects. These *artifacts* differ fundamentally from natural objects like rocks and trees: Their properties are intentionally designed, often for a specific purpose (Kelemen & Carey, 2007). Reasoning about artifacts has enormous practical consequences, as effective use of artifacts as tools is necessary to function in any human society—and has been crucial to human survival and evolution for hundreds of thousands of years (Gibson & Ingold, 1995).

Artifacts are crucial to our lives not only as tools, but also as an ever-present source of social information. Based solely on the artifacts a person owns, people form quick and accurate judgments about a person's traits, interests, and social affiliations (Gosling, 2008; Richins, 1994; Ferraro et al., 2010; Solomon, 1983; McCracken, 1986; Belk, 1988). The artifacts a person *creates* may hold the most social information: Creations like visual art, music, or written text provide rich information about traits, beliefs, and identity (Gosling, 2008).

How do people reason about other individuals from the artifacts they create? Here we explore the nature of this reasoning, which we term *intuitive archaeology*. Just like archaeologists use inanimate objects to make explicit inferences about the people and cultures that created them, people may intuitively infer complex social information from the design of objects. Doing so would require

integrating our mental theories of the physical-mechanical world with our theories of the social world (e.g. Battaglia, Hamrick & Tenenbaum, 2013; Gopnik, 2012; Carey, 2011; Baker, Saxe & Tenenbaum, 2009), to make inferences about the most probable cause of artifact features.

Reasoning about people from their creations fundamentally involves inferring when *social transmission of ideas* has occurred (i.e. social learning, imitation, copying), and when a design was *generated independently* by an individual. Designs that were socially transmitted will be informative about their designers' social history, while independently designed features will not, and may instead lead to different inferences (e.g. about the intelligence of the designer, Gosling, 2008). For instance, if you and another person draw the same detailed drawing, this is unlikely to occur by chance: there are so many degrees of freedom that it is unlikely that two identical drawings were independently generated. Instead, the information is likely to have been socially transmitted, such that one of you copied the other, or you share some social history, such as a common teacher or cultural group (Soley & Spelke, 2016). In this way, reasoning about social transmission serves as the foundation for inferences about important social properties, such as a person's social and cultural group membership.

In the current paper, we explore peoples' inferences about social transmission from the designs of other's artifacts, and propose and test two accounts of the cognitive basis of this inference. In particular, we ask whether participants' judgments are driven by *perceptual heuristics* or instead by *explanation-based reasoning*.

Explanation-based reasoning

We hypothesize that participants will draw conclusions about social transmission through explanation-based reasoning, as an inference to the best explanation (Lipton, 2004; Tenenbaum et al., 2006). Under this account, to explain observed data (e.g. two people have created identical artifacts), people consider multiple hypotheses, then choose the hypothesis that is both plausible a priori and provides a strong explanation of the data (e.g. they copied one another). This type of inferential reasoning occurs in multiple related domains, including causal induction and reasoning about others' mental states (Teglas et al., 2011; Baker et al., 2009).

This account makes specific predictions. People should infer that social transmission has occurred when people create similar artifacts; but should not treat all similar perceptual features as equally strong evidence of copying. If an alternative explanation for the similarity is given, such as a functional constraint (e.g. a barrier constraining which

designs could work) or an independently-shared bias (e.g. the desire to create an *efficient* design), participants should no longer infer that social transmission has occurred, even when two people create the same artifact. These alternative explanations should ‘explain away’ the similarity, making the similarity weaker evidence of copying.

Perceptual heuristics

Alternatively, people may use a simpler cognitive strategy to detect social transmission: heuristics based on perceptual similarity. Observers may simply note the extent to which two artifacts look similar, and use this to judge whether the ideas were copied.

If observers use this perceptual heuristic, we should see specific, revealing errors when conceptual information conflicts with perceptual similarity. In particular, providing an alternative explanation for the common features should have no effect on inferences about copying. Thus the presence of an independently-shared bias (the desire to create an efficient design) or functional constraint (a barrier constraining the available design options) should have no effect. If people are using a perceptual heuristic based on the artifacts’ similarity, they should continue to infer that copying occurred when similar artifacts are created, even under conditions that provide alternative explanations.

The current study

To formally characterize participants’ reasoning, we first develop a Bayesian model of explanation-based inference from artifacts, based on a generative model of artifact features. Broadly speaking, our artifact-generation model has a similar structure to the classic use of Bayes nets in the literature on how children learn about causal structure (e.g. the blinket detector paradigm, Gopnik & Sobel 2000).

The broad idea of the model is to consider the two possible generative processes of an artifact’s design (social transmission of ideas vs. independent creation of ideas) and ask which provides a better explanation of the features of a given artifact. Given this structure, independent creation can “explain away” some aspects of the similarity between artifacts in certain circumstances, reducing the likelihood of inferring social transmission. We particularly test the role of *efficiency* as an alternative explanation: People generally have a strong desire to create efficient designs (Dennett, 1990). Thus, if there is a clear efficient solution, two people may independently create the same artifact design.

We contrast this model with a model of the simpler, perceptual heuristic account, and then tease apart our two models with experimental evidence, asking which model better predicts how people reason about social transmission from the design of artifacts.

A Generative Model of Artifact Design

We model the case of a train track building task with the following structure: Imagine you see a toy train track built by person 1, track t_1 , and a second track, t_2 , built by person

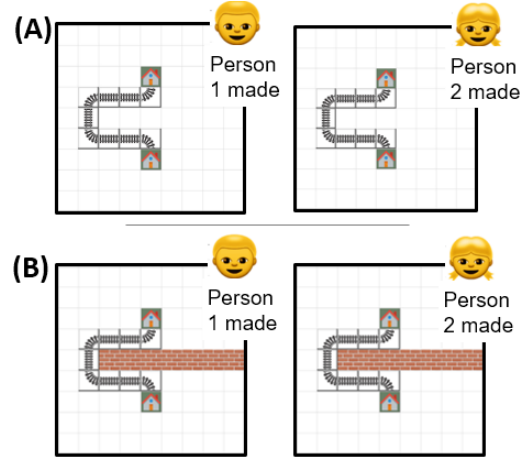


Figure 1: Example stimuli. Did the two builders copy one another, or design their tracks independently? An explanation-based model predicts that identical designs should be seen as stronger evidence of copying in A than B.

2. You wish to infer whether person 2 copied the track’s design from person 1, or independently created it. Each person had 7 straight and 7 curved pieces to build with.

Framed mathematically, this can be understood as an inference, where if c indicates whether person 2 copied person 1, you wish to infer the probability of copying $p(c)$, given the observed object t_1 of person 1 and object t_2 of person 2, along with two additional parameters: a , copying accuracy, and f , the strength of the builders’ efficiency preference. Thus, we wish to make the following inference:

$$p(c | t_1, t_2, a, f) = \frac{p(c)p(t_2 | c, t_1, a)}{p(t_2)}$$

$$p(t_2) = p(c) p(t_2 | c, t_1, a) + (1 - p(c)) p(t_2 | \neg c, f)$$

Here $p(c | t_1, t_2, a, f)$ is the posterior probability on copying after we observe the tracks made by the two people, t_1 and t_2 . This has the structure of a Bayes net, including the key concept of explaining away: A track can be generated either via copying or independently, and evidence for one provides implicit evidence against the other. Thus, if two people create identical track designs, but this design is likely to be created independently, this provides weaker evidence of copying.

To make the model concrete, we must specify three things:

1. $p(c)$ – the *a priori* estimate of how likely person 2 was to have copied (unconditional on the data; i.e., before we see either of the tracks t_1/t_2); this will be set by the description of the context. For example, if person 1 is sitting next to person 2 and encouraged to collaborate, this might be close to 1.0. If person 1 and 2 are in different rooms with no way to see each other, it might be close to 0.0.
2. $p(t_2 | c, t_1, a)$ – the likelihood of the particular track t_2 , given that person 2 was in fact copying person 1’s track (t_1); and given a copying accuracy parameter (a). This graded copying accuracy parameter captures the

intuition that there may be partial social transmission of ideas – participants may be influenced by each other, while not making exactly the same track.

3. $p(t_2|\neg c, f)$ - the likelihood of track t_2 , given that person 2 was NOT copying from person 1's track (t_1), but independently generated the track with no reliance on t_1 . This depends on f , the strength of the efficiency preference – e.g., how likely people are to prefer short, efficient tracks that directly connect the endpoints. A high efficiency parameter (f) should place high likelihood on extremely short tracks, thus making these highly efficient track designs very likely even if the tracks were generated independently – weakening the evidence of copying from certain perceptually identical tracks. This parameter captures the intuition that in this context, participants may be more or less likely to prioritize getting from A to B efficiently, which should change the likelihoods of each track to occur independently and thus affect judgments of whether social transmission has occurred.

To formalize these likelihoods, we treat a and f as the slope of logistic function on similarity between tracks (for a) and on track length (for f). Thus,

$$p(t_2|c, t_1, a) \propto \frac{1}{1 + e^{(a * d_{t_1, t_2})}}$$

where d_{t_1, t_2} is the distance between track t_1 and t_2 , here modeled as the number of non-overlapping piece locations from the two tracks. Probabilities are normalized over the set of t_2 's. Similarly:

$$p(t_2|\neg c, f) \propto \frac{1}{1 + e^{(f * L_{t_2})}}$$

where L_{t_2} is the length of track t_2 (since all track options are required to get from A to B, shorter tracks are by definition more efficient). Probabilities are again normalized over the set of t_2 's.

We also constructed a simpler formal model of the perceptual heuristic account, which makes judgments about copying based solely on the extent of perceptual similarity of the two tracks (using the same metric of difference as the first model: the number of non-overlapping piece locations). This simpler model treats $p(t_2|\neg c)$ as a uniform distribution over all possible tracks given the available pieces.

Given the way these functions are parameterized, larger positive values of a lead to a stronger copying accuracy (e.g., a steeper fall off in how likely a track is to be generated if it is more dissimilar to t_1), and larger positive values of f in the explanation-based model provide a stronger efficiency preference (e.g., a steeper fall off in how likely people are to generate long tracks). Because these parameters are directly interpretable by visualizing how often a given track would be generated under each parameter value, and are largely shared between the two models and with similar effects on each (with the exception of f), we primarily analyzed the model fits with a priori chosen reasonable values, with $p(c)$ (the prior on copying) to 0.10, f (the efficiency parameter) to 1.5, and a (the

copying fidelity parameter) to 2.5. These parameter values make it such that if the builders do copy, they would usually generate almost exactly the same track (thus having a very low rate of error). Similarly, the strength of the efficiency preference (f) provides a moderate-to-strong preference for short tracks, since when people are designing artifacts (i.e. building independently), they generally have a relatively strong preference for efficient designs (Dennett, 1990). However, we also performed model comparison across all reasonable parameter settings (0 to 5 in steps of 0.5 for each of f and a , subject to the constraint that $a > f$ so that copying is more likely to result in the same track than is independent generation).

Testing the models' predictions

This formal model of explanation based reasoning and of perceptual heuristics makes quantitative predictions about the likelihood of copying for any given pair of track designs, in a wide range of contexts. We aimed to test the following three predictions of our models for human behavior, which tease apart our two theoretical accounts.

Firstly, the explanation-based reasoning model predicts that when there is an alternative explanation for the artifacts' similarity, it will 'explain away' the similarity: That is, similarity will be weaker evidence of copying in this case. We test the role of *efficiency* as an alternative explanation, as people generally have a strong, independent desire to create efficient designs (Dennett, 1990); if there is a clear efficient solution, two people may independently create the same artifact design. Thus our explanation-based reasoning model (but not the perceptual heuristic model) predicts that for two identical tracks, people will infer that copying is more likely to have occurred when the two tracks are inefficient paths between the start and end point than when they are efficient paths, even when the tracks are perceptually identical in both cases. To test this, we manipulated the efficiency of track designs, separately from their perceptual overlap.

Our explanation-based reasoning model also predicts that the inferred strength of the designers' efficiency preference (parameter f) should play a role in people's reasoning. To test this, we manipulated whether the builders were instructed to build tracks that would 'get there quickly', or not, to see if this affected participants' judgments about copying in a way predicted by the model.

Lastly, our explanation-based reasoning model predicts that people should take into account the prior likelihood of copying when making their judgments. By manipulating how far from one another the builders were when making their tracks, we manipulated the prior likelihood of copying, to see if judgments changed in the way predicted by the model.

Method

Participants

48 adults were recruited from the undergraduate population at the University of California, San Diego (19 male, 1 gender-fluid, 28 female; mean age = 21.45 years, $SD = 1.75$, Range = 18.9 - 25.7). Participants received course credit for participating.

Procedure and Stimuli

Participants were tested in the lab, with stimuli presented on an iMac computer in a web browser, using lab.tellab.org.

On each trial, participants saw two 9x9 grids, with two locations marked with house icons (here termed A and B). Each grid showed a train track that went from A to B (see Figure 1). Participants were instructed that two different people had built the two different tracks, and were asked to judge whether the builders had copied each other's designs, or created their designs independently.

All subjects completed 6 blocks of 9 trials each, with brief instructions before each block. Trials within each block were presented in one of four pseudo-random orders. Blocks were presented in pairs, with the order of the blocks within each pair counterbalanced between-subjects. The position of the train track stimuli (left vs. right side of the screen) was counterbalanced between subjects.

Manipulating the tracks' efficiency. We manipulated the efficiency of the train track paths in two different ways. First, we manipulated the track length, with shorter paths being more efficient. We contrasted three length conditions: Short (3 track units, highly efficient), medium (7 units) and long (9 units, low efficiency).

Second, we manipulated the tracks' efficiency by introducing a barrier to the grid. The barrier made formerly inefficient tracks more efficient, allowing us to compare tracks with exactly the same perceptual features, under conditions where they were efficient vs. inefficient. We again tested three different track lengths, to parallel the no-barrier condition: Short (9 units), medium (11 units) and long (13 units).

One crucial test case is shown in Figure 1: The long track from the no-barrier condition was perceptually identical to the short track in the barrier condition. With the barrier present, this track became the most efficient track possible, whereas it had previously been highly inefficient.

Manipulating the tracks' perceptual similarity. We also manipulated the perceptual similarity of the two tracks to one another. Each of the main tracks tested (see above) was paired with one of three 'comparison tracks': an identical track, a track with a minor shape difference, or a track with a greater (major) shape difference. This created a 3 (track perceptual similarity) x 3 (track length) x 2 (barrier present/absent) design, with these manipulations tested during the first two blocks of trials (main trials).

Manipulating the expectation of efficiency In two additional blocks of trials, we asked whether participants' judgments of copying changed based on the strength of their expectation of efficiency (parameter f). These trials repeated the same track shapes as the no-barrier trials described above, but manipulated the instructions. On *strong efficiency expectation* trials, the instructions specified that the builders were trying to "get from A to B as quickly as possible", while on *weak efficiency expectation* trials, instructions said they were "just having fun with the train track building task, and just had to make sure their tracks went from A to B".

Manipulating the prior on copying Lastly, we tested whether the prior likelihood of copying affected participants' judgments, in two final blocks of trials. These trials repeated the same track pairs as the no-barrier trials described above, but manipulated the prior via the instructions. On the *high copying prior* trials, instructions said the builders of the paths were "sitting next to each other"; on *low copying prior* trials, instructions said "the builders were seated far apart, facing opposite directions".

Results

Overall, participants inferred that the designs were copied most often when the tracks were identical, and least often when they had major differences (Mean % answering "copied", Identical: 66.2%, $SEM = \pm 2.4\%$; Minor difference: 37.8%, $SEM = \pm 3.0\%$; Major difference: 4.5%, $SEM = \pm 1.2\%$; all contrasts $p < 0.0001$). Thus, participants were attending to the task and taking into account the perceptual similarity of the tracks when inferring whether they were copied.

Our major question of interest was whether participants take into account other factors, like efficiency, that "explain away" some aspects of similarity and reduce the likelihood of copying for some pairs of tracks. To assess this, we first looked at the holistic fit between each of the two formal models (explanation-based model; perceptual heuristics model) and the behavioral data, focusing on the data from the main trials with the normal instructions (see above, 'Manipulating the tracks' efficiency'). Overall model predictions for the explanation-based and perceptual similarity model are plotted in Figure 2. The explanation-based model provided a better match to participants' judgments ($R^2 = 0.93$) than did the perceptual similarity model ($R^2 = 0.69$; difference: $p = 0.0004$; Steiger test for dependent correlations). This difference held across all reasonable parameter values, with an average difference in R^2 of 0.24 in favor of the explanation-based model across all possible parameter settings. This high correlation means that on trials where participants more often say that the builders copied, the explanation-based model also tends to say copying is more likely – the model can predict on a trial-by-trial basis when copying is more likely to be inferred. In contrast, the perceptual heuristics model is significantly less accurate at predicting participants' responses.

Several patterns in participants' responses provide additional empirical support for the explanation-based model. Broadly, the explanation-based model predicts that in addition to using perceptual similarity between the tracks to infer copying, participants should explain away some kinds of perceptual similarity with alternative explanations for independent generation. In particular, because people tend to aim for efficient designs, the likelihood of independently generating efficient tracks is high; the model predicts that similarity between efficient tracks should not be strong evidence of copying. In addition, the model predicts that the extent of this 'explaining away' by efficiency should be modulated by the strength of the builders' preference for efficient tracks: If the builders are all trying hard to make very efficient tracks, then they are likely to generate the same design independently, and similar efficient designs should not provide evidence of copying. In contrast, if people are not trying to be efficient, similar efficient designs should serve as evidence of copying. Finally, the model predicts a strong role for the prior: if *a priori* we believe participants are extremely unlikely to be able to copy each other (e.g., if they are seated very far apart), then we should not infer they copied, even if they produce perceptually identical tracks.

We found evidence consistent with all of these predictions. In particular, efficiency played a significant role in people's judgments above and beyond perceptual similarity. For example, even when pairs of tracks were perceptually identical, people judged the most efficient tracks ($M=32.2\%$, $SEM:\pm 3.7\%$) as copied less often than the least efficient tracks ($M=83.7\%$, $SEM:\pm 2.9\%$; difference: $p<0.0001$). In addition, the exact same track was judged less likely to be copied when a barrier was present, making that track a highly efficient design, given the constraints ($M=56.3\%$, $SEM:\pm 7.2\%$), versus when the barrier was not present, making that track an inefficient design ($M=83.3\%$, $SEM:\pm 3.1\%$; difference: $p=0.0004$; for tracks in Figure 1).

In addition, people were more likely to say that two identical, highly efficient tracks were copied when they thought the builders of those tracks were not aiming for efficiency ($M=45.8\%$; $SEM:\pm 7.2\%$) than when they thought the builders were aiming for efficiency ($M=8.3\%$; $SEM:\pm 4.0\%$; difference: $p<0.0001$). As predicted by the explanation-based model, this 'strength of efficiency preference' manipulation selectively affected judgments about efficient tracks, and not inefficient tracks, for which similarity cannot be explained away by efficiency preferences (medium and long length; strong efficiency preference: $M=83.3\%$, $SEM:\pm 4.3\%$; weak efficiency preference: $M=83.3\%$, $SEM:\pm 4.3\%$; $p>0.10$).

Participants' prior on copying also affected their judgments, as predicted by the explanation-based model: They judged tracks as copied more often when the builders were sitting closer together ($M=80.6\%$; $SEM:\pm 3.3\%$) than when they were far apart ($M=49.3\%$; $SEM:\pm 4.5\%$). When builders were far apart and generated efficient tracks, participants inferred copying extremely rarely, even when

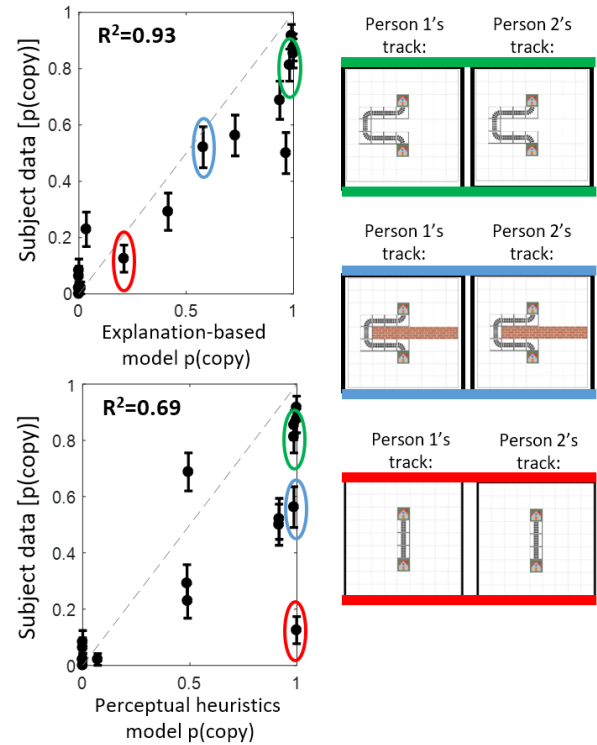


Figure 2: Left: Fit of explanation-based model (top) and perceptual heuristics model (bottom) to participants' likelihood of saying each trial was copied. Each point is a trial; error bars indicate SEM across participants. Right: Three example trials are highlighted in both graphs (by color). Because these 3 trials have perceptually identical tracks, the perceptual heuristics model considers them likely copied. By contrast, the explanation-based model differs markedly in the predicted likelihood of copying across these trials, as do participants' actual judgments.

those tracks were perceptually identical (Identical: $M=16.6\%$; $SEM:\pm 5.4\%$). By contrast, the perceptual heuristics model predicts participants should infer copying nearly 100% of the time in this situation, as the tracks the two builders created were identical.

These effects of efficiency are incompatible with a perceptual similarity model, which views copying as directly related to similarity in the track's structure (for example, see the trials highlighted in Figure 2). However, they are predicted by the explanation-based model, which 'explains away' similarity by taking into account people's general tendency to generate efficient designs. While only model predictions with our *a priori* parameter values are shown in Figure 2, this tendency to discount the likelihood of copying for efficient tracks is true across all reasonable parameter settings, as noted above (e.g. parameter settings where the generative process of "copying" is not so error-prone that it generates less-similar tracks than the generative process of "independently creating").

Discussion

How do people reason about other individuals from the artifacts they create? Reasoning about social transmission of ideas serves as a foundation for inferences about shared social history and group membership. We characterized reasoning about social transmission from artifact features, focusing on copying, an essential and direct form of social transmission. We asked if people use perceptual heuristics or explanation-based reasoning to infer when copying has occurred. We constructed formal Bayesian models to make quantitative predictions, and asked which model better predicted peoples' judgments.

We found clear evidence that participants used explanation-based reasoning, not perceptual heuristics, to draw conclusions about when social transmission of information had occurred. As predicted by the explanation-based reasoning model, it was possible to 'explain away' similarity: Similarity was weaker evidence of copying in the presence of an alternative explanation. This is a key prediction of rational, inferential, explanation-based reasoning, which differentiates it from simpler strategies such as perceptual heuristics (e.g. Tenenbaum et al., 2006). The simpler model of perceptual similarity based heuristics failed to capture human judgments in cases where similarity can be 'explained away', while the explanation-based model successfully predicted judgments even in these cases.

For example, the explanation-based reasoning model predicted that when two people created the same efficient design, this would be seen as weaker evidence of copying than when two people created the same inefficient design. This was the case even when comparing the same exact pair of tracks in different contexts (as in Figure 1). In this case, *efficiency* served as a plausible alternative explanation, showing that people implicitly take into account the idea that other people generally try to create efficient designs (Dennett, 1990). Participants also took into account the strength of the builders' efficiency preference, and their a priori likelihood of copying due to differences in distance from their partner. Each of these patterns are predicted by the model of explanation-based reasoning, and are not captured by a simpler model of perceptual heuristics.

These data provide evidence that people are able to intuitively integrate their mental theories of the physical-mechanical world with their theories of the social world, to draw complex social conclusions from physical object features. In doing so, they use reasoning that is rich, complex, structured, and inferential.

In both content and process, then, this type of reasoning appears analogous to the explicit scientific reasoning of archeologists, who make systematic inferences about people and cultures from the features of inanimate objects. Like intuitive archeologists, people are able to intuitively infer social transmission from design of objects. This ability to make social inferences from objects may provide an important source of social information in everyday life, with the power to inform and support social interaction.

References

- Baker, C.L., Saxe, R., & Tenenbaum, J.B. (2009). Action understanding as inverse planning. *Cognition*, 113, 329-349.
- Battaglia, P. W., Hamrick, J. B., & Tenenbaum, J. B. (2013). Simulation as an engine of physical scene understanding. *Proceedings of the National Academy of Sciences*, 110(45), 18327-18332.
- Belk, R. W. (1988). Possessions and the extended self. *Journal of Consumer Research*, 139-168.
- Carey, S. (2011). *The Origin of Concepts*. Oxford: Oxford University Press.
- Dennett, D. (1990). The interpretation of texts, people, and other artifacts. *Philosophy and Phenomenological Research*, 50, 177-194.
- Ferraro, R., Escalas, J.E., & Bettman, J.R. (2010). Our possessions, our selves: Domains of self-worth and the possession-self link. *J. Consumer Psychology*, 1-9.
- Gibson, T. & Ingold, K.R. (1995). *Tools, Language and Cognition in Human Evolution*. Cambridge: Cambridge University Press.
- Gopnik, A., & Sobel, D. M. (2000). Detecting blickets: How young children use information about novel causal powers in categorization and induction. *Child Development*, 71(5), 1205-1222.
- Gopnik, A. (2012). Scientific thinking in young children. *Science*, 337, 1623-1627.
- Gosling, S. (2008). *Snoop: What your stuff says about you*. New York: Basic Books.
- Kelemen, D., & Carey, S. (2007). The essence of artifacts: Developing the design stance. In E. Margolis & S. Laurence (Eds.), *Creations of the Mind: Theories of Artifacts and Their Representation* (pp. 212-230). Oxford: Oxford University Press.
- Lipton, P. (2004). *Inference to the Best Explanation*, 2nd edition. Routledge Publishers.
- McCracken, G. (1986). Culture and consumption: A theoretical account of the structure and movement of the cultural meaning of consumer goods. *Journal of Consumer Research*, 71-84.
- Richins, M.L. (1994). Valuing things: the public and private meanings of possessions. *Journal of Consumer Research*, 504-521.
- Soley, G., & Spelke, E.S. (2016). Shared cultural knowledge: Effects of music on young children's social preferences. *Cognition*, 148, 106-116.
- Solomon, M.R. (1983). The role of products as social stimuli: A symbolic interactionism perspective. *Journal of Consumer Research*, 319-329.
- Teglas, E., Vul, E., Girotto, V., Gonzalez, M., Tenenbaum, J. B., & Bonatti, L. L. (2011). Pure Reasoning in 12-Month-Old Infants as Probabilistic Inference. *Science*, 332(6033), 1054-1059.
- Tenenbaum, J. B., Griffiths, T., & Kemp, C. (2006). Theory-based Bayesian models of inductive learning and reasoning. *Trends in Cognitive Sciences*, 10(7), 309-318.