

Learning from others: an experimental test of Brownian motion uncertainty models

Journal of Theoretical Politics

2015, Vol. 27(4) 588–612

©The Author(s) 2014

Reprints and permissions:

sagepub.co.uk/journalsPermissions.nav

DOI: 10.1177/0951629814559723

jtp.sagepub.com



David Glick

Boston University, Boston, MA, USA

C Daniel Myers

University of Minnesota, Minneapolis, MN, USA

Abstract

Models of decision-making with outcome uncertainty are common in political science and related fields. Recent work in flagship journals has challenged canonical work by modeling outcome uncertainty as Brownian motion. This theoretical innovation has resonated because it is highly tractable and captures intuitively important realities of many decisions in ways that earlier models cannot. As theoretically attractive as the new models are, they have not yet been evaluated empirically. This is especially important because Brownian motion models place actors in more cognitively demanding situations than previous models. We offer what we believe to be the first experimental test of actors' ability to behave in ways consistent with the Brownian motion model by evaluating subjects' ability to rationally learn from another actor's experiences. We show that subjects adjust their actions in response to the Brownian motion uncertainty. However, they deviate from optimal behavior in important ways, particularly in more complex situations.

Keywords

Brownian Motion; learning; policy diffusion; policy experimentation

Corresponding author:

C Daniel Myers, University of Minnesota, 1414 Social Sciences, Minneapolis, MN 55455, USA.

Email: cdmyers@umn.edu

I. Introduction

Models of decision-making with outcome uncertainty are central to theoretical work in political science and a number of related fields. Situations in which actors cannot perfectly know how their choices produce outcomes are common. They motivate a variety of important questions and answers which span the divide between behavioral and institutional research. For example, theories of delegation to expert committees (e.g. Gilligan and Krehbiel, 1989), learning from others' policy experiments (e.g. Volden, 2006), and forming policy preferences with limited political knowledge (e.g. Berelson et al., 1954; Huckfeldt and Sprague, 1995) all comprise outcome uncertainty. Theoretical work in these areas and beyond is inseparable from questions about how one should conceive of choice–outcome uncertainty and what assumptions one is making about actors' cognitive capacities and decision environments.

Recent work in flagship journals in political science and economics challenges conventional uncertainty assumptions and the canonical models based on them. It does so by developing a new foundation for policy-uncertainty models by utilizing the mathematics of Brownian motion (Callander, 2008, 2011a,b; Glick, 2012). The new uncertainty model has resonated as a challenge to canonical work in large part because it is highly tractable and captures intuitively important realities of many political decisions in ways that earlier models cannot. As theoretically attractive as these models are, they have not yet been evaluated empirically. This next step is especially important because while not dramatically different for an analyst to implement, the Brownian motion model requires actors to incorporate distance and uncertainty into their choices and implies discarding seemingly relevant information in some situations.

In this paper we offer what we believe to be the first experimental test of actors' ability to behave in ways consistent with Brownian motion models. The motivation for the experiment derives from the motivation underlying the theoretical work. As Callander's (2008) article explains, the canonical uncertainty representation is limited because it assumes that the relationship between policies and outcomes is constant across policies and can be represented by a constant linear shock. One result of this assumption is that knowing one policy–outcome pair allows an actor to predict all policy–outcome pairs. In contrast, the Brownian motion innovation allows one to model a situation in which it is reasonable to assume that actors know how moving policy from A to B will move outcomes on average while allowing for new policies to produce unexpected outcomes.

We test this uncertainty representation by evaluating how well subjects learn from observing another actor's policy choice and outcome prior to making their own policy choice, in a situation where the relationships between policies and outcomes are described by a Brownian motion process. As described by Glick (2012), this situation produces a counter-intuitive finding in which some actors are better off mimicking a known policy whose outcome is known, even though they know that this outcome is some distance from their ideal point. Our experiment tests how closely individuals' decisions about (1) when to *mimic* another's policy, (2) when to

modify it, and (3) how much to modify it by, accord with predictions derived from a Brownian motion model (Glick, 2012). The difficulty of these decisions is a function of two factors: the *similarity* between the other's policy outcome and one's own policy goal, and the *complexity* of the policy area. We present subjects with the outcome of one policy decision and then ask them to make their own decision, and vary the similarity and complexity in a 2×2 design. The experiment tests individuals' ability to learn optimally in a richer and more complicated information environment than most existing work assumes.

This empirical test most nearly addresses a different substantive situation than those Callander has applied the basic model to. Specifically, the most direct substantive motivation for our experiment concerns policy diffusion in which governments learn from others' policies. Our work speaks indirectly to questions such as how well can a small state such as New Hampshire learn from a policy that has been implemented in a large state such as California? The diffusion setting is one example of a situation in which different policy uncertainty models with different assumptions imply different optimal behavior when incorporating the information that the first adopter's experiment implies for the second's choice. Thus, our findings contribute both to the literature concerned with uncertainty modeling and behavior, and to the literature concerned with policy diffusion and learning from others. Both the model and our test of it bridge formal work assuming utility maximization and behavioral work that is skeptical of human capacity to act optimally.

This paper proceeds in the following manner. We first review the literature concerning models of policy–outcome uncertainty, learning from others, and experimental tests of decision-making in these settings. We then elaborate on Brownian motion uncertainty and delineate the predictions about 'mimicking' and 'modifying' that we test. Next, we describe our experimental investigation and report the results. We show that subjects incorporate the unique features of Brownian motion uncertainty into their decision-making, but that they also deviate from optimal behavior in important ways, particularly when faced with more complex decisions.

2. Brownian motion and learning

Modeling decisions when there is uncertainty about the relationship between policies and the outcomes they produce requires specifying a 'policy mapping function' by which policies (choices) produce outcomes (Callander, 2008). A common and influential mapping function in continuous space (e.g. Gilligan and Krehbiel, 1987, 1989) is a constant linear shock (denoted by ω in this work) drawn from a known distribution. In this canonical model, famously applied to delegation and committee expertise in Congress, the outcome (y) that a policy (x) produces is $x + \omega$. Importantly, the uncertainty (ω) is originally unknown but constant. This means that the outcome of policy (x) is that policy plus the constant random shock ($x + \omega$) and the outcome of a different policy ($x + n$) is simply the new policy plus the same (and now revealed) random shock such that the outcome is $(x + n) + \omega$. Because the random shock is constant at all locations in policy space, actors who

know one policy and its outcome can ‘invert’ the signal and easily reverse engineer the value of the shock and thus understand the outcome that any other policy will produce (Callander, 2008).

As Callander (2008) argues persuasively, this policy-mapping function inadequately captures policy complexity in many instances. Without rehashing Callander’s (2011b) very thorough critique and explanation, we briefly highlight the intuition. While the canonical model assumes a constant shock, in many policy areas the value of the shock likely varies across the range of possible policies. If this is true, the fact that policy x produces outcome $x + \omega$ does not imply that a different policy denoted $x + n$ produces outcome $(x + n) + \omega$. The assumption of constant shocks seems most likely to break down in areas where policy-making is complicated and small tweaks can lead to large unintended consequences, or when the two policies (x and $x + n$) are particularly far apart.

Callander’s critique is particularly important in situations where actors are trying to learn from either the experience or expertise of others, or from their own past experiences. Knowing the shock that affected one policy may tell one little about the shock that will affect another policy. This makes intuitive sense in many political situations: one competent committee’s markup of a bill does not necessarily shed a lot of light on the outcomes that a different bill on the same topic will produce. One state’s experiment in a challenging policy area does not always eliminate the uncertainty about policies in other states with different goals. A well-informed conservative’s position on a complex political issue does not necessarily tell a poorly-informed liberal what their own ideal position is. In all of these cases, the known information sheds some light on a similar decision with a different goal, but does not remove the uncertainty from it.

Callander (2008) goes well beyond critiquing and proposes an innovative alternative policy–outcome uncertainty representation by modeling the shock as the realization of a Brownian motion. In this new model, uncertainty is represented by a random walk with known variance (denoted by σ^2) around a known linear trend which is called the ‘drift’ (denoted by μ). Actors are assumed to know that the policy process is a Brownian motion and to know its parameters. (See Figure 1 in Callander (2011a), for an example of a Brownian path.) Observing an informed actor’s policy choice reveals one point through which the function passes. This provides some, but not all, information about other policies’ mappings to outcomes. Further, one learns more about policies that are closer to the known policy. One also learns more in less-complex policy-making environments.¹

While seemingly complicated, the math behind the Brownian motion formalization is quite simple and offers a number of attractive properties. Given one known policy–outcome pair, the expected value of the outcome that another policy produces is a linear extrapolation just like in the canonical model, but it is stochastic. The variance around the expected outcome at any point is the product of the Brownian variance and the size of the move from the known policy to the new one. Thus, the variance around this expected value increases the further a new policy is away from the known policy. This uncertainty representation nicely approximates many real-life situations. We often may not know exactly how policies

produce outcomes, but we do have a sense of which direction to tweak an existing policy to achieve different goals. It also approximates many situations in which one can be relatively confident about small changes and relatively uncertain about the consequences of more radical alterations. Finally, one can easily model relatively complex or relatively simple choice environments by changing the random walk's variance.

While not a behavioral model, the Brownian motion framework, and our interest in mimicking and modifying, has interesting overlaps with classic theories of bounded rationality and incrementalism (Lindblom, 1959; Simon, 1978). In some situations, the model predicts the same behavior that a satisficing theory would. Under some conditions, optimizing in the Brownian motion model means mimicking policies which one knows are imperfect.

The Brownian uncertainty model has a number of applications. The substantive context of the predictions we test, and the abstract situation we created in the laboratory, is learning from others' policies. Thus, our experiment speaks in part to the large and growing diffusion literature. (See Karch (2007) and Dobbin et al. (2007) for reviews.) Our focus on micro-mechanisms is unusual in the literature since little work has distinguished mechanisms from each other (Dobbin et al., 2007; Glick and Friedland, 2014) or addressed the challenge of distinguishing diffusion from independent choices (Volden et al., 2008). Thus, an experimental test of a deductive formal model marks a methodological divergence (exceptions include Baybeck et al., 2011; Tyran and Sausgruber, 2005; Volden et al., 2008). The model of mimicking and modifying and the experiment we use to test it, connect formerly disparate ideas in the literature. Our two key variables are similarity (Dolowitz and Marsh, 1996; Grossback et al., 2004; Volden, 2006; Volden et al., 2008; Volden and Makse, 2011) and complexity (Gormley, 1986; Mooney and Lee, 1999; Nicholson-Crotty, 2009; Volden and Makse, 2011). Previous work demonstrated, for example, a link between ideological similarity and policy diffusion (Grossback et al., 2004). We focus on the complexity of policy areas which is at least subtly different than at least some in the literature who focus on the complexity of particular policy instruments (Rice and Rogers, 1980; Volden and Makse, 2011). The difference is that in our approach, complexity is a trait of the challenge the actor is facing. Finally, our focus on modifying builds on ideas of policy 'reinvention' (e.g. Glick and Hays, 1991) including those which consider links between complexity and reinvention (Mooney and Lee, 1999; Rice and Rogers, 1980).

While the predictions and the lab setup are closely related to the diffusion literature, our experiment and results concern human behavior and decision-making capacity more generally. As early as Downs (1957) and Berelson et al. (1954), scholars of political behavior have argued that while most voters are uninformed, they follow informed members of their social groups to make political decisions. Lupia and McCubbins (1998) repeat this claim and argue that it can be rational for voters to do so; some experimental studies explore whether subjects do this in a rational matter (Ahn et al., 2013; Lupia and McCubbins, 1998; Ryan, 2011). However, the cue-taking literature generally focuses on binary decisions and invertible signals.

The model on which we base our test offers a richer conception of the types of decisions and information involved in attempts to learn from others.

Propositions derived from Brownian motion

Callander (2008, 2011a,b) has used the uncertainty representation in models of delegation and of actors experimenting with new policies. Glick (2012) has incorporated the technology into an extensive learning and policy diffusion model that applies to situations in which one actor can learn from another. We focus on two testable predictions concerning ‘mimicking’ (adopting an existing policy) and ‘modifying’ (changing an existing policy for a different context). These predictions are explicit in, and central to, Glick’s model and implicit in Callander’s.

The two propositions we test are formally derived (along with a number of others) in Glick (2012). Rather than reproduce the derivations, we refer to the proposition numbers (1.1 and 3.3, see below) in Glick’s full model that correspond to the two predictions that we test experimentally so that interested readers can see them formally. Nevertheless, we very briefly introduce the basic model and notation in hopes of making the rest of the paper clear. The basic model has two players 1 and 2 with ideal outcomes o_1^* and o_2^* in unidimensional policy space and quadratic loss around those ideal outcomes. The difference in ideal points between the two actors, their goal similarity, is denoted by Δ_{12} . Actor 1 chooses a policy p_1 and then a Brownian motion shock is applied to produce her outcome o_1 . Actors know the parameters of the Brownian motion which are its slope (or drift), denoted μ , and its variance, denoted σ^2 . The first player’s outcome is observable to all and reveals one point through which the Brownian motion passes. It thus provides some useful information about how other policies will map to outcomes. Formally, the expected value (expected outcome) at any other policy (p_2) is just a linear extrapolation based on the slope μ from p_1 to p_2 ($E[p_j] = o_1 + \mu(p_2 - p_1)$). Unlike the canonical model, the uncertainty grows the further policy moves from the known policy–outcome pair. Formally, the variance of policy 2 (p_2) is $|p_2 - p_1|\sigma^2$ such that the variance increases proportional to the distance.

The model investigates player 2’s options after observing player 1’s policy and the outcome it renders while knowing the goal similarity and Brownian motion parameters. The **first proposition** concerns ‘optimal modifying’ (Glick, 2012, proposition 1.1). The question is: when one is modifying a known policy, how much should one modify by? In the Glick model, the second actor arrives at the optimal modified policy after knowing what policy produces the first actor’s ideal point by doing the extrapolation to her own ideal based on the similarity parameter Δ_{12} (as above) and then subtracting $\frac{\sigma^2}{2\mu}$. The ratio of the variance to the slope is indicative of complexity as a higher ratio implies more uncertainty about the outcome that a given policy will produce. Because of the uncertainty, the model predicts moving policy from the observed policy–outcome pair in the direction of one’s own ideal but by less than the gap in the two actors’ ideal points. Moreover, it predicts moving away from the known policy by smaller amounts in more complex issue areas. Substantively, the optimal modified policy should be closer to the well understood

one than it would be without uncertainty and that this closeness should increase with complexity. We test these relationships below by specifying the key parameters in the lab.

The **second proposition** concerns the basic conditions for mimicking (i.e. ideal modification is zero) instead of modifying (Glick, 2012, proposition 3.3). According to the model, one should mimic (not modify at all) when ($\Delta_{12} < \frac{\sigma^2}{2\mu}$). The smaller the difference in the two actors' goals and the larger the complexity, the more likely the mimic condition will be satisfied. Thus, both whether to mimic or modify a new policy, and the magnitude of a modification will vary with the distance to the known policy and vary inversely with task complexity.

These predictions contrast with those that one would derive from traditional models (see above) where the policy mapping function is a constant linear shock. With that alternative uncertainty model, the optimal behavior is to always modify the policy by adjusting for the difference between the two ideal points. This makes the optimal new policy a simple linear extrapolation from the known point to the new ideal point. Note that the difference between the two is not just that players always modify with a simple linear shock. The degree of modification is always greater with the simple linear shock.

These propositions describe optimal behavior when learning from others under Brownian motion uncertainty. These predictions from the model are clear but somewhat counter-intuitive. They are thus ripe for experimental testing. Knowing when to mimic, when to modify, and how much to modify by is a tricky task comprising two major challenges: The challenge posed by the level of (dis)similarity between one's ideal outcome and the outcome produced by the known policy–outcome pair, and the challenge posed by the level of policy complexity. Decreased similarity produces a larger gap between the policy that will in expectation produce the actor's goal and the optimal policy. Increased complexity increases the degree to which an actor should discount information about the distance between the known policy's outcome and the actor's own ideal outcome. These two factors mean that learning from others in such an environment is potentially more cognitively demanding than learning from others if uncertainty is described by the canonical model. We call it cognitively demanding because it requires actors to incorporate additional information and to make potentially counter-intuitive choices even though under some conditions the optimal action is to mimic which may be a very easy choice to make and implement. The key source of additional complexity for decision makers is the fact that optimizing in the Brownian model requires incorporating mean and variance into distance versus uncertainty trade-offs whereas the canonical model only requires distance calculations. The optimal action is not always an intuitive distance adjustment. Even when actors should modify a known policy, they should do so less than the distance between the known policy's outcome and their ideal outcome would suggest. Given the counter-intuitive nature of these theoretical findings, it is far from clear that actors will learn to throw seemingly useful information away or discount it optimally degree.

3. Experimental design

To test the ability of experimental subjects to optimally mimic and modify, we presented them with a simple decision task. In each round they had to pick a number which represented an abstract policy or other input. Subjects were informed that a randomly drawn number, representing a stochastic policy shock, would be added to the number they chose to produce their final outcome. The subject's goal was to end up with a final outcome, representing a policy outcome or other output, that was as close to zero as possible.

Subjects could learn something about the randomly drawn number by observing the action of another player who also chose a policy and had a stochastic shock added to it to produce a final outcome. Specifically, in each round we showed the subject a computer player's choice, the random number that was added to it, and the computer player's final outcome. The difference between the stochastic shock added to the computer player's choice and the stochastic shock added to the human player's choice was the product of a Brownian motion process whose variance depended on the experimental condition and the distance between the numbers picked by the two players. Essentially, subjects faced a decision problem where their ideal outcome was zero, but the further their chosen policy was from the computer player's policy, the less information the computer player's stochastic shock provided about their own stochastic shock (and thus their final outcome).

After seeing the computer player's decision, subjects were asked to choose a number. A random shock was then drawn and added to the subject's choice, and then their final outcome (number) was displayed. A new round then began and each subject was presented with another computer player's number, shock, and final number. Each subject played for 100 rounds.

Subjects were instructed that if they chose the same number that the computer player chose ('mimicking'), they would receive the same stochastic shock, and thus the same outcome as the computer player. This means we are assuming that the two decision contexts are identical such that the mimicked policy produces the same outcome for the second actor as it did for the first.² They were told that if they chose a different number, it would be subject to a different shock, which was defined by the Brownian motion process. We did not attempt to explain the concept of Brownian motion to subjects. Instead, we explained the stochastic component by explaining that the variability of the random number would grow the further away their pick was from the computer players' pick.³ To illustrate how the variance of the shock depended on the distance between their pick and the computer player's pick, we showed them two sets of pictures, examples of which are in Figure 1. First, we showed them four histograms approximating the normal distributions from which their random shock was to be drawn. These four histograms corresponded to picking a number four different distances away from the computer player's choice. The variance of the normal distributions increases as the distance between the subject's pick and the computer player's pick grew. We also showed them a scatter plot with two thousand points drawn from the distribution of shocks that would be drawn in their session. The distance from the computer

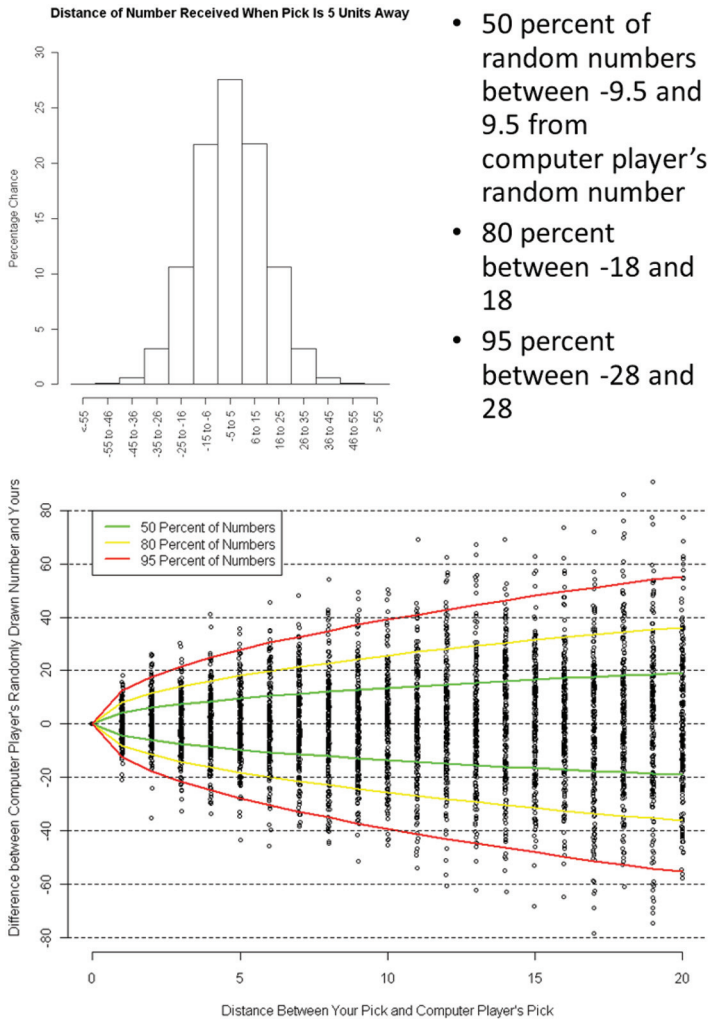


Figure 1. Examples of information provided to subjects to explain the stochastic process to them in the case where the variance was equal to 40.

player was on the x -axis and the shock was on the y -axis. Subjects were given handouts with these figures, and could consult them during the experiment.

Since the experimental task was relatively complicated, we included an additional check on subjects' understanding of the process that turned the number they picked into their final outcome. After the experiment was completed, each subject responded to the prompt 'Please write how you understand that the random number was drawn'. Two independent coders who were unfamiliar with the study's goals coded these responses on a 1–4 scale, where a higher number indicated greater understanding. A majority of subjects' responses were coded 4, and 85% of

Table 1. Optimal actions.

		Complexity	
		High ($\sigma^2 = 40$)	Low ($\sigma^2 = 20$)
Similarity	High ($E(\Delta_{12}) = 10$)	(1) Mimic	(2) Mimic if draw yields $\Delta_{12} < 10$ Modify by $\Delta_{12} - 10$ if draw yields $\Delta_{12} > 10$
	Low ($E(\Delta_{12}) = 20$)	(3) Mimic if draw yields $\Delta_{12} < 20$ Modify by $\Delta_{12} - 20$ if draw yields $\Delta_{12} > 20$	(4) Modify by $\Delta_{12} - 10$

subjects’ responses were coded 3 or 4 by both coders. Removing subjects who scored lower than 3 does not substantively change the results reported in the next section. These open-ended responses give us confidence that the findings we report are not the result of subjects who did not understand the experimental task. Details of the coding procedure and examples of response that received each code are available in Appendix E.

Experimental predictions

In Glick’s (2012) article discussed above, two factors influence a player’s decision to mimic or modify: the complexity of the decision task and the degree of similarity between the player’s goals and the goals of the informed player. These are represented mathematically as the variance (simply σ^2 after assuming a drift of 1) and the distance between the player’s goal and the known policy outcome (denoted Δ_{12}). To test learning with Brownian motion uncertainty, we varied these two factors as experimental treatments in a 2×2 design. The subject’s goal (optimal policy outcome) was fixed at zero. The distance between the subject’s goal and the known policy’s outcome was either small ($\Delta_{12} = 10$) or large ($\Delta_{12} = 20$). The variance of Brownian motion process was also manipulated to make the task either complex ($\sigma^2 = 40$) or simple ($\sigma^2 = 20$). This operationalization of complexity fits our concept of ‘policy area’ complexity because it affects how well actors can predict the consequences of their policy choices. Table 1 shows these four experimental conditions. To ensure that subjects did not face seemingly identical tasks each time, the distance between the subject and the computer player was drawn from a uniform distribution on $(E(\Delta_{12}) - 5, E(\Delta_{12}) + 5)$.⁴ In all cases, the underlying trend or drift (μ) of the Brownian motion was 1.

The parameters were chosen to test the model’s predictions that players should mimic when tasks are more complex or when similarity is high, and modify when tasks are less complex or when similarity is low. The optimal policy choices for the experimental treatments are summarized in Table 1. We cross these two manipulations to produce four experimental conditions. We denote each condition using an uppercase ‘H’ or ‘L’ for the level of complexity and a lower case ‘h’ or ‘l’ for the

level of similarity. Thus the ‘high-complexity–high-similarity’ condition is denoted by ‘H-h’, the ‘high-complexity–low-similarity’ condition is denoted by ‘H-l’, and so on.

In the H-h condition (1), subjects should always mimic; the high degree of similarity means that the computer player’s outcome is close to the subject’s ideal outcome while the high degree of complexity means that modifying introduces a great deal of risk. In the L-l condition (4), players should always modify because the computer player’s outcome is far from the subject’s ideal outcome and the low level of complexity reduces the risk introduced by modifying. In the L-h condition (3) and H-l condition (2) the expected value of Δ_{12} is on the cut point between modifying and mimicking. Depending on whether the draw of Δ_{12} is above or below $E(\Delta_{12})$ in a given round, players should either mimic or modify. Thus, optimal behavior would lead them to mimic half of the time and modify half of the time. When modifying, the optimal average modification would be five units. Also of interest is how subjects would behave if they ignored the complexity of the decision, that is, the aspect that makes the Brownian motion representation unique, and acted as though uncertainty was described in the same way as the canonical model. The canonical constant shock model would predict the same strategy (modifying by Δ_{12}) in all four conditions. This strategy represents the prediction that subjects will simply subtract the difference between their goal and the other player’s outcome from the other player’s optimal action to achieve their goal of zero.

Comparing actual behavior with the optimal behavior as described in Table 1, as well as comparing this to how subjects would act under the canonical model, will show whether subjects are able to adjust their behavior to the kind of uncertainty described by the Brownian motion process. Comparing how close subjects come to optimal behavior across experimental conditions will show how subjects respond to different kinds of difficulty in learning from others. Specifically, we can determine how the complexity of a task affects optimal decision-making by comparing the H-h condition with the L-h condition and the H-l condition with the L-l condition. We can determine how the similarity (or lack thereof) between decision makers affects optimal decision-making by comparing the H-h condition with the H-l condition and the L-h condition with the L-l condition. Finally, comparisons between subjects in the L-h condition and the H-l condition will allow us to tell whether subjects respond in different ways to these two types of task difficulty.

We completed six sessions of the experiment with 86 subjects, who were undergraduates at a private northeastern US research university. Each session was randomly assigned to be a ‘high-complexity’ session or a ‘low-complexity’ session, and subjects within each session were randomly assigned to either high-similarity or low-similarity. Five subjects’ data were dropped for failing to comply with experimental instructions leaving 81 subjects in four experimental conditions. A total of 24 were in the H-h condition, 17 in the L-h condition, 20 in the H-l condition and 24 in the L-l condition. Subjects played the game for 100 rounds. We selected one round at random for payment. Subjects were paid US\$10 minus the squared distance (in cents) between their outcome in the randomly selected round and 0.⁵ They also received a US\$5 show-up fee, and could earn up to US\$12 from a risk preferences elicitation.⁶

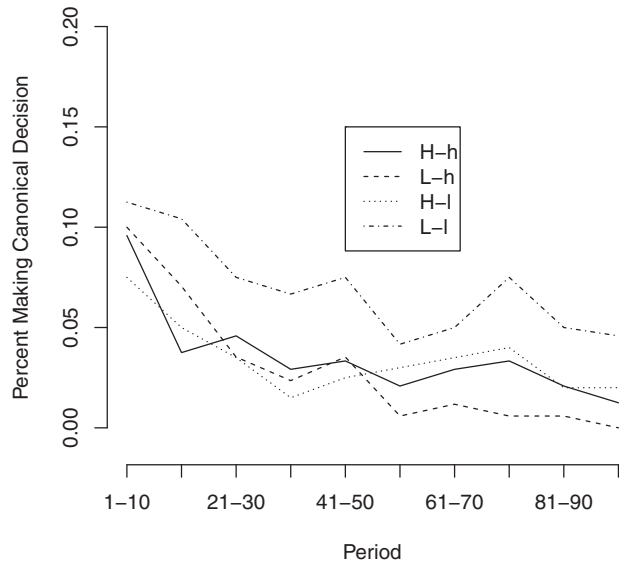


Figure 2. Percentage of decisions matching canonical model by 10-period blocks.

Average earnings were US\$16.81 including average earnings of US\$5.23 from the risk preference elicitation. The experiment was programmed and conducted with the software z-Tree (Fischbacher, 2007); sessions lasted approximately 1 hour.

Results

We first establish that subjects did not behave as they would under the canonical model. Figure 2 shows the percentage of decisions that, in expectation, would produce the subject’s ideal outcome; that is, decisions that would be optimal if uncertainty was characterized as in the canonical model, rather than via Brownian motion. In all four conditions, this percentage is around 10% in the first 10 rounds, but drops rapidly to less than 5% in all four conditions. Most subjects appear capable of understanding that optimal behavior with Brownian motion uncertainty is different from the more intuitive canonical model.

The fact that subjects did not behave as though they were facing canonical uncertainty does not, of course, mean that their actions conform to the predictions of the Brownian motion model. We use two metrics to measure how subject behavior compares with the theoretical predictions. The first is deviation from the optimal action: the number of units and direction that the subject’s action falls from the optional action as defined in Table 1. This number can be positive or negative. A positive deviation indicates over-modifying the action whose outcome was known, while a negative deviation indicates that the subject was too cautious and under-adjusted from the known action. One way to think of this measure is as a measure of bias in the statistical sense of the term where an unbiased subject would take, on average, the optimal action. Similarly, subjects in an experimental condition

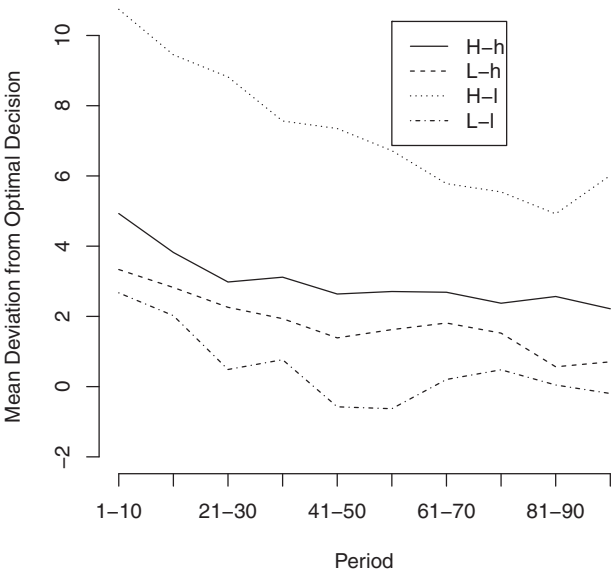


Figure 3. Mean deviation from optimal action by 10-period block.

Table 2. Mean values by condition, with 95% confidence intervals.

	L-l	L-h	H-l	H-h
Percentage mimicking	0.09_a(0.02, 0.29)	0.42 _b (0.23, 0.61)	0.18_a(0.07, 0.29)	0.57_b(0.42, 0.72)
Correct mimicking	0.91_a(0.84, 0.98)	0.69_b(0.58, 0.79)	0.56_c(0.50, 0.62)	0.57_{b,c}(0.42, 0.72)
Average deviation	-0.08 _a (-2.42,2.27)	0.64 _a (-0.02,1.29)	5.48_b(3.13, 7.80)	2.39_c(0.88, 3.90)

Note: **BOLD** indicates statistically different from optimal (*t*-test, $\alpha=0.05$) Means that share a subscript cannot reject the null hypothesis of no difference (*t*-test, $\alpha=0.05$). (For example, difference between percentage mimicking in L-l and L-h is statistically significant, difference between L-l and H-l is not.)

that produces unbiased results select the correct action on average, even if some generally over-adjust and some under-adjust. To give a sense of how subjects behaved over time, Figure 3 shows the average deviation from the optimal modification for each condition grouped into 10-period blocks. To conduct statistical tests, we analyze subject-level data by computing each subject’s average deviation from the optimal action across the final 20 rounds of play.⁷ We then compare the distributions of these variables to the model’s predictions and across experimental conditions. Figure 6 shows these results graphically; Table 2 reports a full range of statistical tests.

The second metric is how frequently the subject mimicked the actions of the computer player. According to Table 1, subjects should always mimic in the H-h

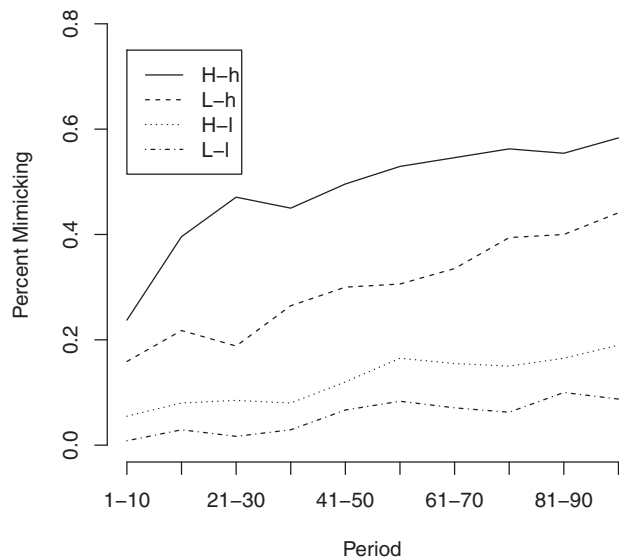


Figure 4. Percentage of decisions mimicking by 10-period block.

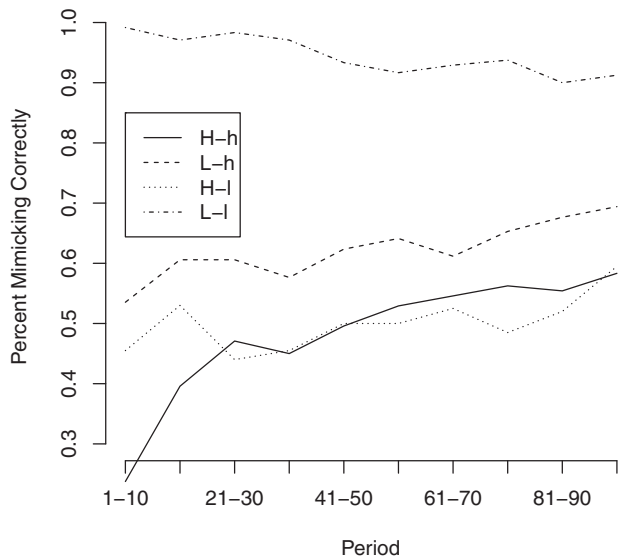


Figure 5. Percentage of decisions correctly mimicking by 10-period block.

condition, should mimic when the computer player’s outcome is less than 10 in the L-h condition, should mimic when the computer player’s outcome is less than 20 in the H-l condition, and should never mimic in the L-l condition. This means that we should expect subjects to mimic 50% of the time in the L-h condition and in

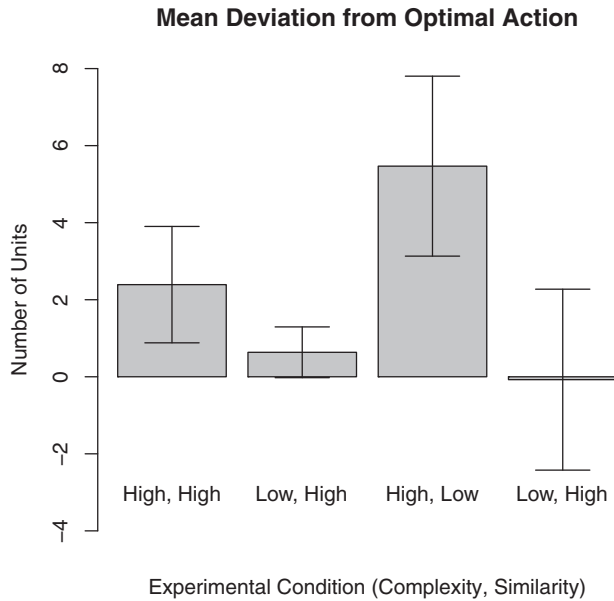


Figure 6. Subjects' mean deviation from optimal action, last 20 rounds.

the H-l condition, every time in the H-h condition and never in the L-l condition. Translated into comparative static predictions, subjects should mimic the most in the H-h condition, mimic less frequently in the L-h and the H-l conditions, and mimic the least in the L-l condition. Figure 4 shows the over time trends in the percentage of mimicking decisions in each condition. We evaluate whether each decision to mimic or not was correct; that is, whether the subject mimicked if the optimal action was to mimic and modified if the optimal action was to modify. Figure 5 shows the percentage of correct decisions over time. For hypothesis testing we compute each subject's percentage of times mimicking as well as the percentage of times each subject made the correct mimic/do-not-mimic decision in the final 20 rounds of play. Figures 7 and 8 show the results of this test graphically. In general, average subject behavior deviates from optimal behavior in two ways: a tendency to over-modify in high-complexity situations, and a bias towards modifying even when mimicking is the optimal choice.

We start by examining the L-l condition. This condition offers the cleanest test of subjects' ability to behave in ways consistent with the Brownian motion model as the optimal action is always to modify, but by less than one would under the canonical model.⁸ Specifically, the optimal action is to modify the first player's choice by 10, halfway between the known policy and the choice that would be ideal under the canonical model. It does not produce the subject's ideal outcome in expectation, but it is the best she can do given the uncertainty introduced by the Brownian motion process. Looking at overtime trends in Figures 3 and 4, we see

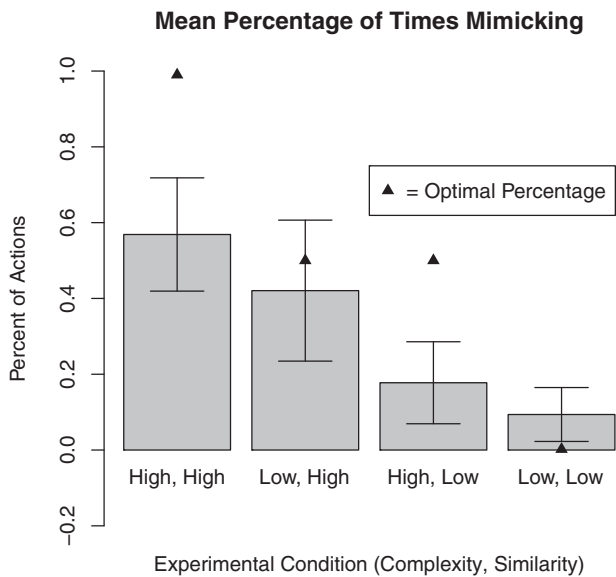


Figure 7. Subjects' percentage of decisions that mimic, last 20 rounds with 95% confidence intervals.

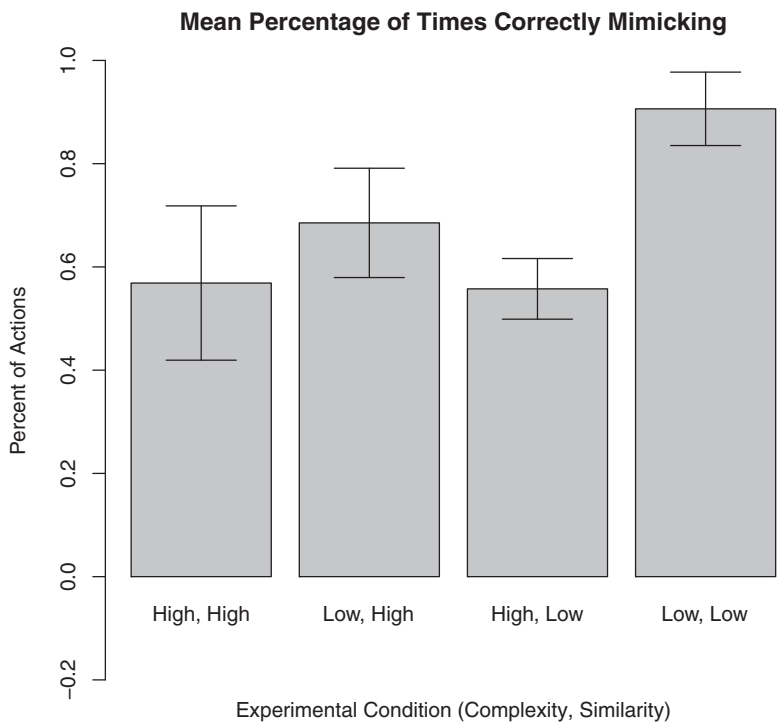


Figure 8. Subjects' percentage of percentage of correct mimicking decisions, last 20 rounds with 95% confidence intervals.

that while subjects begin by over-modifying, by the 50th round the mean deviation approaches zero. Subjects in this condition almost never mimic at any point in the experiment. While the percentage mimicking increases over time, it never goes above 9% of subject decisions.

Looking at subject-level data aggregated over the final 20 rounds, we see a similar pattern that is largely in line with the model's predictions. The average subject's mean deviation from the optimal action is very close to, and not statistically different from, zero ($p = 0.951$, two-sided t -test).⁹ The average subject mimicked nine percent of the time over these rounds, a statistically significant difference from the predicted zero ($p = 0.017$), but one that still shows subjects modifying in all but a substantively small number of cases.

Moving to the L-h condition, we again see mean results that are generally in line with the Brownian Motion model. In this condition, the higher similarity means that subjects should mimic half of the time. In the over-time results, we see the percentage of times mimicking increases over the course of the experiment but never exceeds 44%. Similarly, the mean deviation from optimal action in this condition decreases over time and approaches, though does not quite reach, zero. In the final 20 rounds, the average subject's mean deviation was 0.63 units, a deviation that is marginally statistically different from zero ($p = 0.078$ two-sided t -test). The average subject mimicked 42% of the time in this period, a percentage that is not statistically different from the expected rate of 50% ($p = 0.415$, two-sided t -test); subjects made the correct mimicking decision in 69% of cases, a rate of correct decisions notably smaller than in the L-l condition, but higher than in either of the high-complexity conditions. Increasing the similarity between the subject's goal and the known policy makes the decision task slightly more difficult, but does not appear to seriously degrade subjects' performance.

In contrast, average behavior in the H-l condition shows that increasing the complexity of the decision causes subjects to deviate significantly from optimal behavior. Like in the L-h condition, subjects should mimic 50% of the time in this condition. Here, it is the higher degree of complexity that makes mimicking optimal half the time rather than the higher degree of similarity (see above). Figure 3 shows that decision-making in this condition improved little after the 40th round, with the average decision modifying by three units more than optimal. Similarly, Figure 4 shows that the rate of mimicking in the H-l condition is never above 19%, and that it stops increasing around round 60. Turning to subject-level behavior in the last 20 rounds, we find that the average subject deviated by 5.47 units, significantly different from optimal action ($p < 0.000$). This deviation is also significantly greater than the deviation in the L-l condition ($p = 0.002$) and in the L-h condition ($p = 0.001$). The average subject mimicked 17.7% of the time over this period, significantly less often than the optimal level of 50% ($p < 0.000$) and less often than subjects in the L-h condition ($p = 0.036$), but not significantly different from the 9% rate in the L-l condition. The average subject made the correct mimic/do not mimic decision 56% of the time, significantly less than the 69% figure for the L-h condition ($p = 0.049$). This low rate of correct decisions was not the result of mimicking too frequently, as 91% of these mistakes were from subjects not mimicking when mimicking was

optimal. When faced with increased complexity, subjects mimic less and deviate more from the optimal action than in analogous low-complexity conditions.

Results from the H-h condition also suggest that subjects struggled with highly complex Brownian motion processes. In this condition, subjects should always mimic. Figure 4 shows that subjects mimicked in a higher percentage of decisions in this condition than any other, and that the percentage of decisions in which they mimicked increased over time. However, even in the final 10 rounds only 58% of decisions were mimicking, and the average decision was a modification of 2.2 units. The average subject mimicked 56% of the time in their final 20 decisions, significantly different from optimal rate of 100% ($p < 0.000$). The average subject's mean modification in this period was 2.4, also significantly different from zero ($p = 0.005$). As in the H-l conditions, subjects faced with high-complexity Brownian Motion processes deviate significantly from optimal action even after 100 rounds of play and mimic significantly less than is optimal.

Finally, in addition to results based on average behavior, we examine heterogeneity of behavior across subjects. Figures 9 and 10 present histograms showing each subject's average deviation during the final 20 rounds, percentage of times mimicking in the final 20 rounds, and percentage of times mimicking correctly in the final 20 rounds. We remove a few extreme values for graphical clarity. As these figures show, there was a great deal of heterogeneity in subjects' actions even in experimental conditions where mean behavior was relatively close to theoretical predictions.

We highlight two aspects of heterogeneity across subjects. First, subjects in low-similarity conditions showed considerably more heterogeneity in their average deviation from optimal behavior than did subjects in high-similarity conditions. The average deviations by subjects in high-similarity conditions were tightly clustered. In the L-h condition 8 of 17 subjects averaged within 1.5 units of the correct action, and 15 of 17 subjects in this condition averaged within 2 units. In contrast, most subjects in the low-similarity conditions either consistently over-modified or consistently under-modified. Thus, in the L-l condition, 6 of 24 subjects (20%) over-modified by an average of five units or more, while another 6 subjects under-modified by an average of 5 units or more. The H-h condition showed a similarly degree of heterogeneity though in this case the distribution was not symmetric around zero, and 8 of 20 (40%) subjects over-modified by at least 5 units.

Second, subjects in most conditions displayed a great deal of heterogeneity in their tendency to mimic. In particular, the bias towards action that we found looking at average rates of mimicking seems to affect some subjects more than others. Half of all subjects in the H-l condition and 35% of subjects in the L-h condition never mimicked, despite the fact that mimicking was called for in an average of half of all cases. Even in the H-h condition, where mimicking was always the optimal choice, 13% of subjects never mimicked. In general, Figure 9 shows that subjects in the high-similarity conditions adopted a much greater variety of mimicking strategies than subjects in the low-similarity conditions. However, this mostly reflects more subjects in the low-similarity conditions adopting a 'never mimic' strategy, a strategy that produced a very high rate of correct mimicking decisions in the L-l

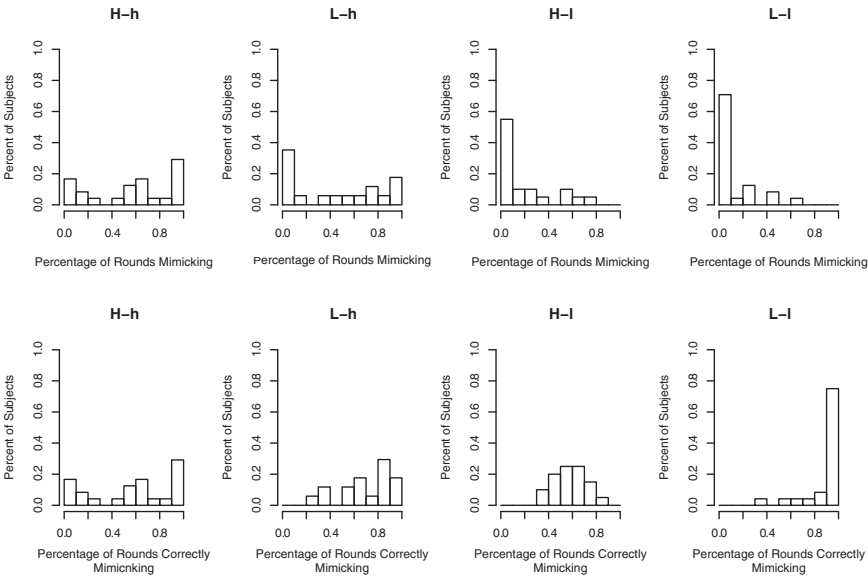


Figure 9. Histogram of subjects' mimicking behavior, last 20 rounds with 95% confidence intervals.

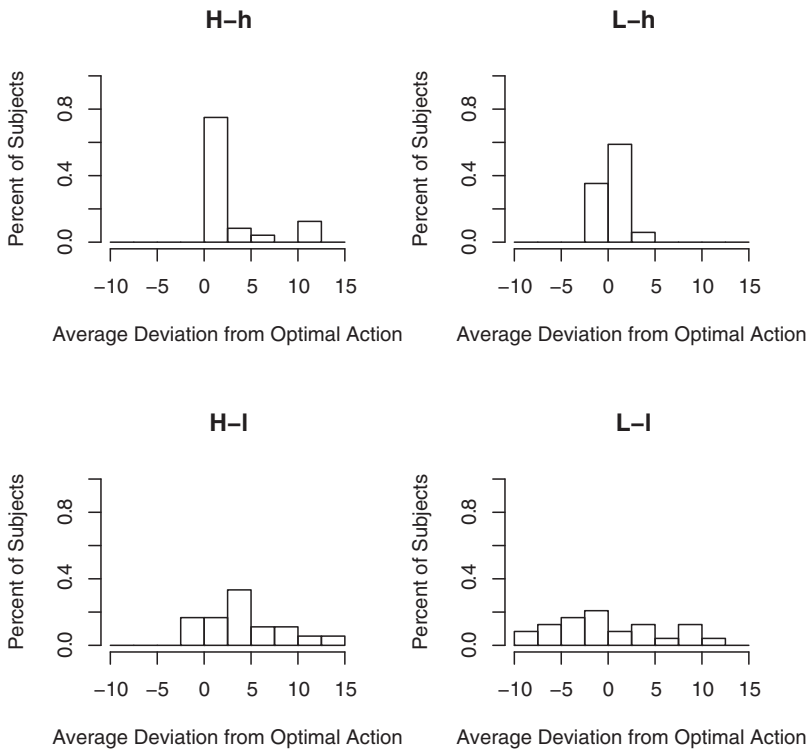


Figure 10: Histogram of subjects' mean deviation from optimal action, last 20 rounds with 95% confidence intervals.

condition but fewer correct decisions in the H-I than in either of the high-similarity conditions.

The one condition without a great deal of heterogeneity in mimicking behavior was the L-I condition, a fact that provides further evidence of a bias towards action. In this condition, mimicking is never optimal, so a bias towards action will push subjects towards acting more correctly. Thus, subjects who are biased towards action and those who are not biased towards or against action will make the correct decision to modify. Indeed, in this condition 17 out of 24 (71%) of subjects never mimicked in the final 20 rounds, and only 12.5% mimicked in more than 25% of the final 20 rounds.

Discussion

This study presents subjects with a form of outcome uncertainty characterized by Brownian motion, a form that is more complicated than the outcome uncertainty in canonical models. While the canonical model presents subjects with a fairly simple decision task, Brownian motion presents them with two more complicated decisions: whether to modify a policy with a known but sub-optimal outcome and, if modifying, how much to modify by. Can subjects navigate these more complicated decisions? The broad answer is yes. Subjects do respond to this more complicated decision environment sensibly, although not quite optimally. Few subjects ignore the more complicated decision environment and act as though they are dealing with the canonical model. Further, their actions trend towards optimal behavior over time, and in the low-complexity conditions their behavior approximated, on average, optimal behavior.

Nevertheless, even after an extended period of learning, subjects do not learn from others in fully optimal ways. The results from this experiment suggest two things about learning from others in situations where policy uncertainty is described by a Brownian motion process. First, the degree to which subjects' behavior is described by the Brownian motion model depends on the complexity of the decision they face. In situations where complexity was low, subject's behavior was generally in line with optimal behavior. However, when complexity was higher, subjects deviated from the predictions of the model in important systematic ways. Subjects underestimated the effect that this complexity has on their decisions and over-modified the policy that they were learning from. This may not come as a big surprise. Complexity, which refers to the amount of variance in the Brownian process, is the element that separates uncertainty as described by Brownian motion from uncertainty as described in the canonical model. The more complex a problem is, i.e. the more uncertainty there is between policies and outcomes, the larger the gap in optimal action in the two models. Unfortunately, this suggests that situations where a Brownian motion conception of complexity offers the greatest theoretical advance over the canonical model are also those where the Brownian motion model does the worst job of predicting behavior.

Second, particularly when complexity was high, subjects showed a bias towards action. In high-complexity conditions subjects make fewer correct mimicking

decisions because they do not mimic in cases where mimicking is optimal. In the low-complexity condition where the optimal rate of mimicking was non-zero subjects still made correct mimicking decisions only two-thirds of the time and mimicked less often than optimal, although not by a statistically significant margin. This may be evidence of action bias (Kahneman and Tversky, 1984; Patt and Zeckhauser, 2000), in which a norm of action causes decision makers to take an action even when inaction is optimal. In the Brownian motion context, decision makers should sometimes mimic a known policy (inaction) despite knowing that it does not produce their ideal outcome. We suggest that subjects in this experiment are operating under a norm of adjusting their behavior in response to seemingly useful information, and against simply discarding this information. The sense that they will regret violating this norm in the case of a bad outcome creates a bias towards action. This suggests that the predictions of Brownian motion models may fare particularly poorly in situations where inaction is a major prediction. When it is optimal to discard information the temptation to tinker may be too great to resist. Alternately, the bias towards action may be highly dependent on the prevailing norms of the decision environment. In some decision scenarios, a norm of inaction or copying¹⁰ might help subjects to act more rationally.

Conclusion

Brownian motion offers a more nuanced, flexible, and theoretically attractive representation of uncertainty. It has applications to institutions and individuals in a wide variety of political settings, of which our test of learning from others is only one. However, behaving consistently with equilibrium predictions predicated on this underlying uncertainty model is more difficult than behaving consistently with the canonical model of uncertainty. Despite this difficulty, our experiment finds that individuals did a respectable job of learning from others' choices in a rich and challenging choice environment.

At the same time, we find that subjects' choices deviate from full rationality in systematic ways. Most notably, people are better at working with variations in goals than they are with highly uncertainty policy outcomes. These results do not suggest rejecting the use of Brownian motion to model uncertainty in decision-making. Instead, they suggest caution in interpreting the predictions of these models in certain situations. The biases our results demonstrate represent systematic deviations from the behavior predicted by Brownian motion models in situations where the Brownian motion process is highly complex. Taking these biases into account when interpreting the predictions of these models will help empirical and formal scholars use them with proper care.

Our findings offer twists on two cognitive biases that affect decision-making: action bias (as discussed above) and anchoring bias. Anchoring describes the tendency of decision makers to choose actions that are closer to a numerical 'anchor' than they would if the anchor, however random, was not present (Tversky and Kahneman, 1974; Wilson et al., 1996). In most anchoring experiments, subjects are asked a question requiring a numerical answer after being presented with a

numerical value that is unrelated to the question. These studies find that subjects' answers are influenced by the value of the 'anchoring' number. In the present study, the computer player's decision, which is presented immediately before the subject makes a decision, provides a potential anchor. Anchoring is usually thought of as a bias that leads actors to act irrationally by giving weight to an irrelevant anchor. However, since rational behavior in the context of Brownian motion uncertainty demands making a smaller modification to a known policy than might seem intuitive at first glance, anchoring bias may actually help actors behave optimally. Indeed, the original mimicking and modifying model may be characterized as a theory where a meaningful rather than random anchor makes anchoring rational.

Instead, we find that subjects deviate too far from this informative anchor. Given that people anchor when the initial values are completely arbitrary, the fact that they deviate too far from informative starting points is noteworthy and suggests interesting overlaps between behavioral economics and rational choice work using sophisticated uncertainty models. Alternately, subjects may be anchoring on something other than the computer player's decision, for example, the distance between the computer player's ideal outcome and their own.¹¹ In either case, while this study's aim was not to test the effect of psychological biases these results suggest that the application of anchoring theory to action in the context of Brownian motion uncertainty is not straight-forward.

One important substantive application of our findings is the diffusion of policies among the states and in other contexts. One important diffusion question is how well policy makers can learn from existing policies and implement their lessons in different contexts; for example, how a state such as New Hampshire can learn from the policies of a state such as California. With only some extrapolation to policy diffusion, the lab results suggest that policy makers in New Hampshire likely do a good job adjusting for the fact that they may have very different policy needs than California. However, in complex policy areas these different needs may lead them to over-adjust a policy, while a bias towards action may lead them to change a policy even when simply copying the policy is the best course of action. This is potentially problematic since complex policy areas may be the areas in which smaller states most rely on learning from bigger and better resourced ones. Needless to say, attempting to verify our findings in applied policy settings with observational data would be ideal but challenging. Nevertheless, at a minimum, one could use interviews, surveys, and/or survey experiments with actual policy makers to test whether they perceive these challenges, and whether they have more confidence in adjusting for goal similarity than for complexity.

Our initial empirical investigation of learning and Brownian motion suggests several other lines of future work. One angle would be to test how performance varies with substantively important changes in the environment. One could, for example, run the same experiment on a pool of individuals with experience making policy choices, allow people to work in teams like they often do in policy choice settings, and/or add a level of realism by adding an accessible concrete story to the choices they are making. Moreover, recent theoretical work concerning learning in

plausible real-world situations with more than one existing policy to learn from can easily provide additional hypotheses to test. One could use a very similar experiment to test how well individuals incorporate information from two known policies (Callander, 2011b; Glick, 2012), or how well they balance similarity and competence tradeoffs when choosing from multiple actors to learn from (Glick, 2012). As in this paper, testing these predictions would provide insight into decision tactics with real applications and into the validity of the underlying uncertainty innovations. In addition, our model uses a simple decision-theoretic environment where subjects do not interact. An important next step is to investigate behavior in game-theoretic environments where actors' beliefs about how others will respond to Brownian motion uncertainty are important.

Acknowledgements

A previous version of this paper was presented at the 2010 MPSA Conference. We'd like to thank Nolan McCarty, Adam Meirowitz, Graeme Boushey, the Princeton Laboratory for Experimental Social Science, and participants at the Research in American and Comparative Politics Colloquium at Boston University. All errors are, of course, our own.

Declaration of Conflicting Interests

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Funding

The author(s) disclosed receipt of the following financial support for the research, authorship, and/or publication of this article: This project received funding from the Princeton Laboratory for Experimental Social Science, Boston University, and the Robert Wood Johnson Foundation Scholars in Health Policy Program.

Notes

1. These invertibility assumptions and variations are also critical to models of social learning and information transfer (e.g. information cascade models (Bikhchandani et al., 1992, 1998)). This work demonstrates how invertibility assumptions affect predictions. In the cascade models signals become less and less invertible as more actors make decisions. The Brownian models are more concerned with signals which are always partially invertible and less restricted to a string of sequential choices.
2. This simplifying assumption mirrors one in the model of Glick (2012). Exploring the implications of relaxing it would make for interesting subsequent work, but testing how people deal with the basic mimic/modify tradeoffs does not require doing so.
3. For complete instructions, see Appendix B. For instructional handouts, see Appendix C.
4. Appendix A shows how subject behavior varies across this uniform distribution of draws.
5. To avoid the possibility of negative earnings (Morton and Williams, 2010, p. 367), subjects for who US\$10 minus the squared distance in cents was less than 0 earned US\$0. This accounted for a relatively small number of cases, 6.2% of subject decisions.

6. See Appendix F for a discussion of the results of this elicitation.
7. We have also examined the final 10 and 30 rounds of play, with substantively similar results.
8. We thank an anonymous reviewer for the suggestion to present the results in this way, and particularly for suggesting this manner of discussing the L-I condition first.
9. We report two-sided *t*-tests in this paper, but have also conducted Wilcoxon exact tests. We report any cases where these results find meaningfully different levels of significance.
10. Consider, for example, the IT professional adage ‘No one ever got fired for buying IBM’.
11. We thank an anonymous reviewer for suggesting this possibility.
12. Full results are available from the authors upon request.

References

- Ahn TK, Huckfeldt R, Mayer A and Ryan JB (2013) Expertise and bias in political communication networks. *American Journal of Political Science* 57(2): 357–373.
- Baybeck B, Berry WD and Siegel DA (2011) A strategic theory of policy diffusion via intergovernmental competition. *The Journal of Politics* 73(1): 232–247.
- Berelson B, Lazarsfeld P and McPhee W (1954) *Voting: A Study of Public Opinion Formation in a Presidential Campaign*. Chicago, IL: University of Chicago Press, 1954.
- Bikhchandani S, Hirshleifer D and Welch I (1992) A theory of fads, fashion, custom, and cultural change as informational cascades. *Journal of Political Economy* 100(5): 992.
- Bikhchandani S, Hirshleifer D and Welch I (1998) Learning from the behavior of others: Conformity, fads, and informational cascades. *Journal of Economic Perspectives* 12: 151–170.
- Callander S (2008) A theory of policy expertise. *Quarterly Journal of Political Science* 3(2): 123–140.
- Callander S (2011a) Searching and learning by trial-and-error. *American Economic Review* 101: 2277–2308.
- Callander S (2011b) Searching for good policies. *American Political Science Review* 105(2): 643–662.
- Dobbin F, Simmons B and Garrett G (2007) The global diffusion of public policies: Social construction, coercion, competition, or learning. *Annual Review of Sociology* 33: 449–472.
- Dolowitz D and Marsh D (1996) Who learns what from whom: A review of the policy transfer literature. *Political Studies* 44: 343–357.
- Downs A (1957) *An Economic Theory of Democracy*. New York: Harper, 1957.
- Eckel CC and Grossman PJ (2008) Forecasting risk attitudes: An experimental study using actual and forecast gamble choices. *Journal of Economic Behavior and Organization* 68(1): 1–17.
- Fischbacher U (2007) z-tree: Zurich toolbox for ready-made economic experiments. *Experimental Economics* 10(2): 171–178.
- Gilligan TW and Krehbiel K (1987) Collective decisionmaking and standing committees: An informational rationale for restrictive amendment procedures. *Journal of Law, Economics, and Organization* 3(2): 287–335.
- Gilligan TW and Krehbiel K (1989) Asymmetric information and legislative rules with a heterogeneous. *American Journal of Political Science* 33(2): 459–490.
- Glick DM (2012) Learning by mimicking and modifying: A model of policy knowledge diffusion with evidence from legal implementation. *Journal of Law, Economics, and Organization*, in press. DOI: 10.1093/jleo/ews041.

- Glick DM and Friedland Z (2014) How often do states study each other? evidence of policy knowledge diffusion. *American Politics Research*, in press.
- Glick HR and Hays SP (1991) Innovation and reinvention in state policymaking: Theory and the evolution of living will laws. *The Journal of Politics* 53(3): 835–850.
- Gormley WTJr (1986) Regulatory issue networks in a federal system. *Polity* 18(4): 595–620.
- Grossback LJ, Nicholson-Crotty S and Peterson DAM (2004) Ideology and learning in policy diffusion. *American Politics Research* 32(5): 521–545.
- Hayes AF and Krippendorff K (2007) Answering the call for a standard reliability measure for coding data. *Communication Methods and Measures* 1(1): 77–89.
- Huckfeldt R and Sprague JD (1995) *Citizens, Politics, and Social Communication: Information and Influence in an Election Campaign*. Cambridge: Cambridge University Press, 1995.
- Kahneman D and Tversky A (1984) Choices, values, and frames. *American Psychologist* 39(4): 341.
- Karch A (2007) Emerging issues and future directions in state policy diffusion research. *State Politics and Policy Quarterly* 7(1): 54–80.
- Lindblom CE (1959) The science of muddling through. *Public Administration Review* 19(2): 79–88.
- Lupia A and McCubbins M (1998) *The Democratic Dilemma: Can Citizens Learn What They Need to Know?* Cambridge: Cambridge University Press, 1998.
- Mooney CZ and Lee M-H (1999) Morality policy reinvention: State death penalties. *The Annals of the American Academy of Political and Social Science* 566(1): 80–92.
- Morton RB and Williams KC (2010) *Experimental Political Science and the Study of Causality: From Nature to the Lab*. Cambridge: Cambridge University Press, 2010.
- Nicholson-Crotty S (2009) The politics of diffusion: Public policy in the american states. *The Journal of Politics* 71(1): 192–205.
- Patt A and Zeckhauser R (2000) Action bias and environmental decisions. *Journal of Risk and Uncertainty* 21(1): 45–72.
- Rice RE and Rogers EM (1980) Reinvention in the innovation process. *Science Communication* 1(4): 499–514.
- Ryan JB (2011) Social networks as a shortcut to correct voting. *American Journal of Political Science* 55(4): 753–766.
- Simon HA (1978) Rationality as process and as product of thought. *The American Economic Review* 68(2): 1–16.
- Tversky A and Kahneman D (1974) Judgment under uncertainty: Heuristics and biases. *Science* 185(4157): 1124–1131.
- Tyran J-R and Sausgruber R (2005) The diffusion of innovations - an experimental investigation. *Journal of Evolutionary Economics* 15(4): 423–442.
- Volden C (2006) States as policy laboratories: Emulating success in the children's health insurance program. *American Journal of Political Science* 50(2): 294–312.
- Volden C and Makse T (2011) The role of policy attributes in the diffusion of innovations. *The Journal of Politics* 73(1): 108–124.
- Volden C, Ting MM and Carpenter DP (2008) A formal model of learning and policy diffusion. *American Political Science Review* 102(3): 319–332.
- Wilson TD, Houston CE, Etling KM and Brekke N (1996) A new look at anchoring effects: Basic anchoring and its antecedents. *Journal of Experimental Psychology: General* 125(4): 387–402.