

Detecting social transmission in the design of artifacts via inverse planning

Ethan Hurwitz (ehurwitz@ucsd.edu),
Timothy F. Brady (timbrady@ucsd.edu),
Adena Schachner (schachner@ucsd.edu)

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

Abstract

How do people use human-made objects (artifacts) to learn about the people and actions that created them? We test the richness of people's reasoning in this domain, focusing on the task of judging whether social transmission has occurred (i.e. whether one person copied another). We develop a formal model of this reasoning process as a form of rational inverse planning, which predicts that rather than solely focusing on artifacts' similarity to judge whether copying occurred, people should also take into account availability constraints (the materials available), and functional constraints (which materials work). Using an artifact-building task where two characters build tools to solve a puzzle box, we find that this inverse planning model predicts trial-by-trial judgments, whereas simpler models that do not consider availability or functional constraints do not. This suggests people use a process like inverse planning to make flexible inferences from artifacts' features about the source of design ideas.

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

Introduction

We live surrounded by human-made objects, or artifacts. These artifacts are crucial to our lives not only as tools, but also as an omnipresent source of social information. Based on the objects a person owns, people make quick and accurate judgments about a person's traits, interests, and social affiliations (Gosling, 2008; Richins, 1994). The artifacts a person creates - like novel tools, art, music, or text - provide particularly rich information about the person and actions that created them (Gosling, 2008).

How do people reason about other individuals from the artifacts they create? Here we explore the nature of this reasoning, a form of *intuitive archeology*. In the same sense that archeologists use objects to make inferences about the people and cultures that created them, we propose that people also infer complex social-causal information from the design of artifacts, by integrating their mental theories of the physical-mechanical world with their theories of the social world (e.g. Battaglia, Hamrick & Tenenbaum, 2013; Gopnik, 2012; Baker, Saxe & Tenenbaum, 2009) to infer the most probable explanation for an objects' features.

Intuitive Archeology as Inverse Planning

Previous work in the domain of action understanding has proposed that people make inferences about the goals of others' actions based on a process of 'inverse planning'

(Baker, Saxe, & Tenenbaum, 2009; Liu, Ullman, Tenenbaum & Spelke, 2018). The idea of inverse planning is that people have knowledge of the generative process behind actions from planning their own - and this planning process allows them to know what a rational agent would do, given the same goals and environmental constraints. Therefore, when reasoning about others' actions, people invert this generative process to infer the goals of another agent from its observed behaviors. Here we propose that a fundamentally similar inverse planning processing explains how we reason about the artifacts people create: People use their own generative model of how they would construct an artifact under a given set of constraints to infer the goals and decisions that led another person to create this artifact and its features. Such a reasoning process would allow people to flexibly infer a variety of social-causal information about others from the physical features of artifacts they create.

We focus on a foundational inference in this domain: Inferring whether *social transmission of ideas* has occurred (i.e. imitation, copying), or whether a particular aspect of a design was *generated independently* by an individual. The interaction of these two basic processes, termed imitation and innovation, account for cultural evolution of artifacts' designs over human history (Henrich, 2015; Tomasello, 1999; Legare & Neilsen, 2015). This inference also has real-world applications for understanding plagiarism detection - and what can be reasonably expected of jurors in plagiarism cases as they consider two designs and determine the likelihood that copying has occurred. Lastly, this inference is foundational to understanding how people infer social-causal information from artifacts, since designs that were created independently license different inferences than those that were copied. For example, a highly functional, complex design that was independently generated may tell you about the intelligence or creativity of a designer (Gosling, 2008), whereas a design that was copied may instead be informative about the designer's social history and cultural group (their source of shared knowledge; e.g. Schachner et al., 2018; Soley & Spelke, 2016). Thus, in the current work, we model and test how people infer whether or not copying (social transmission) occurred in the design of an artifact.

Inverse Planning, Or a Simpler Cognitive Process?

A natural alternative theory exists to the rich and structured explanation-based reasoning process proposed by inverse planning models. People may infer that copying occurred

using a simple heuristic based on perceptual similarity: If two things are more perceptually similar, then copying is more likely to have occurred. Notably, past work on detection of copying in music has relied on this type of simple similarity metric in formal models, to predict jury decisions in music plagiarism cases (Savage, Cronin, Müllensiefen, & Atkinson, 2018).

In contrast to these straightforward similarity-based models, other work has provided initial evidence that people detect copying via a more complex process of inverse planning or explanation-based reasoning (Schachner et al., 2018). In particular, this work found that people expect others to have a preference for efficiency, and factor this in when making inferences about copying. Thus, when two characters create identical train track designs that are also highly efficient ways to achieve the intended goal, observers use efficiency to ‘explain away’ the similarity – and thus judge copying less likely for identical efficient tracks than they would otherwise.

While this work is suggestive of a system of inverse planning, it is possible (and even plausible) that understanding of efficiency is unique and privileged in people’s reasoning. Reasoning about efficiency, and expecting others to act rationally by moving efficiently toward their goals, is thought to be foundational to cognition: It develops early in infancy (Gergely, Nádasdy, Csibra, & Bíró, 1995; Skerry, Carey & Spelke, 2013), is shared with other species (Hauser & Wood, 2010), and is a foundation for the entire domain of action understanding (Dennett, 1987; Baker et al. 2009). Thus, rather than showing a rich and flexible process of reasoning that takes into account a wide variety of alternative explanations (as proposed by inverse planning models), the evidence thus far is consistent with a much simpler system, in which similarity metrics are selectively overridden by privileged efficiency-based explanations.

The Current Work

In the current work, we test whether people use a rich and flexible process of inverse planning that takes into account alternative explanations that go beyond efficiency. In particular, we ask whether people rationally consider two factors: the range of materials available to build with, which we term the *availability constraint*; and whether each of the available materials would function or fail to function to solve the problem at hand, which we term the *functional constraint*. Rationally speaking, if a larger set of materials are available to choose from, similarity should be seen as stronger evidence of copying than if there is a smaller set of materials available to choose from (as the probability of selecting the same item by chance is lower; similar to the suspicious coincidence mechanism sometimes referred to as the ‘size principle’; Tenenbaum & Griffiths, 2001). Similarly, if many of these materials would solve the problem, similarity is more indicative of copying than if only one or a few of the options would solve the problem at hand – as clearly non-functional materials are unlikely to be used. We first formalize these

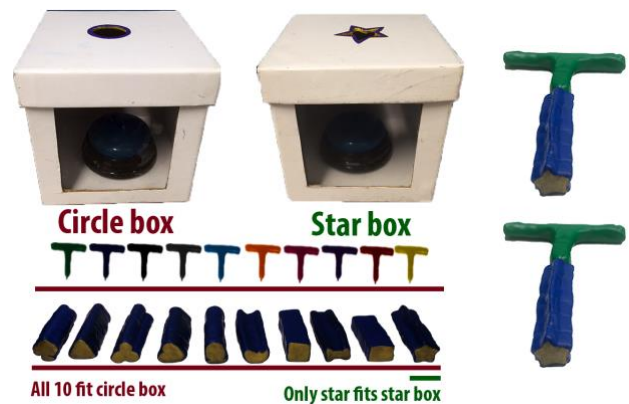


Figure 1: Left: Tool selection task with example handles (which differ in color), and rods (which differ in shape and therefore functionality). Right: Example of two identical tools people might be shown on a particular trial.

constraints and then experimentally test their usage when people make copying inferences.

An Inverse Planning Model of Copy Detection

To provide a clear test of the inverse planning account, and tease it apart from simpler alternatives, we model and test a simple artifact-building task which crucially involved both availability and functional constraints. Consider a scenario where one is asked to solve a puzzle: A button is out of reach in a box, with the front covered by glass, so only the hole in the top allows access. You must build a tool to reach the button. To do so, you are given two sets of pieces: 10 handles, which differ by color; and 10 rods, which differ by shape. You can connect one handle to one rod to form a two-part tool (see Figure 1).

You may be asked to solve one of two puzzle boxes, which differ in one respect: How many of the rods would work to solve them. In particular, for one box, all of the 10 rods would fit through the box’s circular hole and solve the puzzle (*unconstrained; circle box*). In the other case, only 1 of the 10 rods fits (only the star-shaped rod fits into the star-shaped hole), and so only 1 of the 10 rods can be used to solve the puzzle (*constrained; star box*). This box thus introduces a functional constraint that applies selectively to rods, and not handles (which would all function in both cases).

Now, you observe two tools that other people have made: for example, two people built the same tool, choosing the same star-shaped rod and the same red handle. How likely are they to have copied each other? This task provides a simple instantiation of relevant issues people confront when making complex decisions about copying through inverse planning: Reasoning about the range of materials available to the builders; which pieces would work; and a multi-part decision process (choose a handle, choose a rod).

Formally, we can think of this task as having the following structure: You see a tool built by person 1, and a second tool built by person 2, in order to solve a puzzle box. You wish to

infer whether person 2 copied the tool's design from person 1, or independently created it.

Each tool consists of two pieces linked together – a rod, r , and a handle, h – each of which was selected from the set of available options. Formally, you are asked to make an inference, where if c indicates whether person 2 copied person 1, you wish to infer the probability of copying

$P(c|r_1, h_1, r_2, h_2)$, given the observed rod and handle of person 1's tool (r_1, h_1) and the observed rod and handle of person 2's tool (r_2, h_2). Taking only the case of a rod being copied, and assuming copying judgments depend only on the rod and handle being identical or different (e.g., a binary notion of similarity), the posterior on copying is:

$$P(c | r_1, r_2) = \frac{P(c_r) P(r_2 == r_1 | c)}{P(r_2 == r_1)}$$

This is the probability that copying has occurred, given your prior likelihood on copying and the relative likelihoods that such an overlapping design would be generated under each of the possible mechanisms (copying, c , vs. independent creation, $\neg c$), where:

$$P(r_2 == r_1) = P(c_r) P(r_2 == r_1 | c) + (1 - P(c_r)) (r_2 == r_1 | \neg c)$$

In the current task this depends not only on the rod but on both the rod and handle, such that, when the rod is identical but the handle is not identical, this posterior on copying depends on $P(r_2 == r_1 | c, r_1)$, $P(r_2 == r_1 | \neg c)$, $P(h_2 \neq h_1 | c)$, and $P(h_2 \neq h_1 | \neg c)$. This has the structure of a Bayes net, including the key concept of explaining away: A given aspect of the design can be generated either via copying or independently, and evidence for one provides evidence against the other. Thus, if two people create identical tool designs, but this design is also likely to be created independently (due to either availability constraints or functional constraints), this provides weak evidence of copying despite the identical tools.

To make this model concrete, we need to specify 5 things:

(1) $P(c_r), P(c_h)$ - the a priori estimate of how likely person 2 was to have copied either the rod or handle (unconditional on the data; i.e. before we see either of the built objects). This depends for example on how close or distant the two people are from one another (Schachner et al., 2018). We assume the chance of copying is identical and independent for both rods and handles, e.g. $P(c_r) == P(c_h)$, and refer to this as $P(c)$, the prior on copying.

(2) $P(r_2 == r_1 | c)$ - the likelihood of the particular rod being used by person 2 matching that of person 1, given that person 2 was in fact copying from person 1's object. We formalize this as perfect copying plus a small error rate term, e , to account for the rate at which an individual might intend to copy but ultimately select a different rod: $P(r_2 == r_1 | c) = 1 - e$. Therefore $P(r_2 \neq r_1 | c) = e$.

(3) $P(r_2 == r_1 | \neg c)$ - the likelihood of rod r_2 being the same as r_1 , given that person 2 was NOT copying from person 1's object, and independently generated the object with no reliance on r_1 . When all pieces would function, this is simply $1/R$, where R is the total number of rod choices available. However, functional constraints also affect this

factor: When only a subset of pieces will function, this effectively reduces the number of reasonable options. Accordingly, in the context of a functional constraint, the model treats only the functional pieces as options, reducing the value of R to the number of functional options (if only one rod functions, $R=1$).

4) $P(h_2 == h_1 | c)$ - the likelihood of the particular handle being generated by person 2, given that person 2 was in fact copying from person 1's object, and given h_1 . This again is based on the same error rate e .

(5) $P(h_2 == h_1 | \neg c)$ - the likelihood of handle h_2 being the same as h_1 , given that person 2 was NOT copying from person 1's object, and independently generated the object with no reliance on h_1 . In contrast to the rods above, the handles differ only in color rather than shape; thus, all handles function equally well in both the *unconstrained* (circle box) condition, and the *functionally constrained* (star box) condition. This is therefore simply $1/H$, where H is the number of handle options.

Comparing to Simpler Alternatives

This model of inference as inverse planning posits that people consider both the number of available options and the functional constraint of the puzzle box when judging whether copying occurred. To test whether each of these components are needed to predict participants' judgments, we compared this model to three simpler models.

These models followed a 2x2 structure, either taking into account or not taking into account the availability constraints (+/- availability) or the functional constraints (+/- functional). For example, the model that considers availability constraints but ignores functional constraints does not take into consideration the functional constraint of the star box, e.g., assumes people choose among all rods even in the star box condition. The model which ignored availability constraints did not take into account the number of pieces available in a flexible way. Instead, this model posited that people had a fixed a-priori idea of the number of pieces available to choose from, and that this number did not change based on the situation presented. Thus, rather than choose a rod with $1/R$, where R is the number of options, a parameter N quantified this fixed number of imagined choices (e.g., regardless of how many were present). This model *did* take into account the functional constraint of the star box (assuming people only choose the star rod in this case). A final simplified model ignored both functional and availability constraints, and thus effectively instantiated a simple perceptual similarity heuristic. This model only took into account the extent to which the pieces were similar, without taking into consideration either functional constraints or availability constraints.

Testing the Models' Predictions

These models make quantitative predictions about the likelihood of copying for any given pair of tool designs, in a wide range of contexts. We next aimed to test how well the various models predict human behavior. The inverse

planning model predicts that for two identical tools, people will infer that copying is more likely to have occurred when (a) there were more pieces available as options to build with, thus creating more of a suspicious coincidence that the same piece was chosen twice; (b) there were no functional constraints on which pieces would work or not work, thus allowing all of the available pieces to serve as equally good options. By contrast, the simplest perceptual similarity model predicts that any identical objects will lead people to infer copying. Thus, we focused our data collection on these and other particularly informative trials.

Method

Full study design/analysis plan including model code was preregistered on the Open Science Framework (OSF), and is available at <https://osf.io/y8u7t>.

Participants

Using a pre-registered design, $N=108$ adults from the U.S. (57 male, 50 female, 1 other gender identity; M age=37.9, $SD=10.9$, range=20-72) were recruited through Amazon's Mechanical Turk. Sample size was preregistered and determined from power analysis of a pilot dataset with a slightly different design ($N=20$; tested a subset of the current test trials; with each subject completing all trials). The R "pwr" package was used to conduct a paired t-test power calculation on participant-level BICs with the goal of 90% power (Champely et al., 2018). Based on pre-registered exclusion criteria, additional participants were excluded due to: 1. Appearing to be non-native English speakers or a bot ($n=13$; determined by 2 independent coders' rating of free-response text answers) 2. Incorrectly answering any memory check question ($n=49$) 3. Incorrectly answering 50% or more of the attention check questions ($n=12$). The number of participants failing the preregistered memory check questions was higher than expected, thus we reanalyzed the data with these participants included, and found that our model results and conclusions remain unchanged in this case (see *Results*).

Design

Participants were shown tools that two target individuals designed, and were asked to judge whether or not one of those individuals copied the other's tool. Across trials we manipulated (1) the number of rod options available (2 versus 10); (2) the number of handle options available (2 versus 10); (3) The presence or absence of a functional constraint, i.e. whether they were trying to solve the circle or star puzzle box; (4) The extent of similarity of the two tools that were built (both rod and handle identical, one part identical and one part different, or both rod and handle different). As all designers were assumed to have successfully solved the puzzle, we did not include trials in the star box condition which had different rods, as this would involve building a tool that would not function. Thus in total there were 24 unique test trials. Because of the possibility of demand characteristics if all participants saw the full design, each participant completed only a randomly-selected subset

of 4 trials, resulting in 18 unique participants completing each trial.

Procedure

Participants first received instructions regarding the puzzle-box task, and that they would see pairs of tools that people had built to reach the button. Instructions described an ambiguous situation, where copying may or may not have occurred ("While designing the tools the people were in the same room, facing away from each other"). They were instructed that different pairs of people had different numbers of handles and rods to choose from (10 or 2), received either the circle box or star box to solve, and that only one of the rod pieces could fit into the star-shaped opening.

On each trial, participants saw (1) the two tools that the two people had built; (2) which puzzle box the people were trying to solve; (3) the materials they had available to build with. Participants were asked to judge as a 2-alternative forced choice: Do you think someone copied, or they made them independently?

After each trial, an attention check question asked either what puzzle box was present, the number of rod options, or number of handle options. At the end of the task, memory check questions asked participants to select which rods would work, and which handles would work, to successfully solve each of the two puzzle boxes. Lastly, participants were asked to describe what they did in the experiment and guess the point of the study in free-response format, and complete demographics questions.

Analysis Plan

For each model, the best fitting parameters and likelihood of our data given those parameters were assessed via maximum likelihood estimation (MLE). We decided a priori that the prior on copying (range: 0-1) and number of imagined choices (for models that do not use the real number that participants were presented with; range 0-infinity) should be fully free to vary, while the copying error rate e was bounded from 0 to a maximum of 0.1. For all models, using this a priori specification, the MLE-derived value for the copying error rate was at max (0.1). To make sure this boundedness was not responsible for our findings, we also reran analyses letting the error rate parameter vary (0-1), and found the same results for comparative model fits in this case. To compare models, we use BIC (Schwarz, 1978), which penalizes models for complexity according to their number of parameters. We used bootstrapping to calculate standard errors (SEs) for each BIC.

Results

We first checked that participants took into account the perceptual similarity of designs in their assessments of copying, as predicted by all four models. As expected, participants inferred copying most often when the two tool designs were identical ($M=51.4\%$, $SEM=9.8\%$), and least often when the two tools were most different ($M=5.6\%$, $SEM=2.3\%$; $p<.01$).

Table 1: Maximum likelihood parameters for each model

Model	Copying Prior, $p(c)$	Error Rate	Imagined # Options
+Availability +Functional	0.09	0.10	
+Availability -Functional	0.06	0.10	
-Availability +Functional	0.09	0.10	5.31
-Availability -Functional	0.11	0.10	2.76

We next compared the fit of the four alternative models. The full model out-performed all competing models, with a difference in BIC of 35 ($\mp$ SEM: 11-25) in comparison to the next-best-fit model and >400 to the other models (Table 2). Approximately the same results held when including individuals who failed the memory check: difference in BIC of 38 to next-best-fit model and >700 to the other models. In addition, the full model provided a good overall fit to participants' responses across trials ($R^2 = 0.75$, Fig. 2A).

Note that while the model is relatively straightforward to specify, the predictions it makes are quite nuanced: because the model weighs and combines several factors, it predicts a continuous gradient of how likely copying should be, rather than simply saying people should never assume copying took place if there is any alternative explanation. The model thus goes well beyond verbal theories.

Use of Availability Constraints

Participants' judgements showed sensitivity to availability constraints (i.e. the number of pieces available to build with), and the use of availability constraints as an alternative explanation for similarity. For example, on trials where two people made identical tools and no functional constraint was present, participants judged copying more likely as the number of available options increased (circle box condition: 2 rods; 2 handles: 33% judged copied; 2 rods, 10 handles: 72%; 10 rods, 2 handles: 72%; 10 rods, 10 handles: 83%).

Use of Functional Constraints

Participants also showed sensitivity to functional constraints, and used functional constraints as an alternative explanation for similarity. In particular, on trials where two people used identical *rods*, participants judged copying less likely on trials where they were solving the star box (which added a functional constraint; Mean copied=21.5%), vs. when they were solving the circle box (Mean copied=52.8%, $p=0.02$, 2 tailed t-test). In contrast, on trials where two people used identical *handles*, participants' judgements did not differ for the star vs. circle box (Star box: Mean copied=36.8%, Circle Box: Mean copied=37.5%; $p=0.97$, 2 tailed t-test), as predicted since all handle pieces would function equally well for both puzzle boxes. Although the model without functional constraints did not perform that poorly as measured by BIC, it did systematically miss this aspect of the data (see also deviations of this model in Figure 2).

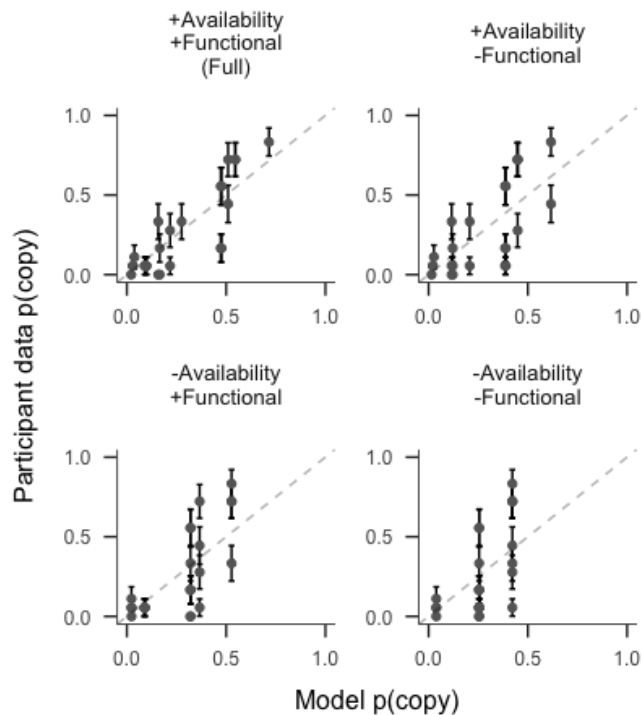


Figure 2: Fit of models' predictions to participants' ratings of whether copying occurred; each point represents one trial. The full inverse planning model appears top left; other plots show three simpler alternative models that do not consider either the availability constraints (-availability) or the functional constraints (-functional).

Table 2: Difference in BIC from best fitting model (higher BIC indicates worse fit)

Model	BIC Δ to full model	$\mp$ SEM
+Availability -Functional	35	11 - 25
-Availability +Functional	467	422 - 490
-Availability -Functional	491	468 - 512

Participants' judgments deviated slightly from the full model's predictions in one regard: Participants appeared to under-weight the similarity of the handles, relative to the rods. For instance, the largest deviations between participants' judgements and the full model's predictions came on trials when the tools had different rods, but the same handle. To demonstrate this differential weighing of the rod vs. handle, consider trials where there are an equal number of rod and handle options, no functional constraint, and the built tools had only one similar piece. On these trials, people were considerably more likely to say the design was copied if the rod was similar than if the handle was (2 options: 0% vs. 17%; 10 options: 17% vs. 56%). Thus, participants seemed to overweight evidence from the functionally-relevant component of the tool, even when functional constraints were not present.

Overall, however, the good fit of the inverse planning model – and the continuous range of predictions it makes –

supports the idea that participants use an inverse planning strategy in judging copying from artifacts.

Discussion

We find strong evidence that when reasoning about artifacts, people use a rich, flexible system of explanation-based reasoning to infer whether a design idea was copied or generated independently. We formalized such reasoning in a Bayesian model as a form of inverse planning. We compared this model to three simpler alternatives in a task where participants had to judge whether a pair of artifacts was copied or designed independently, to test whether each component of the full model was needed to predict judgments.

We found that the full inverse planning model best predicted participants' judgments of whether copying had occurred. In line with the model, we found that people considered two broad classes of alternative explanations for artifacts' similarity: the range of materials available to build with (availability constraints), and which of these materials would work to solve the problem (functional constraints). Both of these constraints 'explained away' similarity, making similarity weaker evidence of copying. This pattern of responses is the signature pattern of a Bayesian reasoner, in which a design can have different alternative explanations, and evidence for one provides evidence against the other (e.g., Gopnik et al. 2004).

The success of this model suggests people use a process of inverse planning to infer the source of design ideas from artifacts' features. In other words, people consider the generative processes involved in building the artifacts, including what the goal would be, what constraints they would be subject to, and what (as a result) they would be likely to build. By inverting this generative process, people rationally infer the source of other people's design ideas, taking into account goals and multiple kinds of constraints.

These findings show that inferences about the source of design ideas do not boil down to various simpler heuristics, or more limited systems of reasoning. First, copying judgments are not just based on the extent of perceptual similarity of the two objects, but take into account rational explanations for this similarity. This has implications for understanding how laypeople detect plagiarism in court cases, which has been previously formalized as a process of simple similarity detection (Savage et al., 2018).

Second, we show that this system of reasoning goes beyond efficiency: People can take into account multiple types of constraints as explanations for similarity, and are not limited only to reasoning about design efficiency as the only, privileged type of alternative explanation. This simpler efficiency-only account was consistent with previous findings, and plausible given the foundational role of efficiency in reasoning about intentional action (Schachner et al., 2018). The current data falsify this simpler account, showing that people flexibly take into account the materials available and the functional constraints of the puzzle boxes,

which do not map to an efficiency metric (e.g. the length of a train track from A to B, used in Schachner et al., 2018).

More broadly, we provide evidence for a novel theoretical and formal framework for artifact cognition, as a form of inverse planning. Previous work has shown that people use inverse planning to understand the causal processes underlying others' actions (Baker et al., 2009; Liu et al. 2018). The current work extends this framework by conceptualizing artifacts as the products of intentional action. We suggest that people use fundamentally the same inverse planning process to understand artifacts as they do to understand actions themselves. Specifically, they rationally take into account people's goals and constraints not only when observing actions, but also when observing artifacts generated by these actions – even when the actions themselves are not observed. This work thus links together artifact cognition and theories of action understanding in a new way, points to a deep connection between reasoning about actions and artifacts, and provides a foundation for formalizing the processes underlying a domain of 'intuitive archeology' – social-causal reasoning about artifacts, as products of intentional action.

Acknowledgements

We thank Michelle Lee for her help with qualitative coding and stimulus design, as well as Carissa Jantz and Kimberly McGee for stimulus preparation. This material is based upon work supported by the National Science Foundation Grant No. BCS-1749551 to AS and TFB.

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.
- Dennett, D.C. (1987). *The Intentional Stance*. MIT Press, Cambridge, MA.
- Gergely, G., Nádasdy, Z., Csibra, G., & Bíró, S. (1995). Taking the intentional stance at 12 months of age. *Cognition*, 56(2), 165-193.
- Gopnik, A., Glymour, C., Sobel, D. M., Schulz, L. E., Kushnir, T., & Danks, D. (2004). A theory of causal learning in children: causal maps and Bayes nets. *Psychological Review*, 111(1), 3-32.
- Gopnik, A. (2012). Scientific Thinking in Young Children: Theoretical Advances, Empirical Research, and Policy Implications. *Science*, 337(6102), 1623–1627.
- Gosling, S. (2008). *Snoop: What your stuff says about you*. Profile Books.
- Hauser, M., & Wood, J. (2010). Evolving the capacity to understand actions, intentions, and goals. *Annual Review of Psychology*, 61, 303-324.

- Henrich, J. (2015). *The secret of our success: How culture is driving human evolution, domesticating our species, and making us smarter*. Princeton University Press.
- Legare, C. H., & Nielsen, M. (2015). Imitation and innovation: The dual engines of cultural learning. *Trends in Cognitive Sciences*, 19(11), 688-699.
- Liu, S., Ullman, T. D., Tenenbaum, J. B., & Spelke, E. S. (2017). Ten-month-old infants infer the value of goals from the costs of actions. *Science*, 358(6366), 1038-1041.
- Richins, M. L. (1994). Valuing things: The public and private meanings of possessions. *Journal of Consumer Research*, 21(3), 504-521.
- Savage, P. E., Cronin, C., Müllensiefen, D., & Atkinson, Q. D. (2018). Quantitative evaluation of music copyright infringement. In *Proceedings of the 8th International Workshop on Folk Music Analysis (FMA2018)*, 61-66.
- Schachner, A., Brady, T.F., Oro, K., & Lee, M. (2018). Intuitive archeology: Detecting social transmission in the design of artifacts. In C. Kalisch, M. Rau, T. Rogers, & J. Zhu, *Proceedings of the 40th Annual Conference of the Cognitive Science Society*.
- Schwarz, G. (1978). Estimating the Dimension of a Model. *The Annals of Statistics*, 6(2), 461-464.
- Skerry, A. E., Carey, S. E., & Spelke, E. S. (2013). First-person action experience reveals sensitivity to action efficiency in prereaching infants. *Proceedings of the National Academy of Sciences*, 110(46):18728-33.
- Soley, G., & Spelke, E. S. (2016). Shared cultural knowledge: Effects of music on young children's social preferences. *Cognition*, 148, 106-116.
- Tenenbaum, J. B., & Griffiths, T. L. (2001). Generalization, similarity, and Bayesian inference. *Behavioral and Brain Sciences*, 24(4), 629-640.
- Tomasello, M. 1999. *The Cultural Origins of Human Cognition*. Cambridge, MA: Harvard University Press.