



Win-Stay, Lose-Sample: A simple sequential algorithm for approximating Bayesian inference



Elizabeth Bonawitz^{a,*}, Stephanie Denison^b, Alison Gopnik^c,
Thomas L. Griffiths^c

^a Department of Psychology, Rutgers University – Newark, Newark, NJ USA

^b Department of Psychology, University of Waterloo, Waterloo, ON Canada

^c Department of Psychology, University of California, Berkeley, USA

ARTICLE INFO

Article history:

Accepted 27 June 2014

Keywords:

Bayesian inference

Algorithmic level

Causal learning

ABSTRACT

People can behave in a way that is consistent with Bayesian models of cognition, despite the fact that performing exact Bayesian inference is computationally challenging. What algorithms could people be using to make this possible? We show that a simple sequential algorithm “Win-Stay, Lose-Sample”, inspired by the Win-Stay, Lose-Shift (WSLS) principle, can be used to approximate Bayesian inference. We investigate the behavior of adults and preschoolers on two causal learning tasks to test whether people might use a similar algorithm. These studies use a “mini-microgenetic method”, investigating how people sequentially update their beliefs as they encounter new evidence. Experiment 1 investigates a deterministic causal learning scenario and Experiments 2 and 3 examine how people make inferences in a stochastic scenario. The behavior of adults and preschoolers in these experiments is consistent with our Bayesian version of the WSLS principle. This algorithm provides both a practical method for performing Bayesian inference and a new way to understand people's judgments.

© 2014 Elsevier Inc. All rights reserved.

* Corresponding author. Address: 334 Smith Hall, Department of Psychology, Rutgers University – Newark, Newark, NJ 07102, USA.

E-mail address: elizabeth.bonawitz@rutgers.edu (E. Bonawitz).

1. Introduction

In the last decade a growing literature has demonstrated that both adults and preschool children can act in ways that are consistent with the predictions of Bayesian models in appropriate contexts (e.g., Griffiths & Tenenbaum, 2005, 2009; Goodman, Tenenbaum, Feldman, & Griffiths, 2008; Gopnik & Schulz, 2004; Gopnik et al., 2004). Bayesian models provide a framework for more precisely characterizing intuitive theories, and they indicate how a rational learner should update her beliefs as she encounters new evidence. The Bayesian approach has been used to explain the inferences adults make about causal relationships, words, and object properties (for a review see Tenenbaum, Griffiths, & Kemp, 2006; for critiques of this approach and responses, Jones & Love, 2011; Chater et al., 2011; Bowers & Davis, 2012; Griffiths, Chater, Norris, & Pouget, 2012; McClelland et al., 2010; Griffiths et al., 2010). The approach can also explain the inferences that children make in these domains (for reviews see Gopnik, 2012; Gopnik & Wellman, 2012; Xu and Kushnir, 2012). Taken together, these findings suggest that Bayesian inference may provide a productive starting point for understanding human learning in these domains.

Bayesian models have typically been used to give a “computational level” analysis of the inferences people make when they solve inductive problems (Marr, 1982). That is, they focus on the form of the computational problem and its ideal solution. Bayesian models are concerned with the computational problem of how a learner should update her beliefs given new evidence. The ideal Bayesian solution selects the most likely hypothesis after taking the evidence into account. However, it is extremely unlikely that people are performing exact Bayesian inference at the *algorithmic* level as such inferences can be extremely computationally challenging (e.g., Russell & Norvig, 2003). In particular, it would be computationally costly to always enumerate and search through all the hypotheses that might be compatible with a particular pattern of evidence. Indeed, it has been generally acknowledged since the pioneering work of Simon (1955, 1957) that human cognition is at best “boundedly rational”.

We might expect that rather than explicitly performing Bayesian inference, adults and children use “fast and frugal” heuristics (Gigerenzer & Gaissmaier, 2011; Gigerenzer & Goldstein, 1996) that approximate Bayesian inference. Such heuristic procedures might lead to apparently irrational inferences when considered step by step, even though in the long run those inferences would lead to a rational solution. An algorithmic account of learning could help explain both how people might produce inferences that can approximate rational models and why their step-by-step learning behavior might appear irrational or protracted.

Thus, algorithmic accounts should approximate Bayesian inference, but they also need to be computationally efficient. The need for efficiency is particularly important when we consider children’s learning. Young children have particularly limited cognitive resources, at least in some respects (e.g., German & Nichols, 2003; Gerstadt, Hong, & Diamond, 1994; Siegler, 1975), but are nonetheless capable of behaving in a way that is consistent with optimal Bayesian models (Gopnik & Tenenbaum, 2007). Children must thus be especially adept at managing limited resources to approximate Bayesian inference. At the same time, many of the most interesting cases of belief revision happen in the first few years of life (Bullock, Gelman, & Baillargeon, 1982; Carey, 1985; Gopnik & Meltzoff, 1997; Wellman, 1990). Understanding more precisely how existing beliefs and new evidence shape children’s learning given limited cognitive resources is a key challenge to proponents of Bayesian models of children’s cognition.

Here we investigate the algorithms that learners might be using to solve a particular kind of inductive problem – causal learning. Children learn a great deal about causal structure from very early in development (e.g. Gopnik & Meltzoff, 1997; Gopnik & Schulz, 2007; Wellman & Gelman, 1992) and causal knowledge and learning have been a particular focus of formal work in machine learning and philosophy of science (Pearl, 2000; Spirtes, Glymour, & Scheines, 2000). Causal learning has also been one of the areas in which Bayesian models have been particularly effective in characterizing human behavior, from adults learning from contingency data (Griffiths & Tenenbaum, 2005) to children reasoning about causal systems and events (Bonawitz, Gopnik, Denison, & Griffiths, 2012; Bonawitz et al., 2011; Griffiths, Sobel, Tenenbaum, & Gopnik, 2011; Gweon, Schulz, & Tenenbaum, 2010; Kushnir & Gopnik, 2007; Schulz, Bonawitz, & Griffiths, 2007; Seiver, Gopnik, & Goodman, 2012).

Most studies of causal learning by children and adults simply relate the total amount of evidence that a learner sees to the final inferences they make. This technique is well-suited to analyzing the correspondence between the learner's behavior and the predictions of computational-level models. However, a more effective way to determine the actual algorithms that a learner uses to solve these problems is to examine how her behavior changes, trial-by-trial, as new evidence is accumulated. We describe a new “mini-microgenetic” method (cf. [Siegler, 1996](#)) that allows us to analyze the changes in a learner's knowledge as she gradually accumulates new pieces of evidence. We can then compare this behavior to possible algorithmic approximations of Bayesian inference.

Computer scientists and statisticians also face the challenge of finding efficient methods for performing Bayesian inference. One common strategy for performing Bayesian inference is to use approximations based on sampling, known as Monte Carlo methods. Under this strategy, computations involving a probability distribution are conducted using samples from that distribution – an approach that can be much more efficient. We introduce a new sequential sampling algorithm inspired by the Win-Stay, Lose-Shift (WSLS) strategy,¹ in which learners maintain a particular hypothesis until they receive evidence that is inconsistent with that hypothesis. We name our strategy Win-Stay, Lose-Sample to distinguish it from the original strategy. We show that our WSLS algorithm approximates Bayesian inference, and can do so more efficiently than enumerating hypotheses or performing simple Monte Carlo search. Previous work in cognitive psychology demonstrates that people use a simple form of the Win-Stay, Lose-Shift strategy in concept learning tasks ([Levine, 1975](#); [Restle, 1962](#)).

In previous experiments that compared human performance to the predictions of Bayesian models, researchers have typically looked at responses at just one point – after the evidence has been accumulated. In these tasks, the distribution of these responses often matches the posterior distribution that would be generated by Bayesian inference. That is, the proportion of people who endorse a particular hypothesis is closely related to the posterior probability of that hypothesis (e.g., [Denison, Bonawitz, Gopnik, & Griffiths, 2013](#); [Goodman et al., 2008](#); [Vul, Goodman, Griffiths, & Tenenbaum, 2009](#); [Xu & Tenenbaum, 2007](#); though, see also [Bowers & Davis, 2012](#); [Jones & Love, 2011](#); [Marcus & Davis, 2013](#)). However, the psychological signature of WSLS algorithms is a characteristic pattern of dependency between people's successive responses as the evidence comes in. In general terms, a learner who uses a WSLS algorithm will be “sticky”; that is, they will show a tendency to prefer the hypothesis currently held to other possible hypotheses. What we demonstrate mathematically is that, despite an individual's tendency toward “stickiness”, the overall proportion of individuals selecting a hypothesis at a given point in time will match the probability of that hypothesis given by Bayesian inference at that point in time.

To test whether the signature pattern of dependence produced by our WSLS algorithm is present in people's judgments, we systematically compare this algorithm to an algorithm that uses independent sampling. In the independent sampling algorithm, the learner simply resamples a hypothesis from the new distribution yielded by Bayesian inference each time additional evidence is acquired. Independent sampling is similar to our WSLS algorithm in that it would also approximate the predictions of a Bayesian model. However, by definition, it would not have the characteristic dependence between responses, as would WSLS. This contrast to independent sampling allows us to test systematically whether learner's responses are indeed dependent in the way that our WSLS algorithm would predict.

The plan for the rest of the paper is as follows. Next, we introduce the causal learning tasks that will be the focus of our analysis and summarize how Bayesian inference can be applied in this task. We then introduce the idea of sequential sampling algorithms, including our new WSLS algorithm. This is followed by a mathematical analysis of the WSLS algorithm, showing that it approximates Bayesian inference in deterministic and non-deterministic cases. The remainder of the paper focuses on experiments in which we evaluate how well both Bayesian inference, in general, and the WSLS algorithm, in particular, capture judgments by children and adults in two causal learning tasks in which the learner gradually acquires new evidence.

¹ Also known as Win-Stay, Lose-Switch strategy ([Robbins, 1952](#)).

2. Analyzing sequential inductive inferences

While the algorithms that we present in this paper apply to any inductive problem with a discrete hypothesis space, we will make our analysis concrete by focusing on two simple causal learning problems, which are representative of the larger causal learning problems that children and adults face. In the experiments that appear later in the paper we provide actual data from adults and children solving these problems, but we will begin by presenting a Bayesian computational analysis of how these problems might be solved ideally. The first of these two problems involves deterministic events, and the second involves non-deterministic events, where object behavior is stochastic. The distinctive feature of these learning problems is that participants are asked to provide predictions over multiple trials, as new evidence is observed. This allows us to model potential dependencies between an individual learner's responses as her beliefs are updated following each successive piece of new evidence.

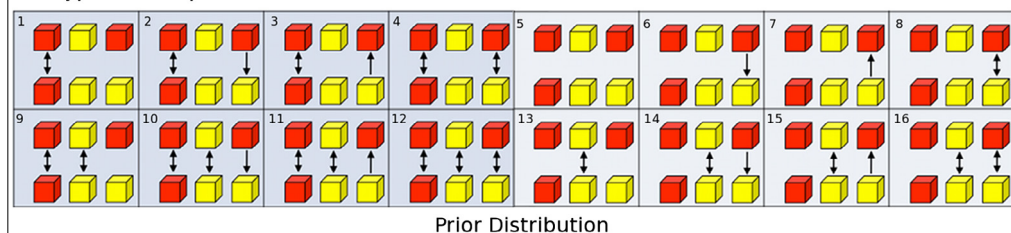
2.1. Deterministic events

Deterministic events provide a simple starting point for considering how a learner might update his or her beliefs. In the deterministic case the evidence a learner observes will necessarily rule out certain hypotheses. As more and more evidence accumulates fewer hypotheses will remain.

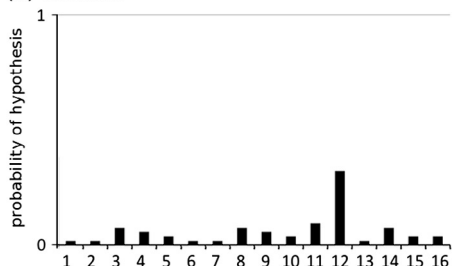
For example, consider the following causal learning problem: there are two categories of objects (yellow and red blocks). Some of these blocks light up when they come into contact with other blocks. Blocks of the same color all behave similarly, and will either light or not light when they interact with other blocks. Given these constraints, we can generate a hypothesis space of 16 possible causal structures, as illustrated in Fig. 1(a).

Using this hypothesis space, we can consider how an ideal learner should update his or her beliefs in light of evidence. Assume that the learner begins with a prior distribution over hypotheses, $P(h)$, reflecting her degree of belief that each hypothesis is true before seeing any data. Given some

(a) Hypothesis Space



(b) Children



(c) Adults

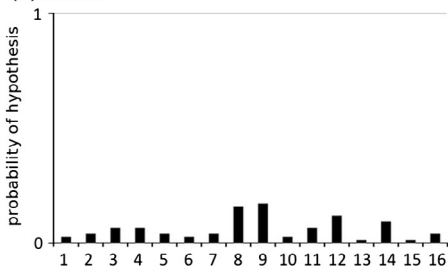


Fig. 1. A causal induction problem. (a) The sixteen possible hypotheses illustrated with the red (dark) and yellow (light) blocks; arrows indicate the causal direction of lighting. (b) The children's prior distribution over the sixteen hypotheses. (c) The adult's prior distribution over the sixteen hypotheses. (For interpretation of the references to color in this figure legend, the reader is referred to the web version of this article.)

observed data, d , on how the blocks interact, the learner obtains a posterior distribution over hypotheses, $P(h|d)$, via Bayes' rule:

$$P(h|d) = \frac{P(d|h)P(h)}{\sum_{h' \in H} P(d|h')P(h')} \quad (1)$$

where $P(d|h)$ is the likelihood, indicating the probability of observing d if h were true, and H is the hypothesis space.

In our causal induction problem, we estimate a prior distribution over the 16 hypotheses from the initial guesses that learners make about the causal structure (see Results of Experiment 1 for details). This distribution is shown in Fig. 1(b). The data that learners observe consists of a single interaction between a red and a yellow block. No observations are made about red–red or yellow–yellow interactions, and so a number of hypotheses remain consistent with this evidence. We assume a deterministic likelihood, such that if a causal relationship exists it always manifests, with each block from the relevant class lighting up each block from the other class. Consequently, we have

$$P(d|h) = \begin{cases} 1 & h \text{ consistent with } d \\ 0 & \text{otherwise} \end{cases} \quad (2)$$

where the likelihood is simply 1 or 0, depending on whether or not the data d could be generated by h .

Often, Bayesian inference must be performed sequentially. The learner makes a series of observations, one after another, and the posterior distribution is updated after each observation. In our causal learning problem, learners could receive a sequence of observations of blocks interacting, rather than a single observation. Letting $d_1, \dots, d_n$ denote observations after n trials, we are interested in the posterior distribution $P(h|d_1, \dots, d_n)$. This can be computed via Equation (1), substituting $d_1, \dots, d_n$ for d . However, it can be simpler to follow a sequential updating rule, which allows us to compute the posterior after observing $d_1, \dots, d_{n+1}$ from the posterior based on $d_1, \dots, d_n$. Formally, we assume that the observations d_i are conditionally independent given h (i.e., that the probability that blocks interact in a particular way on a given trial is independent of their interaction on all other trials, once the nature of the underlying causal relationship is known). In this case, we can perform Bayesian inference sequentially using the update rule

$$P(h|d_1, \dots, d_{n+1}) = \frac{P(d_{n+1}|h)P(h|d_1, \dots, d_n)}{\sum_{h' \in H} P(d_{n+1}|h')P(h'|d_1, \dots, d_n)} \quad (3)$$

where the posterior after the first n observations plays the role of the prior for observation $n + 1$.

With deterministic events, sequentially performing Bayesian inference reduces to narrowing down the hypothesis space, with the posterior distribution being the prior renormalized over the remaining hypotheses. That is, with each observation, the learner eliminates some hypotheses and reconsiders the remaining hypotheses. This can be seen by substituting the likelihood from Eq. (2) into the update rule given in Eq. (3). On observing the first piece of data, d_1 , the likelihood will give hypotheses inconsistent with the data a posterior probability of zero. The hypotheses consistent with the data will remain, with probability proportional to their prior probability, and the sum in the denominator ranges over just those hypotheses. The same process takes place with each subsequent piece of data, with the posterior at each point being the prior renormalized over the hypotheses consistent with the data observed up to that point.

2.2. Stochastic events

Determinism is a reasonable assumption in some causal learning contexts (Griffiths & Tenenbaum, 2009; Schulz, Hoopell, & Jenkins, 2008; Schulz & Somerville, 2006), but learners also need to be able to deal with stochastic data. The stochasticity may be due to noisy observations – a result of noise in the demonstration or noise in the child's observation – or it could reflect the presence of a non-deterministic causal mechanism.

We designed a stochastic causal learning task to explore the inferences that adults and children make for non-deterministic systems. In this task, there are three categories of objects: red, green,

and blue blocks. Each of these kinds of blocks activates a machine with different probability when they are placed on the machine. The red blocks activate the machine on five out of six trials, the green blocks on three out of six trials, and the blue blocks on just one out of six trials. A new block is then presented that has lost its color and needs to be classified as either a red, green, or blue block. The learner observes what happens when the block is placed on the machine over a series of trials. What should learners infer about the category and causal properties of this block as they gradually acquire more evidence about its effects?

As with the deterministic case, given this hypothesis space, we can compute an ideal learner's posterior distribution over hypotheses via Bayes' rule (Eq. (1)). In this case, the likelihood takes on a wider range of values. If we take our data d to be the activation of the machine, the probability of this event depends on our hypothesis about the color of the block. For the red block, we might take $P(d|h)$ to be 5/6; for the green block, 3/6; and for the blue block, 1/6. Importantly, observing the block light up or not light up the detector no longer clearly rules out any of these hypotheses. Bayesian inference thus becomes a matter of balancing the evidence provided by the data with our prior beliefs. We can still use the sequential updating rule given in Eq. (3), but now each new piece of data is only going to steer us a little more toward one hypothesis or another.

2.3. Toward algorithms for Bayesian inference

This analysis of sequential causal learning in deterministic and nondeterministic scenarios allows us to think about how learners should select hypotheses by combining prior beliefs with evidence. However, assuming that learners' responses do reflect this optimal trade-off between priors and evidence, this analysis makes no commitments about the algorithms that might be used to approximate Bayesian inference. In particular, it makes no commitments about how an individual learner will respond to data, contingent on her previous guess, other than requiring that both responses be consistent with the associated posterior distributions. Given a potentially large hypothesis space, searching through that space of hypotheses in a way that yields something approximating Bayesian inference is a challenge. In the remainder of the paper, we investigate how a learner might address this challenge.

3. Sequential sampling algorithms

The Bayesian analysis presented in the previous section provides an abstract, “computational level” characterization of causal induction for these two causal learning tasks, identifying the underlying problem and how it might best be solved (Marr, 1982). We now turn to the problem of how to approximate this optimal solution.

One way a learner might compute an optimal solution is by simply following the recipe provided by Bayes' rule. However, enumerating all hypotheses and then updating their probabilities by multiplying prior and likelihood, quickly becomes computationally expensive. We thus consider the possibility that people may be approximating Bayesian inference by following a procedure that produces samples from the posterior distribution.

3.1. Approximating Bayesian inference by sampling

Sophisticated Monte Carlo methods for approximating Bayesian inference developed in computer science and statistics make it possible to draw a sample from the posterior without having to calculate the full distribution (see, e.g., Robert & Casella, 2004). Thus, rather than having to evaluate every single hypothesis and work through Bayes' rule, these methods make it possible to consider only a few hypotheses that have been sampled with the appropriate probabilities. These few hypotheses can be evaluated using Bayes rule to generate samples from the posterior distribution. Aggregating over these samples provides an approximation to the posterior distribution. That is, for approximate algorithms, as the number of samples they generate increase, the distribution of samples converges to the normative distribution. These sampling algorithms are consistent with behavioral evidence that

both adults and children select hypotheses in proportion to their posterior probability (Denison et al., 2013; Goodman et al., 2008).

Note that these sampling algorithms also address a conflict between rational models of cognition and the finding that people show significant individual variability in responding (e.g. Siegler & Chen, 1998). A system that uses this sort of sampling will be variable – it will entertain different hypotheses apparently at random from one time to the next. However, this variability will be systematically related to the probability distribution of the hypotheses, as more probable hypotheses will be sampled more frequently than less probable ones. And, on aggregate, the distribution of hypotheses from many learners (or one learner across multiple trials), will approximate the probability distribution. Recent research has explored this “Sampling Hypothesis” generally and suggested that the variability in young children’s responses may be part of a rational strategy for inductive inference (Bonawitz, Denison, Griffiths, & Gopnik, 2014; Denison et al., 2013).

The idea that people might be generating hypotheses by sampling from the posterior distribution still leaves open the question of what kind of sampling scheme they might be using to update their beliefs. Independent Sampling (IS) is the simplest strategy, and is thus a parsimonious place to start considering the algorithms learners might use. In Independent Sampling, a learner would generate a sample from the posterior distribution independently each time one was needed. This can be done using a variety of Monte Carlo schemes such as importance sampling (see Neal, 1993, for details) and Markov chain Monte Carlo (see Gilks, Richardson, & Spiegelhalter, 1996, for details). Such schemes have previously been proposed as possible accounts of human cognition, (Shi, Feldman, & Griffiths, 2008; Ullman, Goodman, & Tenenbaum, 2010). However, this approach can be inefficient in situations where the posterior distribution needs to be approximated repeatedly as more data become available. Independent Sampling requires that the system recompute the approximation to the posterior after each piece of data is observed.

People might adopt an alternative strategy that exploits the sequential structure of the problem of updating beliefs. The problem of sequentially updating a posterior distribution in light of evidence can be solved approximately using sequential Monte Carlo methods such as particle filters (Doucet, de Freitas, & Gordon, 2001). A particle filter approximates the probability distribution over hypotheses at each point in time with a set of samples (or “particles”), and provides a scheme for updating this set to reflect the information provided by new evidence. The behavior of the algorithm depends on the number of particles. With a very large number of particles, each particle is similar to a sample from the posterior. With a small number of particles, there can be strong sequential dependencies in the representation of the posterior distribution. Recent work has explored particle filters as a way to explain patterns of sequential dependency that arise in human inductive inference (Levy, Real, & Griffiths, 2009; Sanborn, Griffiths, & Navarro, 2010).

Particle filters have many degrees of freedom, and many different schemes for updating particles are possible (Doucet et al., 2001). They also require learners to maintain multiple hypotheses at each point in time. Here, we investigate a simpler algorithm that assumes that learners maintain a single hypothesis, only occasionally resampling from the posterior and shifting to a new hypothesis. The decision to resample and so shift to a new hypothesis is probabilistic. The learner resamples with a probability dependent on the degree to which the hypothesis is contradicted by data. As the probability of the sampled hypothesis given the data goes down, the learner is more likely to resample. This is similar to using a particle filter with just a single particle. The computationally expensive resampling step is more likely to be carried out as that particle becomes inconsistent with the data. This algorithm tends to maintain a hypothesis that makes a successful prediction and only tries a new hypothesis when the data weigh against the original choice. Thus, this algorithm may have some advantages over other forms of Monte Carlo search that resample after every observation. The central idea of maintaining a hypothesis provided it successfully accounts for the observed data leads us to refer to this strategy as the Win-Stay, Lose-Sample (WSLS) algorithm.

3.2. History of the WSLS principle

The WSLS principle has a long history in computer science and statistics, where it appears as a heuristic algorithm in reinforcement learning (Robbins, 1952). Robbins (1952) defined a heuristic

method for maximizing reward in a two-armed bandit problem based on this principle, in which an agent continues to perform an action provided that action is rewarded. These early studies employed very specific types of WSLs procedures, but in fact, the principle falls under a more general class of algorithms that depend on decisions from experience (see [Erev, Ert, Roth, et al., 2010](#), for a review).

The principle also has a long history in psychology. The basic principle of repeating behaviors that have positive consequences and not performing behaviors that have negative consequences can be traced back to [Thorndike \(1911\)](#). [Restle \(1962\)](#) did not explicitly call his model a WSLs algorithm, but he proposed a simple model of human concept learning that employed a similar principle, in which people were assumed to hold a single hypothesis in mind about the nature of a concept, only rejecting that hypothesis (and drawing a new one at random) when they received inconsistent evidence. This model, however, assumed that people could not remember the hypotheses they had tried previously or the pattern of data that they had observed. In fact people do use this information in concept learning, and Restle's version of the WSLs algorithm was subsequently shown to be a poor characterization of human concept learning ([Erikson, 1968](#); [Trabasso & Bower, 1966](#)).

A form of the WSLs strategy has also been shown to be used by children, most notably in probability learning experiments using variable reinforcement schedules ([Schusterman, 1963](#); [Weir, 1964](#)). [Schusterman \(1963\)](#) asked children to complete a two-alternative choice task in which they could obtain prizes by searching in one of two locations across many trials. Each location contained the prize on 50% of the trials, but in one condition, the probability that a reward would appear on a given side was 64% if the reward had occurred on that side in the previous trial, while in another condition, this probability was 39%. The three-year-olds very strongly followed a WSLs strategy, even in the condition where this strategy would not maximize the reward (i.e., the 39% condition). The five-year-olds showed a more sophisticated pattern, using a WSLs strategy more often in the 64% than in the 39% condition.

[Weir \(1964\)](#) also used a probability learning task to examine the development of hypothesis testing strategies throughout childhood. Over many trials, a child had to choose to turn one of three knobs on a panel to obtain a marble. Only one knob produced marbles and it did so probabilistically (e.g., on 66% of the trials). Results were analyzed in terms of a variety of strategies, including WSLs. In this task, the most common pattern followed by three- to five-year-old children was a win-stay, lose-stay strategy, (i.e., persevering with the initial hypothesis) which maximizes reward in this task. The 7- to 15-year-old children were more likely than the other age groups to show a WSLs strategy, but children in this age range also commonly used a win-shift, lose-shift (i.e., alternation) strategy.

A parallel form of WSLs has also been analyzed as a simple model of learning that leads to interesting strategies in game theory, where it is often used to describe behavior in deterministic binary decision tasks (see [Colman, Pulford, Omtzigt, & al-Nowaihi, 2010](#), for a more complete historical overview). For example, [Nowak and Sigmund \(1993\)](#) showed that WSLs outperforms tit-for-tat (TFT) strategies in extended evolutionary simulations of the Prisoner's Dilemma game when the simulations include these additional complexities. The use of WSLs, as opposed to TFT, results in much quicker recovery following errors and exploitation of unconditional cooperators – behaviors that correspond well with typical decisions of both humans and other animals.

3.3. Analyzing the WSLs algorithm

The WSLs principle provides an intuitive strategy for updating beliefs about hypotheses over time. This raises the question of whether there are specific WSLs strategies that can actually approximate the distribution over hypotheses that would be computed using full Bayesian inference. That is, are there specific variants of WSLs that, like independent sampling, converge to the posterior distribution as the number of samples increase? These kinds of WSLs strategies would have to be more sophisticated than the simple strategies we have described so far.

In exploring the question of whether an algorithm based on the WSLs principle can approximate Bayesian inference, we will also extend our analysis of WSLs to include cases where data are probabilistic. Previous work discussing WSLs strategies would only allow for deterministic shifts from one hypothesis to another. This involves a simple yes–no judgment about whether the data confirm or refute a hypothesis. Bayesian inference, in contrast, is useful because it allows a learner to weight

evidence probabilistically. Can an algorithm that uses the basic WSLs principle, but includes more sophisticated kinds of probabilistic computation also be used to approximate Bayesian inference?

Our WSLs algorithm assumes that learners maintain their current hypothesis provided they see data that are consistent with that hypothesis, and generate a new hypothesis otherwise. This is closest to the version explored by Restle (1962), but diverges from it in that instead of randomly choosing a hypothesis from an a priori *equally* weighted set of hypotheses (consistent with the data), a learner *samples* a hypothesis from the posterior distribution. That is, the important difference between these models is that in contrast to choosing a hypothesis at random from a stable, uniform distribution, in our model a learner chooses the next hypothesis with probability proportional to its posterior probability, which is continually being updated as new data are observed. The posterior distribution in deterministic cases is simply the prior distribution renormalized by the remaining consistent hypotheses.

It is relatively straightforward to show that this algorithm can approximate Bayesian inference in cases where the likelihood function $p(d_i|h)$ is deterministic, giving a probability of 1 or 0 to any observation d_i for every h , and observations are independent conditioned on hypotheses. More precisely, the marginal probability of selecting a hypothesis h_n given data, $d_1, \dots, d_n$, is the posterior probability $P(h|d_1, \dots, d_n)$, provided that the initial hypothesis is sampled from the prior distribution $P(h)$ and hypotheses are sampled from the posterior distribution whenever the learner chooses to shift hypotheses.² That is to say, like other approximation algorithms such as Independent Sampling, if you were to look at the distribution of responses of a group of learners who observed the same evidence and who were each following this WSLs strategy, the distribution of answers would approximate the posterior distribution predicted by Bayesian inference. This algorithm is thus a candidate for approximating Bayesian inference in settings similar to the deterministic causal induction problem that we explore in this paper. The proof is given in Appendix A and provides the first of two specific WSLs algorithms that we will investigate.

More interestingly, the WSLs algorithm can be extended to approximate Bayesian inference in stochastic settings by assuming that the probability of shifting hypotheses is determined by the probability of the observed data under the current hypothesis. A proof of the more general case, for stochastic settings, is provided in Appendix B. The proof provides a simple set of conditions under which the probability that a learner following the WSLs algorithm entertains a particular hypothesis matches the posterior probability of that hypothesis. Essentially, when data are more consistent with the current hypothesis we treat it as a “win”; this means a learner is more likely to stay with her hypothesis. When data are less consistent with the hypothesis, we treat it as a “loss”; this means the learner is more likely to sample from the updated posterior.

There are two interesting special cases of the class of WSLs algorithms identified by the conditions laid out in our proof. The first special case is a simple algorithm that makes a choice to resample from the posterior based on the likelihood associated with the current observation d_i for the current h , $p(d_i|h)$. With probability proportional to this likelihood, the learner maintains the current hypothesis; otherwise, she samples a new hypothesis from the posterior distribution. The second special case is the most efficient algorithm of this kind, in the sense that it minimizes the rate at which sampling from the posterior is required. Resampling is minimized by considering the likelihood for the current hypothesis relative to the likelihoods for all hypotheses when determining whether to resample – if the current hypothesis assigns highest probability to the current observation, then the learner will always maintain that hypothesis. The proof demonstrates that following one of these algorithms results in the same overall aggregate result as learners always sampling a response from the updated posterior, and thus approximates Bayesian inference. Our exploration of the WSLs algorithm in stochastic settings will focus on the first of these cases, where the choice to resample is based just on the likelihood for the current hypothesis, as this minimizes the burden on the learner.

² It is worth noting that the WSLs algorithm still requires the learner to generate samples from the posterior distribution. However, this algorithm is more efficient than Independent Sampling, since samples are only generated when the current hypothesis is rejected. Other inference algorithms, such as Markov chain Monte Carlo, could be used at this point to sample from the posterior distribution without having to enumerate all hypotheses – a possibility we return to later in the paper.

Our WSLS algorithm differs in an important way from other variants of WSLS strategies when used in a stochastic setting, in that it is possible that the learner ends up with the same hypothesis again after “switching” (because that hypothesis is sampled from the posterior). This contrasts with variants of the WSLS rule that require the learner to always shift to a new hypothesis (see [Colman et al., 2010](#)). While such an algorithm would behave the same as our algorithm in a deterministic setting if it was appropriately augmented to sample from the posterior on switching, it would behave differently in a stochastic setting. It is an interesting open question whether this (potentially more efficient) “must switch” algorithm can be generalized to approximate Bayesian inference in a stochastic setting.

Both Independent Sampling and WSLS can approximate the posterior distribution, but there is an important difference between them: WSLS introduces dependencies between responses and favors sticking with the current hypothesis. We can characterize this difference formally as follows. In IS there is no dependency between the hypotheses sampled after data points n and $n + 1$, h_n and h_{n+1} , but there is for WSLS: if the data are consistent with h_n , then the learner will retain h_n with probability proportional to $p(d|h_n)$ rather than randomly sampling h_{n+1} from the posterior distribution. We can use this difference to attempt to diagnose whether people use a WSLS type algorithm when they are solving a causal learning problem.

4. Evaluating inference strategies in children and adults

We now turn to the question of whether people’s actual responses are well captured by the algorithms described in the previous section, which we explore using mini-microgenetic studies with both adults and preschool children. Young children are a particularly important group to study for two reasons. First, their responses are unlikely to be influenced by specific education or explicit training in inference. Second, if these inductive procedures are in place in young children, they could contribute to the remarkable amount of causal learning that takes place in childhood.

In particular, we explore children’s responses in a causal learning task in which they must judge whether a particular artifact, such as a block, is likely to have particular effects, such as causing a machine or another block to light up. In earlier studies using very similar methods (e.g. “blicket detector” studies) researchers have found that preschool children’s inferences go beyond simple associative learning and have the distinctive profile of causal inferences. For example, children will use inferences about the causal relation of the block and machine to design novel interventions on the machine – patterns of action they have never actually observed – to construct counter-factual inferences and to make explicit causal judgments including judgments about unobserved hidden features of the objects (e.g. [Bonawitz, van Schindel, Friel, & Schulz, 2012](#); [Gopnik & Sobel, 2000](#); [Gopnik et al., 2004](#); [Schulz, Gopnik, & Glymour, 2007](#); [Sobel, Yoachim, Gopnik, Meltzoff, & Blumenthal, 2007](#)).

Although in many contexts young learners demonstrate sophisticated casual reasoning, there are numerous findings that children (and adults alike) have difficulty with explicit hypothesis testing (e.g., [Klahr, Fay, & Dunbar, 1993](#); [Kuhn, 1989](#)). Although young children clearly can sometimes articulate hypotheses explicitly, and can learn them from evidence, they have much more difficulty articulating or judging **how** hypotheses are justified by evidence. Our task is designed to explore children’s and adults’ reasoning about a causal system, but does not require the learner to be meta-cognitively aware of their own process of reasoning from the data.

In each experiment, we first determine whether the responses of both adults and children in these tasks are consistent with Bayesian inference at the computational level. Then we explore whether their behavior is consistent with particular learning algorithms, with special focus on the WSLS algorithm and Independent Sampling. In particular, we might expect that if participants behave in ways consistent with the WSLS algorithm we should observe dependencies between their responses. Experiment 1 begins with an empirical investigation of the deterministic causal learning scenario we described earlier. Experiments 2 and 3 examine how people make inferences in the stochastic scenario.

5. Experiment 1: Deterministic causal learning

5.1. Methods

5.1.1. Participants and design

Participants were 37 four- and five-year-olds (mean 56 months, range 52–71 months) recruited from a culturally diverse daycare and local science museum and 60 adult undergraduate volunteers recruited by flyers. Participants did not receive any compensation. Participants were randomly assigned to one of two conditions: 18 children and 30 adults participated in an *Evidence Type 1* condition; 19 children and 30 adults participated in an *Evidence Type 2* condition. One child in the *Evidence Type 2* condition was excluded from analysis due to technical malfunction of the stimuli during the evidence presentation. There were no differences in children's age between conditions ($t(34) = 0.09$, $p = 0.93$).

5.1.2. Stimuli

Stimuli were six (6 cm^3) cubic blocks. Three blocks were colored red and three colored yellow on all but one side. A small light switch was surreptitiously controlled at the back of the block, which caused the block to appear to glow through the uncovered side when activated.

5.1.3. Procedure

The procedure was identical for both children and adults except as noted below. Participants were first introduced to the blocks and told that sometimes blocks like this might light up and sometimes they might not, but it was the participant's job to help figure out how these particular red and yellow blocks work. Then participants were asked to predict what would happen if two red blocks were bumped together, "Will they both light or will they both not light?" Participants were given descriptive pictures as they generated their responses so that they could point to the pictures or respond verbally. This was done to provide children with an alternative if they were too shy to respond verbally. Participants were then asked what would happen if two yellow blocks were bumped together: "Will they both not light or will they both light?" Finally, participants were asked what would happen if a red and yellow block bumped into each other: "Will just the red light? Will just the yellow light? Will they both light? Or will they both not light?" The order of presentation of questions was counter-balanced across participants. These initial responses gave a measure of participants' prior beliefs about the blocks – their belief in the likely rule that governed how the blocks should behave, before observing any information about the pattern of lighting between blocks.

After they provided an initial response, participants were then shown one of two patterns of evidence. One group (in the *Evidence Type 1* condition) was shown that when a red and yellow block bumped together, just the yellow block activated. The other group (in the *Evidence Type 2* condition) was shown that when a red and yellow block bumped together, both blocks activated.³ Both groups of participants were then asked again what would happen if a red and red block bumped together and what would happen if a yellow and yellow block bumped together. If participants said "I don't know" they were prompted to make a guess.

5.2. Results

Children's responses were video-taped and coded online by a research assistant who observed the experiment either through a one-way mirror or sitting in the same room as the experimenter, facing the child. A second research assistant recoded responses. Responses were unique and unambiguous, and coder agreement was 100%; both coders were blind to hypotheses. During adult testing, the experimenter wrote out the responses during the testing session and later recorded them electronically. A research assistant, blind to study hypotheses, checked the data entry to confirm that there were no errors in transcription.

³ This evidence necessarily ruled out 12 of the 16 hypotheses, which were inconsistent with the observed evidence.

To assess which theory each learner held during each stage of the experiment, we looked at the causal judgments about each set of blocks interactions and selected the hypothesis consistent with participants' predictions. For example, if the participant responded that two red blocks would light each other, two yellow blocks would not light each other, and a red block would light a yellow block (but not vice versa) then the response was coded as theory 2; see Fig. 1.

5.2.1. Prior distribution

5.2.1.1. Children. Because the initial phase of the experiment was identical across conditions, we combined the children's responses in the initial phase to look at the overall prior distribution over hypotheses. Children had a slight bias for theories that included a greater number of positive, lighting relations. For both IS and WSLS models we estimated an empirical prior from this initial distribution produced by the children. The frequencies of each causal structure were tabulated from the children's responses. The relatively large number of causal structures to choose from compared to the number of participants resulted in a few structures that were never produced. To minimize bias from these smaller sample sizes, smoothing was performed by adding 1 "observation" to each theory before normalizing, consistent with using a uniform Dirichlet prior on the multinomial distribution over causal structures (e.g., Bishop, 2006). The smoothed estimated prior distribution over children's hypotheses is shown in Fig. 1(b).

5.2.1.2. Adults. We also combined responses from both conditions in order to investigate the overall prior distribution over adults' responses and tabulated frequencies adding 1 to all bins before normalizing. The estimated prior distribution over adult's hypotheses is shown in Fig. 1(c).

5.2.2. Aggregate distribution

We then considered how responses would change when the learner received new evidence. We computed the updated Bayesian posteriors for both conditions and for the children and the adults separately using the prior distributions above. Comparing the proportion of participants who endorsed each theory following the evidence to the Bayesian posterior distribution revealed high correlations (*Children*: $r = .92$; *Adults*: $r = .97$). So, at the computational level, the learners seemed to produce responses that are consistent with Bayesian inference in this task. In the *Evidence Type 1* condition there appears to be mild deviations between the model predictions and children's endorsement of Theory 10 and Theory 14; however, comparing the proportion of responses predicted by the Bayesian model for these theories to the proportion given by the children revealed no significant difference, $\chi^2(1, N = 30) = 1.2, p = .27$.

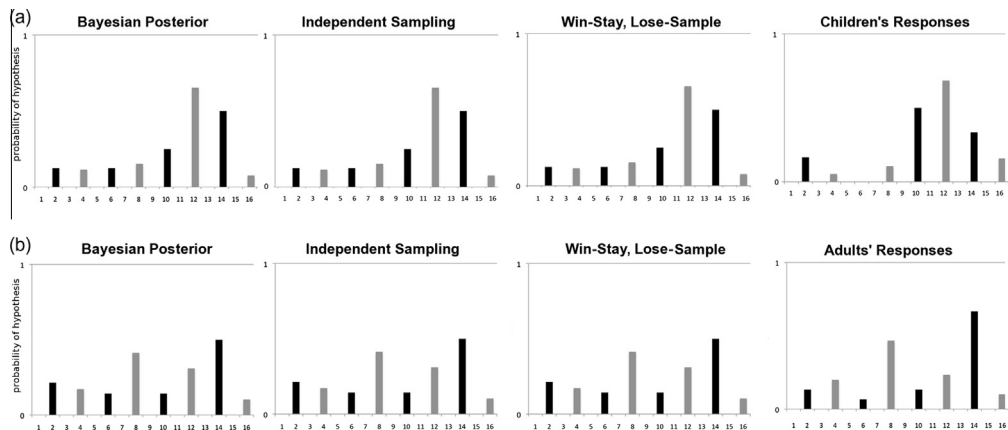


Fig. 2. Results of the Independent Sampling and Win-Stay, Lose-Sample algorithms, as compared to the Bayesian posterior and participant responses. (a) Predictions generated from children's priors and children's responses. (b) Predictions generated from adult's priors and adult's responses. Dark bars represent the *Evidence Type 1* and the lighter lines are *Evidence Type 2*.

We also depict example aggregate results for the WSLS and IS algorithm. As can be seen in Fig. 2, and consistent with the formal results provided as a proof in the appendices, both IS and WSLS produce identical distributions, which also approximate the ideal Bayesian posterior.

5.2.3. Dependency predictions

Given that both children and adult response patterns on aggregate reflect Bayesian posteriors, we can investigate which of the two approximation algorithms (WSLS or IS) best capture individuals' responses. In particular, we investigate whether there are dependencies between learners' responses indicative of the WSLS algorithm.

5.2.3.1. Children. Results suggest that children's responses following the evidence were strongly dependent on their initial responses: If children generated a data-consistent response initially, they were significantly more likely to give the same response again. Indeed, only 2 of the 15 children who initially provided a would-be, data-consistent response, shifted to a different response following the data. For comparison to the models, we generated a representative depiction of a WSLS and IS run (see Fig. 3(a) and (b)). For both models, we initialized first responses to the distribution of responses generated by the children in the Experiment (for maximal ease of visual comparison). Results show a typical run of each model following one of two data observations provided in the experiment (*Evidence Type 1* and *Type 2* conditions).

To more rigorously test whether children's pattern of responding is consistent with the models, we calculated the probability of switching hypotheses, given a first guess on the correct hypothesis, under the IS algorithm. Combining both groups (*Evidence Type 1* and *Type 2*) weighed by the proportion of children starting from correct hypotheses in each of these groups, the probability of switching under IS is $p = .42$. Note that only 2 of the 15 children who started with a correct hypothesis switched on their second query. This is significantly fewer than predicted by the IS algorithm (Binomial, $p < .05$). The pattern of children's responding is thus inconsistent with the IS algorithm, and suggests dependencies indicative of the WSLS algorithm. However, note that the deterministic WSLS algorithm predicts that *all* children should retain their hypothesis if the evidence provided is consistent; in our study, there were two children who observed consistent evidence and nonetheless changed responses on the second query. We come back to these findings in the Discussion below.

5.2.3.2. Adults. Adults also reflected strong dependencies between responses (Fig. 3(d)). Only 3 of the 20 adults who initially provided a response that then turned out to be consistent with the data, shifted to a different response afterward. This result is significantly fewer than predicted by the IS algorithm (IS predicted shifting = 62%; Binomial, $p < .0001$) and is thus inconsistent with the IS algorithm. It does reflect the dependencies between responses predicted by the WSLS algorithm. Thus, both children and adults tended to stick with their initial hypothesis when data were consistent more than would be predicted if sampling was independent.

5.3. Discussion

Our deterministic causal learning experiment demonstrates that learners' responses on aggregate return the posterior distributions predicted by Bayesian models. Importantly, these results show how tracking learning at the level of the individual can help us understand the specific algorithms that child and adult learners might be using to approximate Bayesian inference in deterministic cases. The data from Experiment 1 suggest that preschoolers and adults show dependencies between responses rather than independently sampling responses each time from the posterior. These results extend previous work exploring WSLS strategies. They show that at least one strategy of this kind provides a viable way to approximate Bayesian inference for a class of inductive problems, in particular problems of deterministic causal inference. They provide evidence that WSLS is a more appropriate algorithm for describing learners' inferences in these deterministic settings.

However, the non-zero probability of switching following theory-consistent evidence is not perfectly consistent with WSLS. Forgetting and other attentional factors may have led to some random noise in the participants' pattern of responding. A likely possibility is that learners may have believed

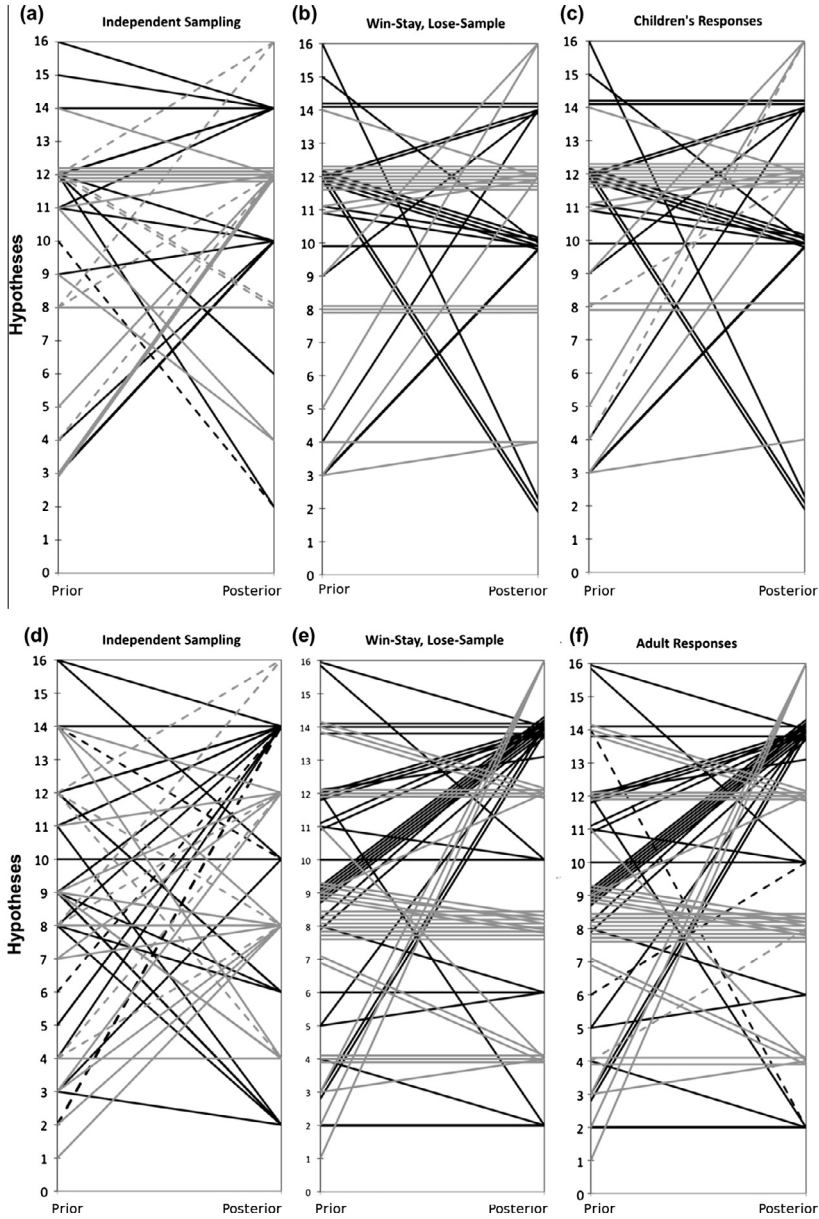


Fig. 3. Predictions of both models with priors generated from children's initial predictions, compared against human judgments (a) Independent Sampling, and (b) Win-Stay, Lose-Sample, and (c) children's responses. Models with priors generated from adults' initial predictions (d) Independent Sampling, and (e) Win-Stay, Lose-Sample, and (f) adults' responses. Dark lines represent the data observation of a red and yellow block interacting and resulting in just the yellow block lighting, leaving only hypotheses 2, 6, 10, and 14 consistent with the data. The lighter lines represent data observations of a red and yellow block interacting and resulting in both the yellow and red block lighting, leaving hypotheses 4, 8, 12, and 16 consistent with the data. For visual aid, dashed lines indicate cases where the initial response provided was consistent with the data, but the second response provided switched to a different hypothesis.

that the blocks lighting behavior was not perfectly deterministic. Thus, our WSL algorithm that permits stochastic data may better account for these results.

Because deviations from deterministic responding were rare, the results of this experiment do not provide a strong basis against which to evaluate the stochastic version of our WSL algorithm. Consequently, in Experiments 2 and 3 we consider cases where the causal models were intentionally designed to be stochastic, and examine how this affects the strategies that children and adults adopt. We also increase the number of evidence and response trials in order to support more sophisticated analyses.

6. Experiment 2: Stochastic causal learning

Experiment 1 investigated a special case of WSL in deterministic settings; however, a potentially more natural case of causal learning involves causes that are stochastic. In our earlier analyses, we presented a special case of WSL in stochastic scenarios where the learner chooses whether to resample a hypothesis based on the likelihood of the data they have just observed. To investigate whether this algorithm captures human behavior, we designed an experiment with causally ambiguous, but probabilistically informative evidence. We asked learners to generate predictions as they observed each new piece of evidence and then compared learners' pattern of responses to our models.

6.1. Methods

6.1.1. Participants and design

Participants were 40 preschoolers (mean 58 months, range 48–70 months) recruited from a culturally diverse daycare and local science museum and 65 undergraduates recruited from an introductory psychology course. The participants were split into two conditions ($N(\text{children}) = 20$, $N(\text{adults}) = 28$ in the *Active then inactive* condition; $N(\text{children}) = 20$, $N(\text{adults}) = 32$ in the *Inactive then active* condition). An additional five adult participants were excluded for not completing the experiment and four children were excluded for failing the comprehension check (see Procedure). There were no differences in children's age between conditions ($t(38) = 0.01$, $p = 0.99$).

6.1.2. Stimuli

Stimuli consisted of 13 white cubic blocks (1 cm^3). Twelve blocks had custom-fit sleeves made from construction paper of different colors: four red, four green, and four blue. An "activator bin" large enough for 1 block sat on top of a [$15'' \times 18.25'' \times 14''$] box. Attached to this box was a helicopter toy that lit up. The machine was activated by the experimenter, who surreptitiously pressed a hidden button in the box as she placed a block in the bin. This led to a strong impression that the blocks had actually caused the effect. There was a set of "On" cards that pictorially represented the toy in the on position, and a set of "Off" cards that pictorially represented the toy in the off position. Because adult participants were tested in large groups, a computer slideshow that depicted the color of the blocks and the cards was used to provide the evidence for them.

6.1.3. Procedure

6.1.3.1. Children. Participants were told that different blocks possess different amounts of "blicketness," a fictitious property. Blocks that possess the most blicketness almost always activate the machine, blocks with very little blicketness almost never activate the machine, and blocks with medium blicketness activate the machine half of the time. A red block was chosen at random and placed in the activator bin. The helicopter toy either turned on or remained in the off position. A corresponding On or Off card was placed on the table to depict the event. The cards remained on the table throughout the experiment. After five repetitions using the same red block for a total of six demonstrations, participants were told that red blocks have the most blicketness (they activated 5/6 times). The same procedure was repeated for the blue and green blocks. The blue blocks had very little blicketness (activating the toy 1/6 times), and the green blocks had medium blicketness (activating 3/6 times). Children were asked to remind the experimenter which block activated the machine the most, least,

and a medium amount to ensure that children understood the procedure and remembered the probabilistic properties of the blocks. Data from children who were unable to provide correct answers to this comprehension check were eliminated.

After the comprehension check, a novel white block that lost its sleeve was presented and children were asked what color sleeve the white block should have (red, green, or blue). This provided a measure of participants' initial beliefs about the intended color of the novel block before observing any evidence about whether or not the block activates the machine – the prior probability that a block of a particular color would be sampled. Children were told they would be asked about the block a few times. The white block was then placed into the bin four times and each time the participant saw whether or not the toy activated. Following each demonstration, the appropriate On or Off card was chosen and participants were asked to provide a guess about the right color for the block. Before each guess the participants were told, “It’s okay if you keep thinking it is the same color and it is also okay if you change your mind.”

We designed the trials to include a greater opportunity for belief revision to most effectively test our models. In the *Active then inactive* condition the toy turned on for the first trial and then did not activate on the three subsequent trials. In the *Inactive then active* condition the toy did not activate on the first trial, but turned on for the three subsequent trials. (See Fig. 4.) Children’s sessions were video recorded to ensure accurate transcription of responses.

6.1.3.2. Adults. The procedure for adults was identical to that for the children with the following exceptions. Participants in each condition were tested on separate days in two large groups. Participants were instructed to record responses using paper and pen and not to change answers once they had written them down. Adults were also shown a computer slide show so they could more easily see the cards depicting the evidence, and the slideshow remained on the screen throughout the experiment. To ensure that participants were paying attention, they were asked to match each color to the proper degree of blcketness (most, very little, medium). As with the children, adults were introduced to the novel white block and were asked to record their best guess as to the color of the block before they observed any demonstrations and after each demonstration.

6.2. Results

We evaluated models of participants’ behavior using three criteria. First, we assessed whether the aggregate responses of children and adults were consistent with Bayesian inference. Both the WSLS and IS algorithm approximate Bayesian inference in aggregate, so this does not discriminate between these accounts, but it does potentially rule out alternative models. Second, we looked at participants’ trial-by-trial data to compute the probabilities with which they switched hypotheses. This allowed us to compare the WSLS and IS algorithms, which make different predictions about switch probabilities.

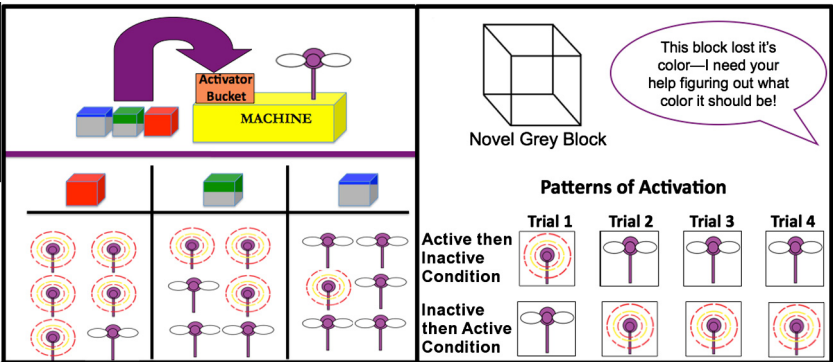


Fig. 4. Depiction of the procedure used in Experiments 2 and 3.

Third, we looked at the probability of the trial-by-trial choices that participants made under the WSLs and IS algorithms. This produces log-likelihood scores for each model, which can be compared directly.

6.2.1. Comparison to Bayesian inference

6.2.1.1. Children. Responses were uniquely and unambiguously categorized as “red”, “green”, and “blue”. A coder blind to hypotheses and conditions reliability coded 50% of the video clips; agreement was 100%. We determined the parameters for the children’s prior distribution and likelihood in two ways, based on initial responses or based on maximizing the fit to the Bayesian posterior. For the first way (“initial responses”) priors were determined by the participants’ predictions about the color of the novel block, prior to seeing any activations. The likelihood of block activation was determined by the initial observations of block activations during the demonstration phase (5/6 red, 1/2 green, 1/6 blue). Children’s initial guesses before seeing the first demonstration (i.e. “on” or “off”) reflected a slight bias favoring the red blocks (50%), with blue (30%) and green (20%) blocks being less favored. That is, prior to observing information about the novel block, children seemed to favor the guess that the block sampled was the one that would activate the machine most often. For the second way (“maximized”) we searched for the set of priors and the likelihood activation weights that would maximize the log-likelihood for the model.⁴

We compared the proportion of children endorsing each block at each trial of the experiment to the Bayesian posterior probability. Using either set of parameters, children’s responses were well captured by the posterior probability (initial responses: $r(22) = .77$, $p < .0001$; maximized: $r(22) = .86$, $p < .0001$, see Fig. 5(a–d)). In fact, because of the general noise in the data as reflected in the relatively high standard deviations for small samples and categorical responding, these fits between the model and the data are actually about as high as we could expect. Indeed, the correlations between the models and the data are not significantly different from those obtained by finding the correlation of a random sample of half of the participant responses compared to the other half ($r(22) = .84$; Fisher r -to- z transformation, initial: $z(24) = -0.65$, $p = ns$; maximized: $z(24) = 0.23$, $p = ns$).

6.2.1.2. Adults. Adult responses were also uniquely and unambiguously categorized as “red”, “green,” and “blue”. There was a slight bias to favor green blocks (60%), with red (25%) and blue (15%) blocks being less favored.⁵ As with the children’s data, we determined the parameters for the prior distribution and likelihood by using initial responses or using the maximized parameters.⁶

We then compared the proportion of adults endorsing each color block at each trial to the Bayesian posterior probability. Using either set of parameters, adult participant responses were well captured by the posterior probability (initial responses: $r(22) = .76$, $p < .001$; maximized: $r(22) = .85$, $p < .0001$, see Fig. 5(e–h)).

Both child and adult responses on aggregate, then, were well captured by the predictions of Bayesian inference. The primary difference between the model and data is that adults converged more quickly to a block choice than predicted by the model. This may be a consequence of pedagogical reasoning, a point we return to in the General discussion. However, given that children and adult responses on aggregate are well captured by the posterior probability predicted by Bayesian inference, we now turn to the question of what approximation algorithm best captures individual responding, and to whether responses showed the distinctive dependency patterns of the WSLs algorithm.

6.2.2. Comparison to WSLs and IS

To compare people’s responses to the WSLs and IS algorithms, we first calculated the “switch” probabilities under each model in the two ways we described previously: using the parameters from

⁴ The maximized priors for children were .38 red, .22 green, .4 blue; these priors correspond to the priors represented by child participants. The maximized likelihood was .73 red, .5 green, .27 blue, which also corresponds to the likelihood given by the initial activation observations.

⁵ Such a bias is consistent with people’s interest in non-determinism; the green blocks were the most stochastic in that they activated on exactly half the trials.

⁶ The maximized priors for adults were .27 red, .48 green, .25 blue; these priors correspond strongly to the priors represented by participants. The maximized likelihood was .85 red, .5 green, .16 blue, which also corresponds strongly to the likelihood given by the initial activation observations.

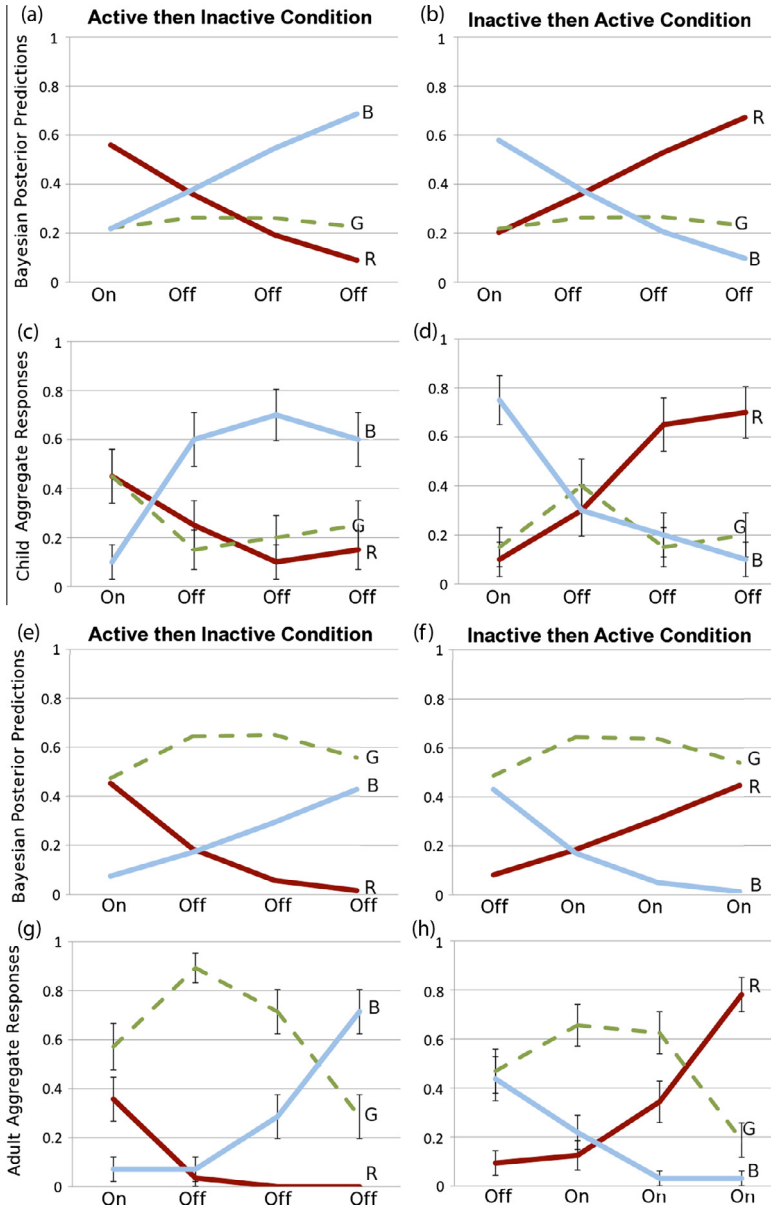


Fig. 5. Bayesian posterior probability and human data from Experiment 2 for each block, red (R), green (B), and blue (B) after observing each new instance of evidence, using parameters estimated from fitting the Bayesian model to the data. WSLs and IS produce the same aggregate result reported here. (For interpretation of the references to color in this figure legend, the reader is referred to the web version of this article.)

the initial responses and using the previously estimated maximized parameters. Calculating switch probabilities for IS is relatively easy: because each sample is independently drawn from the posterior, the switch probability is simply calculated from the posterior probability of each hypothesis after observing each piece of evidence.

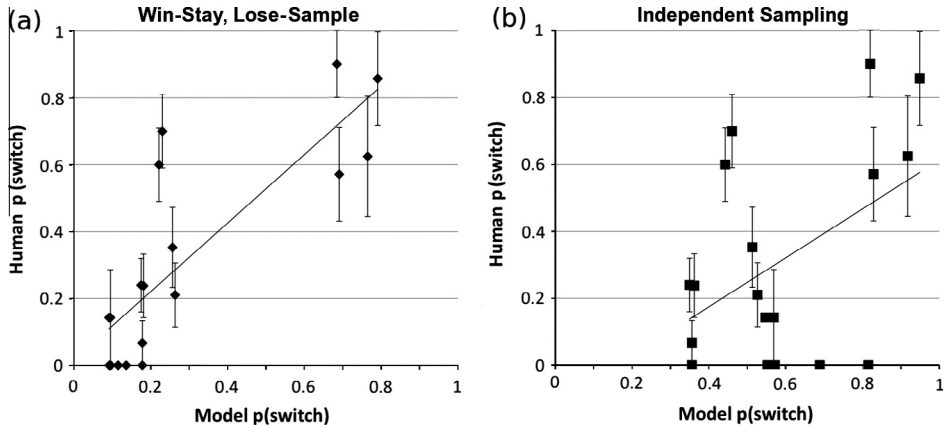


Fig. 6. Correlations between the probability of switching hypotheses in the models given the maximized parameters and the adult data in Experiment 2, for (a) the Win-Stay Lose-Sample algorithm and (b) Independent Sampling.

Recall that in the stochastic case of WSLS, participants should retain hypotheses that are consistent with the evidence, and resample proportional to the likelihood, $p(d|h)$. Thus, switch probabilities for WSLS were calculated such that resampling is based only on the likelihood associated with the current observation, given the current h . That is, with probability equal to this likelihood, the learner resamples from the full posterior, which is computed using all the data observed up until that point. For example, if a participant guessed “red” initially and then observes that the novel block does not activate the toy on trial 1 (e.g. an outcome that would occur 1/6 times given it was actually a red block), then under WSLS she should stay with the guess “red” with probability 1/6 and resample a guess from an updated posterior that takes into account all the data observed so far with probability 5/6.

6.2.2.1. Children. We computed the proportion of children “switching” given each color block at each trial and compared it to the predicted switch probabilities of each model. Children’s responses were equally well captured by the WSLS and IS model when comparisons were made using the maximized parameters (WSLS: $r(22) = .61$, $p < .01$; IS: $r(22) = .61$, $p < .01$; Fisher r -to- z transformation, $z(22) = 0$, $p = ns$) or using parameters given by the initial responses (WSLS: $r(22) = .58$, $p < .01$; IS: $r(22) = .58$, $p < .01$; Fisher r -to- z transformation, $z(22) = 0$, $p = ns$). We computed the log-likelihood scores for both models. The IS model better fit the child data than the WSLS model (initial responses: log-likelihood for WSLS = -229 , log-likelihood for IS = -205 ; maximized: log-likelihood for WSLS = -217 , log-likelihood for IS = -197). These log-likelihood scores can also be compared using Bayes factors, which in this case correspond to the difference in the two log-likelihoods. Following guidelines proposed by Kass and Raftery (1995) the Bayes factors revealed “very strong” evidence suggesting that the IS model better fit the data than the WSLS model, using either the initial responses or the maximized parameters. However, taken together, these results suggest that while the pattern of dependencies between children’s responses is captured marginally better by the IS algorithm, neither model provided a particularly strong fit.

6.2.2.2. Adults. We computed the proportion of adults “switching” given each color block at each trial and compared it to the predicted switch probabilities of each model. Adult responses were much better captured by the WSLS algorithm using the maximized parameters ($r(15) = .81$, $p < .0001$)⁷

⁷ In 5 of the 24 possible condition by color by trial cells, there were fewer than 2 data points possible (e.g. none or only 1 adult had generated a “red” response previously and so switching score from red on that trial to the next could only be 0, or 0 and 1 respectively.) These data cells were dropped from adult correlation scores.

and the parameters given by participant initial responses ($r(15) = .78, p < .001$) as compared to the IS algorithm (maximized: $r(15) = .58, p = .02$; initial responses: $r(15) = .39, p = ns$), although given the small sample size, these correlation coefficients were not significantly different from each other. Scatterplots are shown in Fig. 6. We also computed the log-likelihood scores for both models. The WSLs model better fit the adult data than the IS model (initial responses: log-likelihood for WSLs = -221 , log-likelihood for IS = -262 ; maximized: log-likelihood for WSLs = -215 , log-likelihood for IS = -251). Comparing Bayes Factors based on these log-likelihoods revealed “very strong” evidence in favor of the WSLs model for both sets of parameters (Kass & Raftery, 1995). These results suggest that the pattern of dependencies between adults’ responses is better captured by the WSLs algorithm than by an algorithm such as IS that produces independent samples.

6.3. Discussion

Adult responses in Experiment 2 and in Experiment 1 were best captured by the WSLs algorithm. In Experiment 2, the adults’ pattern of responses were highly correlated with predictions from WSLs (whether using initial responses to set parameters or using maximized values based on correlating parameters to the Bayesian posterior). Both correlation and log-likelihood scores were greater for the WSLs model as compared to the IS model. This suggests that the adult predictions are consistent with the dependencies predicted by the WSLs algorithm. Moreover, there was a close fit to the WSLs model itself.

Although the overall pattern of children’s responding was consistent with Bayesian inference, children’s responses in Experiment 2 were not well correlated with either the WSLs or IS model, and further analyses of the likelihood scores under each model revealed a marginally better fit to the IS model. These results stand in conflict with the results from Experiment 1, which reflected strong dependencies in children’s predictions that are characteristic of the WSLs algorithm.

Why might children have shown less dependency and greater switching between responses in Experiment 2? One major difference between the two experiments is that in Experiment 1, children were not asked about the same set of interactions after having observed new data; in contrast, children were asked about the same block on repeated trials in Experiment 2. Such questioning may have led children to believe that the experimenter was challenging their initial responses and thus may have led to greater switching, which would not be consistent with WSLs. Indeed, recent research suggests that repeated questioning from an experimenter gives children strong pragmatic cues that initial responses were incorrect, leading the child to switch responses, even if they had been confident in their initial prediction (Gonzalez, Shafto, Bonawitz, & Gopnik, 2012). In Experiment 3, we investigate this possibility with a minor modification to the procedure used in Experiment 2.

7. Experiment 3: Stochastic causal learning when exchanging testers

In Experiment 3, we investigate the possibility that the greater than predicted amount of switching by children in Experiment 2 was caused by repeated questioning from the same experimenter. To control for this, we replicated the procedure in Experiment 2 with one minor modification: instead of having the same experimenter ask children about their beliefs after each new observation of data, we had a new experimenter come in for each subsequent trial. As suggested by Gonzalez et al. (2012), the ignorance of each new experimenter to the children’s previous guess should effectively break any possible pragmatic assumptions that may cause children to believe that repeated questioning is an indication that their initial response was incorrect.

7.1. Methods

7.1.1. Participants and design

Participants were 40 preschoolers (mean 58 months, range 48–71 months) recruited from a culturally diverse daycare and local science museum. The participants were split into two conditions (*Active then inactive*: $N = 20$; *Inactive then active*: $N = 20$). An additional eight participants were excluded for

one of three reasons: failing the comprehension check (3), non-compliance with the procedure – due to being uncomfortable with interacting with such a large number of experimenters (4), and experimenter error (1). There were no differences in children's ages between conditions ($t(38) = 0.10$, $p = 0.92$) or across the three experiments ($F(2, 113) = 0.84$, $p = 0.434$).

7.1.2. Stimuli

Stimuli were identical to those used in Experiment 2.

7.1.3. Procedure

The procedure was identical to Experiment 2 with the following exceptions, which occurred after Experimenter 1 obtained prior hypotheses from each child. The experimenter said to the child, “Well, I have to go write something down, but my friend [Experimenter 2's name] is going to come in and take a turn playing with you and these toys. She has never seen this block before and she does not know what color coat it should have so you'll get to help her figure it out.” Experimenter 1 retrieved Experimenter 2 from outside the room and Experimenter 2 said, “I really like to play with blocks. Wow, it's a block without a coat! Should I put it in this machine to see what happens?” Just as in Experiment 2, the machine either activated or did not activate (depending on the condition). The experimenter and child then chose a card to depict the event and the experimenter asked the child what color they thought the block was. Experimenter 2 told the child that another friend who had never seen the block before was going to come and have a turn. Experimenters 3 through 5 followed the same script as Experimenter 2. Many of the children immediately offered a guess to the new experimenter when she entered the room. When this happened the experimenter said, “Okay, you think it should have a [red/green/blue] coat. Well, let's see what happens when I take a turn. It's okay if you keep thinking it should have a [red/green/blue] coat and it's okay if you change your mind and think it's a [red/green/blue] coat.”

7.2. Results

7.2.1. Comparison to Bayesian inference

Responses were uniquely and unambiguously categorized as “red”, “green”, or “blue”. A coder blind to hypotheses and conditions reliability coded 50% of the clips; agreement was 100%. As with Experiment 2, we determined the parameters for the children's prior distribution and likelihood in two ways: values given by initial responses and the maximized priors.⁸ As in Experiment 2, children's initial guesses (before seeing the first demonstration) reflected a slight bias favoring the red blocks (50%), with blue (20%) and green (30%) blocks being less favored. Comparing the proportion of children endorsing each color block at each trial to the Bayesian posterior probability, using either set of parameters, revealed strong correlations (initial responses: $r(22) = .67$, $p < .001$; maximized: $r(22) = .87$, $p < .0001$). These correlations are not significantly different from those obtained in Experiment 2 (Fisher r -to- z transformation, initial: $z(24) = -0.6$, $p = ns$; maximized: $z(24) = 0.25$, $p = ns$).

These results replicate Experiment 2; children's responses on aggregate were well captured by the predictions of Bayesian inference. We turn to the primary question of interest: do children show the dependencies we would predict on the WSLS model?

7.2.2. Comparison to WSLS and IS

To compare children's responses to the WSLS and IS algorithms, we calculated the “switch” probabilities using both sets of parameters (initial, maximized). We computed the proportion of children “switching” given each color block at each trial and compared it to the predicted switch probabilities of each model. Children's responses were better captured by the WSLS algorithm than by the

⁸ The maximized priors for children were .46 red, .21 green, .33 blue; these priors correspond to the priors represented by child participants. The maximized likelihood was .69 red, .5 green, .31 blue, which also corresponds to the likelihood given by the initial activation observations.

IS model when comparisons were made using the maximized parameters (WSLS: $r(14) = .79, p < .001$; IS: $r(14) = .67, p < .01$)⁹ and were also better captured by WSLS when using parameters given by the initial responses (WSLS: $r(14) = .78, p < .001$; IS: $r(14) = .70, p < .01$), although given the small sample size, these correlation coefficients were not significantly different from each other. We computed the log-likelihood scores for both models. The WSLS model fit the data better than the IS model (initial responses: log-likelihood for WSLS = -185 , log-likelihood for IS = -223 ; maximized: log-likelihood for WSLS = -172 , log-likelihood for IS = -200 ; Bayes Factors revealed “very strong” evidence in favor of WSLS for both sets of parameters) and also better than the log-likelihood scores of the children’s responses in Experiment 2. Children showed a better fit to the WSLS model in this experiment than in Experiment 2. In sum, correlations were higher using maximized parameters and the parameters given by initial responses, and the log-likelihood scores in this experiment were also higher than the log-likelihood scores for WSLS in Experiment 2.

7.3. Discussion

As with Experiments 1 and 2, children’s responses on aggregate corresponded strongly to Bayesian inference. However, in this experiment unlike in Experiment 2, the correlation and log-likelihood scores were greater for the WSLS model as compared to the IS model. Moreover, the correlations to the model were high overall, near .8. This suggests that the dependencies in children’s predictions are consistent with the WSLS algorithm when other possible methodological, pragmatic concerns are removed. In turn this suggests that the WSLS algorithm may be an appropriate way to capture dependencies in children’s responses, but also suggests that such dependencies are easily influenced by social information. Children are extremely sensitive to possible cues provided by experimenters and will use those cues when they generate predictions.

8. Consideration of alternative models

One primary goal of this paper is to connect algorithmic and computational level accounts of learning, offering an analysis of the conditions under which WSLS strategies approximate Bayesian inference. This analysis reveals that there are special forms of WSLS strategies that produce behavior that on aggregate approximates Bayesian inference in deterministic and stochastic settings. However, it is worth considering the predictions that alternative models of learning and inference make on these tasks. Here we consider a few alternatives and discuss how these approaches relate to our model.

8.1. Associative models

One might wonder how associative accounts relate to the models proposed here. A traditional associative model, the Rescorla–Wagner model (Rescorla & Wagner, 1972), has a direct Bayesian generalization known as the Kalman filter (Dayan & Kakade, 2001). The main difference between these models and the Bayesian models we consider here is that associative models capture how one might update estimates of the strength of a causal relationship, assuming that causes combine to influence their effects in a linear fashion, while the models that we have focused on infer causal structure (i.e., whether or not causal relationships exist at all, or which class of causal relationship is being exhibited).

We could imagine a learner using an associative updating mechanism for the experiments presented here, considering the degree to which the activation of the novel block is consistent with the behavior of the Red, Green, and Blue block. Such a model, however, must then indicate how a learner chooses among these possibilities. Associative models are known to be able to exhibit probability matching behavior, but this is typically just matching the frequency with which different actions are rewarded (e.g. see Jones & Liverant, 1960; Myers, 1976; Neimark & Shuford, 1959; Vulkan, 2000;

⁹ In 8 of the 24 possible condition by color by trial cells, there were fewer than three data points possible; these data cells were dropped from correlation scores.

though see Denison et al., 2013). The phenomenon that needs to be accounted for in our data is probability matching to the posterior distribution indicated by the Bayesian model, which does not relate in a simple way to a past history of contingency or reward. Even if the model predicted probability matching to the posterior, this would simply implement the Independent Sampling algorithm against which we compared our WSLs algorithm.

8.2. Reducing cognitive load

In Appendix B, we define a second special case of WSLs that approximates Bayesian inference but minimizes the rate at which sampling from the posterior is required. In our paper, we focus on the first case (where the likelihood given the current data is considered) because it does not require the learner to consider the likelihoods of all other hypotheses relative to the currently held hypothesis. However, another way to reduce cognitive load is simply to not consider whether to stay or shift after seeing each piece of data. Instead a learner could let a small amount of data amass over N trials, considering the result of those last N trials when choosing whether to stay or shift.

Under the simple algorithm that has been our focus in this paper, increasing the number of trials used in the decision to stay or shift would exponentially increase the probability that the learner would resample (following the same logic as when the outcome of ten fair coin flips is significantly smaller than a single flip of the coin). If a learner (in effect) “skipped” the resampling step occasionally, and hypotheses were chosen only after each sequence of N trials, using the data from those N trials, this model would approximate Bayesian inference in aggregate – it is just WSLs with a different way of characterizing the “data” observed on each trial. However, the higher rate of switching would mean that the predictions of this WSLs model would be closer to the IS model, and thus more difficult to contrast.

More sophisticated WSLs algorithms, such as those in the second class of algorithms we discuss in Appendix B, could make it possible to reduce the frequency of deciding whether to stay or shift without increasing the rate of resampling, and are worth exploring in future work. As our primary goal was to compare WSLs with IS, we chose to solicit participant predictions after every trial to maximize differences between the models. Querying the participants had the effect of “forcing” them to consider hypotheses after each trial, which might have pre-empted the natural dynamics of hypothesis updating. One might be interested in whether learners actually do consider their hypotheses after each new observation of data, in the absence of querying. Of course, the only way to assess the models is by recording a measure of the learner’s current beliefs; it is not immediately obvious how one might retrieve a learner’s current belief state without “triggering” the sample step. Future work may investigate this separate, but interesting question of whether, and how often, learners choose to reconsider their beliefs when independently exploring in the absence of external queries.

8.3. Memory

8.3.1. The role of memory in the WSLs model

One might be concerned that the WSLs principle is too simple to support intelligent inferences in that it appears there is no “memory” in the system. However there is a sense in which our WSLs algorithm does have memory. It may appear that a learner who sees a hundred trials consistent with a single hypothesis and a learner who observes only one such trial will be equally likely on the subsequent trial to abandon their hypothesis if the data are inconsistent. In fact, however, our model simply predicts that the two learners will be equally likely to *resample* their hypotheses. Since these samples are drawn from the posterior, given all evidence observed so far, this provides a form of memory. The first learner, who has observed a hundred consistent trials, is resampling from a posterior that is weighted very strongly in favor of the hypothesis in consideration, whereas the second learner – who has observed just a single trial – will be sampling from a more uniform distribution. As a result, the first learner is more likely to *resample the hypothesis he was just considering*, while the second learner is more likely to sample an alternate hypothesis. This means that the first is more likely to maintain her hypothesis than the second learner.

8.3.2. *Inhibited memory in WSLS models*

A WSLS model with limited memory might use the same metric for deciding whether to stay or shift, but would instead sample from a posterior distribution based only on the last K observations. This kind of model would be appealing from an algorithmic perspective, because the learner would only need to keep track of the last few observations in updating their posterior. Although this model does not approximate a Bayesian solution for small values of K , it converges to an approximation of Bayesian inference as K becomes large.

We compared this kind of model to our participant data in a few ways. Because there are only four evidence trials in our Experiment 2, and taking $K = 4$ is thus the same as computing the full Bayesian posterior, we computed correlations between this model and participant results for $K = \{1, 2, 3\}$ for our three comparison criteria. First, we implemented this model and computed correlations to the overall aggregate response pattern of adults (Experiment 2) and children (Experiment 3). Second, we computed the switch probabilities for all three of the K model variants. Third, we computed the log likelihood scores for all three of the K models. For both adults and children, we found correlations that were not significantly different across any of these three measures (see [Appendix C, Tables 1a and 1b](#)).

8.3.3. *Empirical investigation of the inhibited memory WSLS models*

Although we did not find significant differences between the inhibited memory WSLS model and our full memory WSLS model in explaining the results of our experiments, the inhibited model seems intuitively like a poor characterization of human learning and performance because it does not account for cases in which data accumulate over numerous trials, resulting in strongly stable beliefs. Our original experiments were not designed to test for the role of mounting evidence because we were interested in cases in which the WSLS and IS models would make very different predictions (which can only be tested at the individual level, as both algorithms converge to Bayesian predictions on aggregate). In the interest of comparing WSLS to these alternative memory models, we ran a simple supplementary experiment. In this experiment, participants first receive mounting evidence for the “Red” block; the machine activates every time the novel test block is placed on the machine. Then there are two evidence trials in which the machine does not activate; if the participant believes the block is “Red” going into the trials, the failed activation is surprising, but not unexplainable.

After encountering trials in which the block does not activate the machine, a WSLS model sampling from the updated distribution predicts continued endorsement for the “Red” block. This is because the initial data weigh the posterior so strongly in favor of “Red”, that a few surprising trials do not significantly change this bias. Thus, “Red” is likely to be resampled when the surprising trial is considered. However, the inhibited memory WSLS model does not maintain this biased posterior. Thus, it will show a sudden, drastic shift following these few anomalous trials, with most of the data redistributing to a Blue or Green block.

We collected data from 20 adult participants over Amazon Mechanical Turk, following an online version of the same probabilistic block paradigm used in Experiment 2. After learning about the machine and the probability of activation for each block color, participants were introduced to the novel block and asked to make a guess about the block’s color. The participants then observed eight trials in which the machine activates, two trials in which the machine does not activate, and finally followed by two trials in which the machine activates again. After each trial, participants are asked about the color of the novel block.

We determined the priors on the red, green, and blue blocks by looking at the distribution of responses for each on the first trial (prior to observing any evidence). Our experiment tested relatively few participants to get a reasonable distribution on the prior; to get a close approximation to the ideal priors, we also included data from an additional 140 participants, who had completed an unrelated study, but for whom the procedure was identical up to the point of the prior solicitation. As with our data from Experiment 2, the online participants in this experiment had a prior preference for green blocks (75% respondents), and weaker preferences for red (17%) and blue (8%) block guesses. On aggregate, participants quickly converged to guessing “Red” following the presentation of repeated activation. By the fourth trial greater than 90% of participants believed the block was red. Importantly, participants continued to produce “Red” responses at near ceiling (>90%) for the remaining 8 trials, even after observing the 9th and 10th trial in which the machine did not activate.

Our first criterion for model comparison is whether the models capture the aggregate data. The Bayesian model captured these aggregate results ($r = .91$, $p < .0001$). We also implemented the inhibited memory model at $K = 2$ (as this was the best performing model given the original data). In contrast to the Bayesian model, the inhibited memory model failed to capture the aggregate results; in fact, there was no relationship between the predictions of the inhibited memory model and the data ($r = .05$, $p = ns$). The Bayesian model significantly outperformed the inhibited memory model, $z(36) = 6$, $p < .0001$. Because the memory model failed on the first criterion – capturing the aggregate data, follow-up analyses of the trial-by-trial switching and log-likelihood scores would also indicate significant outperformance of the WSLs model over the inhibited memory model.

8.4. Simplified WSLs algorithms

The key difference between our WSLs algorithm and previous WSLs algorithms that have been considered in the literature (e.g., [Restle, 1962](#)) is that in our algorithm the distribution from which new hypotheses are sampled is updated over time, given the aggregate data from past experience. This is what allows the algorithm to approximate Bayesian inference. We can also consider what happens to the WSLs algorithm when we adjust the distribution that samples are drawn from. The simplest two alternative models either draw samples from a uniform distribution or from the original prior distribution. Such a model is actually a version of the inhibited memory model, where $K = 0$ and the posterior remains constant. Importantly, this simplified WSLs model diverges from our model, because it does not redraw from a continually updated posterior distribution (thus inhibiting it from approximating Bayesian inference in aggregate). These simpler models instead spend more time (relative to other hypotheses) on hypotheses that are more consistent with the data.

Our experiment was not designed to test for the role of mounting evidence, critical for the relative stability of strongly held beliefs, so implementing this model did not reveal significantly different correlations than our WSLs model obtained (see [Appendix C](#)). However, as illustrated by our additional experiment that induces a strong bias for the red block, such a model seems intuitively like a poor characterization of human learning and performance. Indeed, for the supplementary experiment described in [Section 8.3](#), neither the Sample-from-Uniform ($r = -.03$) nor the Sample-from-Prior ($r = .72$) performed as well as the Bayesian model at capturing the aggregate data (Uniform vs. Bayes: $Z(36) = 6.33$, $p < .0001$; Prior vs. Bayes: $Z(36) = 2.52$, $p = .01$). These results demonstrate significant superiority of our WSLs algorithm over these simpler alternatives.

8.5. Model summary

We designed causal learning tasks that that we believed would likely meet the first criteria for investigating algorithmic difference – that responses on aggregate will approximate Bayesian inference. Our causal learning paradigm was inspired by the relatively recent approaches that solicit intuitive causal judgments and that tend to yield Bayesian consistent responses by children and adults (e.g., [Bonawitz & Lombrozo, 2012](#); [Denison et al., 2013](#); [Griffiths et al., 2011](#); [Sobel, Tenenbaum, & Gopnik, 2004](#)). These tasks are likely more natural for participants than the traditional concept learning, probability matching, and decision making tasks which might require meta-cognitive and pragmatic approaches to solve. Thus, task differences may affect the learning strategies employed, explaining the differences between previous studies supporting alternate forms of the WSLs strategy and our task.

Our primary goal in this paper is to develop cognitively plausible algorithms that efficiently approximate optimal Bayesian predictions. Nonetheless, it is important to understand how additional cognitive limitations such as memory constraints, response noise, and other psychological processes might be incorporated into algorithms that approximate Bayesian inference. A detailed discussion of this approach is outside the scope of this paper, but we recommend [Griffiths, Vul, and Sanborn \(2012\)](#) and [Bonawitz, Gopnik et al. \(2012\)](#) as further reading on connecting the algorithmic and computational levels of cognition.

9. General discussion

Our results show that a mini-microgenetic method can help us understand the specific algorithms that learners might be using to approximate Bayesian inference. First we introduced an algorithm, Win-Stay, Lose-Sample, based on the Win-Stay, Lose-Shift principle, that approximates Bayesian inference by maintaining a single hypothesis over time. We proved that the marginal distribution over hypotheses after observing data will always be the same for this algorithm as for sampling hypotheses independently from the posterior (an algorithm that we refer to as Independent Sampling). That is, both algorithms return a distribution over responses consistent with the posterior distribution obtained from Bayesian inference. Our analysis also made clear that there are important differences in what WSLS and IS predict for the dependency between guesses, making it possible to distinguish these algorithms empirically.

We explored whether people actually use algorithms like these to solve inductive inference problems in two simple causal learning tasks. Our experiments were designed with two sets of analyses in mind. First in our experiments, both children's and adults' overall responses are consistent with Bayesian inference. Second, the algorithmic level analysis revealed that the dependencies between an individual's responses are characteristic of the WSLS algorithm. Participants did not independently sample responses each time from the posterior. These results extend previous work exploring WSLS strategies. They show that at least one strategy of this kind provides a viable way to approximate Bayesian inference. They also provide evidence that WSLS is an appropriate algorithm for describing people's inferences in both deterministic and stochastic causal learning tasks.

In what follows, we turn to a broader discussion of the idea of connecting probabilistic models of cognition with simple cognitive processes by exploring algorithms for approximating Bayesian inference. We then discuss how these approaches can inform developmental psychology and vice versa. We make some proposals about how these algorithms are affected by sampling assumptions, such as pedagogy. Finally, we look ahead to see how future work can inform these questions.

9.1. Connecting algorithms to Bayesian inference

Probabilistic models of cognition have become increasingly prevalent and powerful, but they are limited. Literally following the procedures of Bayesian inference by enumerating and testing each possible hypothesis is computationally costly, and so could not be the actual algorithm that learners use. Instead learners must use some other algorithm that approximates ideal Bayesian inference. Considering the algorithmic level of analysis more seriously can help to address these significant challenges for Bayesian models of cognition.

In fact, some classic empirically generated psychological process models turn out to correspond to the application of Monte Carlo methods, which approximate ideal Bayesian inference. For example, [Shi et al. \(2008\)](#) showed that importance sampling corresponds to exemplar models, a traditional process-level model that has been applied in a variety of domains. [Sanborn et al., \(2010\)](#) used particle filters to approximate rational statistical inferences for categorization. [Bonawitz and Griffiths \(2010\)](#) show that importance sampling can be used as a framework for analyzing the contributions of generating and then evaluating hypotheses.

The work we have presented here provides a new contribution toward this literature on "rational process models". While most of the previous work employing these models has focused on showing how such models can reproduce and explain existing effects, here we have set up an experimental design with the explicit goal of discerning which algorithm learners might be using. As a result, this novel approach helps us to empirically explore the explicit algorithms in human cognition and understand how they connect to computational level explanations.

9.2. Algorithms and development

The computational level has provided an important perspective on children's behavior, affording interesting and testable qualitative and quantitative predictions that have been borne out empirically.

But we have suggested that this is just the starting point for exploring learning in early childhood. By considering other levels of analysis we can help to address significant challenges for Bayesian models of cognitive development. In particular we can begin to address the problem of how a learner might search through a (potentially infinite) space of hypotheses. This problem might be particularly challenging for young children who, in at least some respects, have more restricted memory and information-processing capacities than adults (German & Nichols, 2003; Gerstadt et al., 1994). WSLS is one algorithm that provides a plausible account of belief revision in early childhood and a practical starting point for addressing the concerns raised by Bayesian models of development.

For example, both WSLS and IS algorithms can be seen as extreme versions of particle filters: Independent Sampling is similar to the case where there is a large set of particles drawn from the posterior and one member of the set is drawn at random for each query; WSLS is similar to using a single particle that is resampled from the posterior when the particle becomes inconsistent with the data. There may be some value in exploring algorithms that lie between these extremes, with a more moderate number of particles. Future work could examine the degree to which fewer or greater numbers of particles capture inference and to what degree these constraints change with age and experience. In particular, developmental changes in cognitive capacity might correspond to changes in the number of particles, with consequences that are empirically testable as changing the number of particles will result in different patterns of dependency.

9.3. Pedagogical reasoning in the causal learning tasks

In any Bayesian model, the learner must make assumptions about how data were generated. These assumptions can lead to different predictions, but are not at odds with the WSLS model. For example, the likelihood term which accounts for the switching probability in WSLS can be influenced by assumptions about how the data are sampled – such as in the context of teaching, which can lead to faster convergence on the correct concept.

In Experiment 2, the aggregate distribution of adult responses shifted more dramatically than predicted by the Bayesian model we presented. It is likely, given the context of showing participants a predetermined computer slideshow, that adults were making a pedagogical assumption (Shafto & Goodman, 2008), which would better capture the data. Pragmatically, adults might have wondered, “Why is this experimenter showing me a predetermined slideshow if they don’t already have a goal in mind – to teach me a specific concept.” Additionally, recruiting adult participants from a classroom setting may have pedagogically-primed the participants. Also consistent with the pedagogical framework, the act of showing an initially “misleading” trial may have led the participants to believe that they were being intentionally deceived, which would have led to more rapid convergence to the alternate block.

Children in Experiment 2 showed a greater than predicted propensity to switch responses between guesses. This is consistent with another consequence of pedagogy: the perceived intention of the experimenter can itself be a form of evidence. For example, previous research shows that repeated questioning from a knowledgeable experimenter provides strong cues that the child’s previous guess was incorrect and should be changed (Gonzalez et al., 2012). In Experiment 3, we tested the hypothesis that this may have been a consequence of a pedagogical assumption made by the children. Children might wonder: “Why is this adult asking me the same question over and over if I am giving a reasonable answer? After all, it is her toy.” When this assumption was diminished in Experiment 3, children showed less shifting and their pattern of responding was captured by WSLS. Future work may develop additional WSLS algorithms that take into account these assumptions by the learner about the data.

9.4. Future work

Connecting the computational and algorithmic levels is a significant challenge for Bayesian models of cognition, and this is only a first step in understanding the psychological processes at work in causal inference. We believe that there are several important directions for future research in this area.

First, it would be interesting to test the WSLS algorithm's predictions across various psychological experiments that have relied purely on a Bayesian inference approach. This would allow for a better assessment of the WSLS algorithm's efficiency over a wider range of tasks and would provide a way to investigate the potential scope of its use by human learners. In particular, given previous arguments for use of the WSLS principle in concept learning (Restle, 1962), this may be a place to start a broader investigation. Other work that has found dependencies in human response patterns (e.g., Vul and Pashler, 2008) suggests that strategies that successively evaluate variants on a hypothesis might be relevant in other domains as well.

In addition to considering a wider range of tasks, future work should explore a wider range of algorithms. For example, it will be interesting to explore algorithms that shift from one hypothesis to the next by modifying the current hypothesis in a principled and structured manner such that the learner need not 'track' the posterior distribution during learning. Gibbs sampling (e.g. Robert & Casella, 2004) provides one example of a strategy of this kind that has been explored in the past (Sanborn et al., 2010), but there may be other algorithms in this arena that are worth investigating and attempting to relate to human behavior.

We constrained our space to a modest number of hypotheses, but other work has begun to examine how hypothesis spaces may be learned and simultaneously searched (Bonawitz, Gopnik et al., 2012; Katz, Gooman, Kersting, Kemp, & Tenenbaum, 2008; Kemp, Goodman, & Tenenbaum, 2011; Kemp, Tenenbaum, Niyogi, & Griffiths, 2010; Ullman et al., 2010). This enterprise should be jointly developed with approaches taken here that explore the space of plausible algorithms that capture people's causal inferences.

An individual learner may also employ different strategies in different contexts. For example, changes in development, changes in the complexity of the task, and even changes in temperament may influence which algorithm is selected. For example, in ongoing work, we are looking at whether causal inferences are affected by the learner's current temperament. That is, we put a learner in positive or negative mood, perform the mini-microgenetic method, and compare responses to various algorithms to see whether the best fitting algorithm changes with induced mood. Explaining why and understanding how these different factors influence algorithms remains a challenge for future work.

10. Conclusions

We have demonstrated that a probabilistic version of the WSLS stratagem can be used to construct an algorithm that approximates Bayesian inference. It provides a way to perform sequential Bayesian inference while maintaining only a single hypothesis at a time and leads to an efficient approximation scheme that might be useful in computer science and statistics. We have also shown that a WSLS algorithm seems to capture adult and child judgments in two simple causal learning tasks. Our results add to the growing literature suggesting that even responses by an individual that may appear non-optimal may in fact represent an approximation to a rational process.

Acknowledgments

Thanks to Annie Chen, Sophie Bridgers, Alvin Chan, Swe Tun, Madeline Hansen, Tiffany Tsai, Jan Iyer, Christy Tadros for help with data collection and coding. This research was supported by the McDonnell Foundation Causal Learning Collaborative, grant IIS-0845410 from the National Science Foundation, and grant FA-9550-10-1-0232 from the Air Force Office of Scientific Research.

Appendix A, B, & C. Supplementary material

Supplementary data associated with this article can be found, in the online version, at <http://dx.doi.org/10.1016/j.cogpsych.2014.06.003>.

References

- Bishop, C. M. (2006). *Pattern recognition and machine learning*. New York: Springer.
- Bonawitz, E., & Griffiths, T. L. (2010). Deconfounding hypothesis generation and evaluation in Bayesian models. In *Proceedings of the 32nd annual conference of the cognitive science society*.
- Bonawitz, E., Gopnik, A., Denison, S., & Griffiths, T. L. (2012). Rational randomness: The role of sampling in an algorithmic account of preschooler's causal learning. In J. B. Benson (Serial Ed.) & F. Xu & T. Kushnir (Vol. Eds.), *Advances in child development and behavior: Rational constructivism in cognitive development* (pp. 161–192). Waltham, MA: Academic Press.
- Bonawitz, E., Denison, S., Griffiths, T., & Gopnik, A. (2014). Probabilistic models, learning algorithms, response variability: Sampling in cognitive development. *Trends in Cognitive Science*. <http://dx.doi.org/10.1016/j.tics.2014.06.006>.
- Bonawitz, E., & Lombrozo, T. (2012). Occam's rattle: Children's use of simplicity and probability to constrain inference. *Developmental Psychology*, 48(4), 1156.
- Bonawitz, E. B., Shafto, P., Gweon, H., Goodman, N. D., Spelke, E., & Schulz, L. E. (2011). The double-edged sword of pedagogy: Teaching limits children's spontaneous exploration and discovery. *Cognition*, 120(3), 322–330.
- Bonawitz, E., van Schindel, T., Friel, D., & Schulz, L. (2012). Children balance theories and evidence in exploration, explanation, and learning. *Cognitive Psychology*, 64(4), 215–234.
- Bowers, J. S., & Davis, C. J. (2012). Bayesian just-so stories in psychology and neuroscience. *Psychonomic Bulletin*, 138(3), 389–414.
- Bullock, M., Gelman, R., & Baillargeon, R. (1982). The development of causal reasoning. In W. J. Friedman (Ed.), *The developmental psychology of time* (pp. 209–254). New York: Academic Press.
- Carey, S. (1985). *Conceptual change in childhood*. Cambridge, MA: MIT Press.
- Chater, N., Goodman, N., Griffiths, T. L., Kemp, C., Oaksford, M., & Tenenbaum, J. B. (2011). The imaginary fundamentalists: The unshocking truth about Bayesian cognitive science. *Behavioral and Brain Sciences*, 34(4), 194–196.
- Colman, A. M., Pulford, B. D., Omtzigt, D., & al-Nowaihi, A. (2010). Learning to cooperate without awareness in multiplayer minimal social situations. *Cognitive Psychology*, 61, 201–227.
- Dayan, P., & Kakade, S. (2001). Explaining away in weight space. In T. K. Leen, T. G. Dietterich, & V. Tresp (Eds.), *Advances in neural information processing systems* (vol. 13, pp. 451–457). MIT Press.
- Denison, S., Bonawitz, E. B., Gopnik, A., & Griffiths, T. L. (2013). Rational variability in children's causal inferences: The sampling hypothesis. *Cognition*, 126(2), 285–300.
- Doucet, A., de Freitas, N., & Gordon, N. J. (Eds.). (2001). *Sequential Monte Carlo methods in practice*. Berlin: Springer-Verlag.
- Erev, I., Ert, E., Roth, A. E., et al (2010). A choice prediction competition: Choices from experience and from description. *Journal of Behavioral Decision Making*, 23, 15–47.
- Erikson, J. R. (1968). Hypothesis sampling in concept identification. *Journal of Experimental Child Psychology*, 76, 12–18.
- German, T. P., & Nichols, S. (2003). Children's inferences about long and short causal chains. *Developmental Science*, 6, 514–523.
- Gerstadt, C. L., Hong, Y. J., & Diamond, A. (1994). The relationship between cognition and action: performance of children 3–7 years old on a stroop-like day-night test. *Cognition*, 53, 129–153.
- Gigerenzer, G., & Gaissmaier, W. (2011). Heuristic decision making. *Annual Review of Psychology*, 62, 451–482.
- Gigerenzer, G., & Goldstein, D. G. (1996). Reasoning the fast and frugal way: Models of bounded rationality. *Psychological Review*, 103, 650–669.
- Gilks, W. R., Richardson, S., & Spiegelhalter, D. J. (1996). *Markov Chain Monte Carlo in practice*. Boca Raton, FL: Chapman and Hall/CRC.
- Gonzalez, A., Shafto, P., Bonawitz, E. B. & Gopnik, A. (2012). Are you sure? The effects of neutral queries on children's choices. In *Proceedings of the 34th annual conference of the Cognitive Science Society*.
- Goodman, N. D., Tenenbaum, J. B., Feldman, J., & Griffiths, T. L. (2008). A rational analysis of rule-based concept learning. *Cognitive Science*, 32, 108–154.
- Gopnik, A. (2012). Scientific thinking in young children. Theoretical advances, empirical research and policy implications. *Science*, 337, 1623–1627.
- Gopnik, A., Glymour, C., Sobel, D., Schulz, L., Kushnir, T., & Danks, D. (2004). A theory of causal learning in children: Causal maps and Bayes nets. *Psychological Review*, 111, 1–31.
- Gopnik, A., & Meltzoff, A. (1997). *Words, thoughts, and theories*. Cambridge, MA: MIT Press.
- Gopnik, A., & Schulz, L. (2004). Mechanisms of theory formation in young children. *Trends in Cognitive Science*, 8, 371–377.
- Gopnik, A., & Schulz, L. (Eds.). (2007). *Causal learning: Psychology, philosophy, and computation*. Oxford: Oxford University Press.
- Gopnik, A., & Sobel, D. (2000). Detectingblickets: How young children use information about novel causal powers in categorization and induction. *Child Development*, 71, 1205–1222.
- Gopnik, A., & Tenenbaum, J. B. (2007). Bayesian networks, Bayesian learning, and cognitive development. *Developmental Science*, 10(3), 281–287.
- Gopnik, A., & Wellman, H. M. (2012). Reconstructing constructivism: Causal models, Bayesian learning mechanisms, and the theory theory. *Psychological Bulletin*, 138(6), 1085–1108.
- Griffiths, T. L., Chater, N., Kemp, C., Perfors, A., & Tenenbaum, J. (2010). Probabilistic models of cognition: Exploring representations and inductive biases. *Trends in Cognitive Sciences*, 14(8), 357–364.
- Griffiths, T. L., Chater, N., Norris, D., & Pouget, A. (2012). How the Bayesians got their beliefs (and what those beliefs actually are): Comment on Bowers and Davis (2012). *Psychological Bulletin*, 138(3), 415–422.
- Griffiths, T. L., Sobel, D., Tenenbaum, J. B., & Gopnik, A. (2011). Bayes and Blickets: Effects of knowledge on causal induction in children and adults. *Cognitive Science*, 35(8), 1407–1455.
- Griffiths, T. L., & Tenenbaum, J. B. (2005). Structure and strength in causal induction. *Cognitive psychology*, 51(4), 334–384.
- Griffiths, T. L., & Tenenbaum, J. B. (2009). Theory-based causal induction. *Psychological Review*, 116, 661–716.
- Griffiths, T. L., Vul, E., & Sanborn, A. N. (2012). Bridging levels of analysis for probabilistic models of cognition. *Current Directions in Psychological Science*, 21(4), 263–268.
- Gweon, H., Schulz, L., & Tenenbaum, J. (2010). Infants consider both the sample and the sampling process in inductive generalization. *Proceedings of the National Academy of Science*, 107(20), 9066–9071.

- Jones, M. H., & Liverant, S. (1960). Effects of age differences on choice behavior. *Child Development*, 31(4), 673–680.
- Jones, M., & Love, B. C. (2011). Bayesian fundamentalism or Enlightenment? On the explanatory status and theoretical contributions of Bayesian models of cognition. *Behavioral and Brain Sciences*, 34, 169–231.
- Kass, R. E., & Raftery, A. E. (1995). Bayes factors. *Journal of the American Statistical Association*, 90(430), 773–795.
- Katz, Y., Gooman, N. D., Kersting, K., Kemp, C., & Tenenbaum, J. B. (2008). Modeling semantic cognition as logical dimensionality reduction. In *Proceedings of the thirtieth annual conference of the cognitive sciences society*.
- Kemp, C., Goodman, N. D., & Tenenbaum, J. B. (2011). Learning to learn causal models. *Cognitive Science*, 1–59.
- Kemp, C., Tenenbaum, J. B., Niyogi, S., & Griffiths, T. L. (2010). A probabilistic model of theory formation. *Cognition*, 114(2), 165–196.
- Klahr, D., Fay, A., & Dunbar, K. (1993). Heuristics for scientific experimentation: A developmental study. *Cognitive Psychology*, 25, 111–146.
- Kuhn, D. (1989). Children and adults as intuitive scientists. *Psychological Review*, 96, 674–689.
- Kushnir, T., & Gopnik, A. (2007). Conditional probability versus spatial contiguity in causal learning: Preschoolers use new contingency evidence to overcome prior spatial assumptions. *Developmental Psychology*, 44, 186–196.
- Levine, M. (1975). *A cognitive theory of learning: Research on hypothesis testing*. Hillsdale, NJ: Lawrence Erlbaum.
- Levy, R., Reali, F., Griffiths, T. L., 2009. Modeling the effects of memory on human online sentence processing with particle filters. In: D. Koller, D. Schuurmans, Y. Bengio, L. Bottou (Eds.), *Advances in neural information processing systems* (Vol. 21, pp. 937–944).
- Marcus, G. F., & Davis, E. (2013). How robust are probabilistic models of higher-level cognition? *Psychological Science*, 24(12), 2351–2360.
- Marr, D. (1982). *Vision*. Freeman Publishers.
- McClelland, J. L., Botvinick, M. M., Noelle, D. C., Plaut, D. C., Rogers, T. T., Seidenberg, M. S., et al (2010). Letting structure emerge: connectionist and dynamical systems approaches to cognition. *Trends in Cognitive Sciences*, 14(8), 348–356.
- Myers, J. L. (1976). Probability learning and sequence learning. In W. K. Estes (Ed.), *Handbook of learning and cognitive processes: Approaches to human learning and motivation* (pp. 171–205). Hillsdale, NJ: Erlbaum.
- Neal, R. M. (1993). Probabilistic inference using Markov chain Monte Carlo methods. Technical Report CRG-TR-93-1, Dept. of Computer Science, University of Toronto.
- Neimark, E. D., & Shuford, E. H. (1959). Comparison of predictions and estimates in a probability learning situation. *Journal of Experimental Psychology*, 57(5), 294–298.
- Nowak, M., & Sigmund, K. (1993). A strategy of win-stay, lose-shift that outperforms tit-or-tat in the prisoner's dilemma game. *Nature*, 364, 56–58.
- Pearl, J. (2000). *Causality: Models, reasoning and inference*. Cambridge, UK: Cambridge University Press.
- Rescorla, R. A., & Wagner, A. R. (1972). A theory of Pavlovian conditioning: Variations in the effectiveness of reinforcement and nonreinforcement. In A. H. Black & W. F. Prokasy (Eds.), *Classical conditioning II: Current research and theory* (pp. 64–99). New York: Appleton-Century-Crofts.
- Restle, F. (1962). The selection of strategies in cue learning. *Psychological Review*, 69, 329–343.
- Robbins, H. (1952). Some aspects of the sequential design of experiments. *Bulletin of the American Mathematical Society*, 58, 527–535.
- Robert, C., & Casella, G. (2004). *Monte Carlo statistical methods* (2nd ed.). New York: Springer.
- Russell, S. J., & Norvig, P. (2003). *Artificial intelligence: A modern approach* (2nd ed.). Upper Saddle River, NJ: Prentice Hall.
- Sanborn, A. N., Griffiths, T. L., & Navarro, D. J. (2010). Rational approximations to rational models: Alternative algorithms for category learning. *Psychological Review*, 117(4), 1144–1167.
- Schulz, L. E., Bonawitz, E. B., & Griffiths, T. L. (2007). Can being scared make your tummy ache? Naive theories, ambiguous evidence and preschoolers' causal inferences. *Developmental Psychology*, 43, 1124–1139.
- Schulz, L., Gopnik, A., & Glymour, C. (2007). Preschool children learn about causal structure from conditional interventions. *Developmental Science*, 10(3), 322–332.
- Schulz, L. E., Hoopell, K., & Jenkins, A. (2008). Judicious imitation: Young children imitate deterministic actions exactly, stochastic actions more variably. *Child Development*, 79(20), 395–410.
- Schulz, L. E., & Sommerville, J. (2006). God does not play dice: Causal determinism and children's inferences about unobserved causes. *Child Development*, 77(2), 427–442.
- Schusterman, R. J. (1963). The use of strategies in two-choice behavior of children and chimpanzees. *Journal of Comparative and Physiological Psychology*, 56(1), 96–100.
- Seiver, T., Gopnik, A., & Goodman, N. D. (2012). Did she jump because she was the big sister or because the trampoline was safe? Causal inference and the development of social attribution. *Child Development*, 84(2), 443–454.
- Shafra, P., & Goodman, N. (2008). Teaching games: Statistical sampling assumptions for pedagogical situations. In *Proceedings of the 30th annual conference of the cognitive science society*.
- Shi, L., Feldman, N. H., & Griffiths, T. L. (2008). Performing Bayesian inference with exemplar models. In *Proceedings of the 30th annual conference of the cognitive science society*.
- Siegler, R. S. (1975). Defining the locus of developmental differences in children's causal reasoning. *Journal of Experimental Child Psychology*, 20, 512–525.
- Siegler, R. S. (1996). *Emerging minds: The process of change in children's thinking*. New York: Oxford University Press.
- Siegler, R. S., & Chen, Z. (1998). Developmental differences in rule learning: A microgenetic analysis. *Cognitive Psychology*, 36(3), 273–310.
- Simon, H. (1955). A behavioural model of rational choice. *Quarterly Journal of Economics*, 69, 99–118.
- Simon, H. (1957). *Models of man*. New York: Wiley.
- Sobel, D. M., Tenenbaum, J. B., & Gopnik, A. (2004). Children's causal inferences from indirect evidence: Backwards blocking and Bayesian reasoning in preschoolers. *Cognitive Science*, 28(3), 303–333.
- Sobel, D., Yoachim, C., Gopnik, A., Meltzoff, A., & Blumenthal, E. (2007). The blicket within: Preschoolers' inferences about insides and causes. *Journal of Cognition and Development*, 8(2), 159–182.
- Spirtes, P., Glymour, C., & Scheines, R. (2000). *Causation, prediction and search* (2nd ed). New York, NY: MIT Press.

- Tenenbaum, J. B., Griffiths, T. L., & Kemp, C. (2006). Theory-based Bayesian models of inductive learning and reasoning. *Trends in Cognitive Science*, 10, 309–318.
- Thorndike, E. L. (1911). *Animal intelligence: Experimental studies*. New York: Macmillan.
- Trabasso, T., & Bower, G. H. (1966). Presolution dimensional shifts in concept identification: A test of the sampling with replacement axiom in all-or-none models. *Journal of Mathematical Psychology*, 3, 163–173.
- Ullman, T., Goodman, N., & Tenenbaum, J. B. (2010). Theory acquisition as stochastic search. In S. Ohlsson & R. Catrambone (Eds.), *Proceedings of the 32nd annual conference of the cognitive science society* (pp. 2840–2845). Austin, TX: Cognitive Science Society.
- Vul, E., Goodman, N. D., Griffiths, T., & Tenenbaum, J. B. (2009). One and done? Optimal decisions from very few samples. In *Proceedings of the 31st annual conference of the cognitive science society*.
- Vul, E., & Pashler, H. (2008). Measuring the crowd within probabilistic representations within individuals. *Psychological Science*, 19(7), 645–647.
- Vulkan, N. (2000). An economist's perspective on probability matching. *Journal of Economic Surveys*, 14(1), 101–118.
- Weir, M. W. (1964). Developmental changes in problem solving strategies. *Psychological Review*, 71(6), 473–490.
- Wellman, H. M. (1990). *The child's theory of mind*. Cambridge, MA: MIT Press.
- Wellman, H. M., & Gelman, S. A. (1992). Cognitive development: Foundational theories of core domains. *Annual Review of Psychology*, 43, 337–375.
- Xu, F., & Kushnir, T. (2012). Rational constructivism in cognitive development. *Advances in child development and behavior* (Vol. 43). Waltham, MA: Academic Press.
- Xu, F., & Tenenbaum, J. B. (2007). Word learning as Bayesian inference. *Psychological Review*, 114(20), 245.