

On the Complexity of Forming Mental Models ^{*}

Chad Kendall[†] Ryan Oprea[‡]

October 3, 2022

Abstract

We experimentally study how people form predictive models of simple data generating processes (DGPs), by showing subjects datasets and asking them to predict future outputs. We find that subjects: (i) often fail to predict in this task, indicating a failure to form a model, (ii) often can't explicitly describe the model they've formed even when successful, and (iii) tend to be attracted to the same, simple models when multiple models fit the data. Examining a number of formal complexity metrics, we find that all three patterns are well-organized by metrics suggested by Lipman (1995) and Gabaix (2014) that describe the information processing required to deploy models in prediction.

Keywords: Complexity, mental models, inference, bounded rationality, behavioral economics, economics experiments

JEL codes: C91, D91, C0

^{*}We thank Guillaume Frechette, Cary Frydman, Thomas Graeber, Ariel Rubinstein, Andrew Schotter, Emanuel Vespa, Georg Weizsacker and Sevgi Yuksel for valuable conversations about this research. We are grateful to seminar audiences at New York University, Berkeley Haas, Yale SOM, Stanford, Surrey, Queensland, the Berlin Behavioral Economics Seminar, Texas Tech University, the 6th World Congress of the Game Theory Society, the Workshop on Applications of Revealed Preferences (WARP), and USC. This research was supported by the National Science Foundation under Grant SES-1949366 and was approved by UC Santa Barbara IRB.

[†]Kendall: Department of Finance and Business Economics, Marshall School of Business, University of Southern California, 701 Exposition Blvd, Ste. 231 HOH-231, MC-1422, Los Angeles, CA 90089-1422, chadkend@marshall.usc.edu.

[‡]Oprea (corresponding author): Economics Department, University of California, Santa Barbara, Santa Barbara, CA, 93106, roprea@gmail.com.

1 Introduction

In order to understand the world, we continuously have to form *mental models* of data generating processes (DGPs). That is, we have to form beliefs about the causal dependence of outcomes we care about on other observables (e.g. characteristics, past events) and deploy those beliefs in prediction.¹ These inference problems are fundamental, arising regularly across domains in social and economic life. For instance, in order to strategically interact with others, we have to use the history of play to infer the strategy they are using. In order to predict other people’s behavior, we have to infer the habits, heuristics, norms and demand functions that drive their decision making. In order to make political decisions, we have to form models about how policies and geopolitical events conspire to generate the social or economic outcomes we care about. To understand any natural process or solve an engineering problem, we have to form predictive models of nature’s laws.

Although there is a long experimental literature studying how and why subjects fail at simple statistical inference (i.e. Bayesian inference problems, e.g. Benjamin (2019)), there has been almost no investigation of the equally fundamental question of how (and how well) people form predictive mental models of the sorts of algorithms that connect variables across social and economic life. This is important because even when data generating processes are perfectly deterministic, requiring no statistical reasoning, they may nonetheless be difficult to accurately model because doing so requires elaborate pattern recognition and sophisticated representation of what has been recognized. We call this difficulty “inferential complexity.” What makes a data generating process difficult to infer from data? What makes a coherent model *complex* to construct from the observation of past events? What kinds of models are people good at forming and what kinds of models do people form when many models are available to explain data?

In this paper, we report an experiment designed to take some first steps towards answering these kinds of questions. In order to do this, we reverse the approach typically taken in the inference literature: instead of studying statistical inference in

¹We refer to “mental models” (rather than “models”) to emphasize that they need not be formal, explicit or even consciously available to decision makers (DMs) in order to guide choice.

stochastic but algorithmically trivial environments as the literature generally does, we study deterministic but algorithmically complex inference settings. The experiment consists of a series of tasks each consisting of three parts. In Part 1, we show subjects a dataset – a sequence of twelve *inputs* (“a” or “b” repeated in some sequence) and, following each, an *output* (e.g. “x” or “y”) – generated by a deterministic, but unknown, algorithm that may (or may not) respond to the history of inputs. We tell subjects virtually nothing about the true data generating algorithm other than that it may or not respond to inputs and that it is not random. The subject’s job in Part 2 is to predict the outputs that this same algorithm will produce after each of twelve further inputs, presented to the subject in sequence. In order to guess effectively, the subject must “extract” a model – a rationalizing algorithm – from the Part 1 data (even if unconsciously). Based on their choices in Part 2 we can infer whether or not they were successful at forming an accurate model. In Part 3, we incentivize subjects to explicitly “articulate,” in words, the model they have extracted.

In our main treatment (called “Unique”), we apply this three-stage process to a variety of DGPs, each a distinct, deterministic algorithm connecting inputs to outputs. Specifically, we vary the underlying DGP in a series of eight distinct tasks. Even though we constrain ourselves to the simplest algorithms, we are able to confront subjects with a rich variety of data generating processes that arise regularly in social life. For instance, in “autonomous” processes outputs cycle independently of the inputs, mirroring e.g. seasonal oscillations in traffic patterns or the daily habits of a co-worker. By contrast, with “instruction” processes, the current input directly determines the output, as in the way income predictably affects consumer demand, or the way someone playing a tit-for-tat strategy directly responds to their opponent’s past action in a repeated prisoner’s dilemma game. In “switch” processes, the input instead causes a *change* in the output e.g. as in consumers that switch brands when they receive poor customer service, or voters who oust the current incumbent whenever the economy is bad.² Our algorithms include these DGPs and more elaborate

²Of course, our interest in how successful people are at recognizing and distinguishing between these and other DGPs. To take the last example, the question is how difficult it is for a politician to discover that voters change their vote in response to a bad economy (a “switch” process) when other possibilities exist *ex ante*. Voters may instead simply tire of incumbents so that the process is “autonomous.” Or, they may vote for one party when the economy is good and the other when it is bad, following an “instruction” rule. Which is the politician most likely to recognize? When

DGPs that generate outputs based on more intricate patterns in the history of inputs.

We find strong evidence that extracting a correct model from data and using it to guide prediction is quite difficult for the average subject. Subjects are only able to extract a model that rationalizes the Part 1 input string about 40% of the time.³ What’s more this ”extraction rate” is highly variable across the types of DGPs we study, with some inducing extraction rates as high as 76% and others as low as 5%. Crucially, we find that this variation has structure both within- and between- subjects: subjects that extract models of on-average difficult DGPs tend to extract on-average simpler DGPs too. This regularity suggests that a latent complexity ordering over DGPs exists, governing failures to extract models across subjects.

The richness of the set of algorithms we study allows us to look for formal structure in this complexity ordering by evaluating the predictive power of a number of notions of complexity culled from various disciplines. We cast a wide net, collecting notions and metrics from computer science, information theory, algorithmic information theory, and economics. Some of these notions are rooted in characteristics of the DGPs themselves and the implementation costs they exact – for instance the minimum number of states (state complexity) or transitions (transition complexity) in the simplest description of the algorithm, or the minimal resource burden of implementing the algorithm on physical hardware. Others, such as measures of the Kolmogorov Complexity or the Shannon entropy of the output or input string, are instead rooted in characteristics of the datasets subjects observe. Still others such as the mutual information between inputs and outputs, the sensitivity of outputs to past inputs, and measures of the fineness with which decision makers must partition datasets (”partition complexity”) or (relatedly) the number of inputs/outputs that must be attended to (”sparsity complexity”) in order to apply a model, are rooted in both the DGP and the kinds of datasets they produce.

multiple processes are consistent with the data, which model is the politician most likely to adopt?

³In the Unique treatment, we say a subject has extracted a model if she guesses correctly in Part 2 of the experiment. Of course, a subject may fail in this task simply because she forms a model that is more complicated than the true one, leading to mistakes. This possibility is one of the key reasons we asked subjects to verbally articulate their model in part 3 of each task. Examining this data, we find no evidence that subjects extract models that rationalize the Part 1 dataset but that require more states than the true DGP.

We find that two closely related notions that describe the information processing required to apply a model – “partition complexity” (adapted from Lipman (1995)) and “sparsity complexity” (adapted from Gabaix (2014)) – are the best at predicting extraction, with impressively strong correlations in excess of 2/3. By contrast, some other measures (e.g. mutual information, entropy, Kolmogorov Complexity) do quite poorly, with correlations with extraction of well below 0.5. Others, such as computational complexity and especially the number of transitions in the simplest description of the DGP as a finite state machine, do somewhat better. Perhaps surprisingly, the number of states in the simplest finite state machine description of the DGP – a commonly used measure of the complexity of algorithms in other contexts – does a poor job of organizing the difficulty of forming models.

In addition to studying people’s ability to “extract” a model to forecast, we also study people’s ability to consciously “articulate” a model. We find strong evidence that subjects are often able to extract models that they are unable to explicitly describe. That is, a lot of the learning involved in model formation is implicit rather than explicit. Nonetheless, we find that the explanatory power of complexity notions is nearly identical for articulation and extraction: our information processing measures (partition and sparsity complexity) are highly correlated with articulation rates across DGPs, just as they are with extraction rates.

In the final step in our investigation, we study what types of models people are drawn to when multiple models are consistent with the data. In our “Multiple” treatment, we ran subjects through an identical series of tasks to those in the Unique treatment, featuring the exact same DGPs. However, unlike in Unique, we deliberately showed subjects Part 1 datasets that are consistent with *multiple* distinct DGPs and therefore are consistent with multiple distinct mental models.⁴ Structurally estimating subjects’ models from their Part 2 choice behavior, we find evidence that subjects are systematically drawn to *simple* models. Under almost any complexity notion available, subjects tend to “select” the simplest model that rationalizes the data set at a rate higher than random. Of these, once again our information pro-

⁴Technically, any dataset is consistent with an infinite set of mental models. By “unique”, we mean unique within the class of models that can be described by finite automata with fewer than four states. Our results from the articulation task show that this refinement is not restrictive – we find no evidence from verbal descriptions that subjects form rationalizing models with more states.

cessing notions (partition and sparsity complexity) – the notions which best predict rates of extraction and articulation – also best predict model selection when multiple data-consistent models are available. Conditional on extracting a model, subjects extract the partitions/sparsity-simplest model 2/3 of the time (compared to a 1/3 rate predicted for a random model selector).

These findings are relevant to a growing literature in economics on “mental models.” Most closely related perhaps is Aragonés et al. (2005) which points out that forming a model from a noisy dataset (“fact free learning”) is computationally complex (“NP Hard”) and therefore likely to be intractable for decision makers as datasets grow large (as the number of variables increase). We show that model formation can be intractably difficult even under deterministic and minimally sized (two variable) datasets, and that this is because factors other than the dimensionality of the dataset (that can vary even in the simplest datasets) generate severe complexity burdens of their own. Our work is also related to recent theoretical (Esponda and Pouzo (2016), Bohren and Hauser (2021), Fudenberg, Romanyuk and Strack (2017), Heidhues, Koszegi and Strack (2018), and Gagnon-Bartsch, Rabin and Schwartzstein (2021)) and empirical (Hanna, Mullainathan and Schwartzstein (2014), Handel and Schwartzstein (2018), Esponda, Vespa and Yuksel (2021), Enke (2020), Graeber (2021)) work showing that decision makers can get “stuck” in misspecified mental models, and have difficulty revising these models with learning. Our work complements this literature by showing that there is predictable structure in the types of models people are likely to get “stuck” in in the first place. Our finding that decision makers are initially “biased” towards models that are simple (in the sense that they require little information processing) may be useful in guiding the construction of models in this literature. Finally, there is a recent theoretical literature studying the implications of people being unable to form accurate causal models (Spiegler (2016), Eliaz and Spiegler (2020)), focusing on DGPs in which the direction of causality must be inferred (i.e. describable by Directed Acyclic Graphs). We study simpler DGPs in which the direction of causality is known, but we emphasize that our methods can (and should) be extended to these more difficult settings.⁵

⁵Related to Eliaz and Spiegler (2020) is Schwartzstein and Sunderam (2021), who study the ability of senders to persuade receivers by supplying them models. Our interest lies instead in how people form models that are not supplied to them by others.

Our work also relates to a literature in economics studying the effects of complexity on human decision making. Most relevant is the “automata” literature, which models decision rules and procedures (e.g. strategies in games) as finite automata, simple descriptions of algorithms from computer science (see Lipman (1995) and Chatterjee and Sabourian (2009) for reviews). The main strand of this literature studies the effect of “implementation complexity” – the subjective costs of employing decision rules describable by complex automata – on decision making (Rubinstein (1986), Abreu and Rubinstein (1988), Kalai and Stanford (1988)). Oprea (2020) measures these implementation complexity costs experimentally by paying subjects to follow rules describable by automata and then eliciting their willingness pay to avoid being assigned these same rules again in the future. In contrast, our paper builds on a second strand of the automata literature which studies not the implementation complexity of following automata, but rather the “computational complexity” of *inferring* automata from data (Gilboa (1988), Ben-Porath (1990)). This literature shows theoretically that automata are complex to infer in the sense that they become computationally intractable to detect as the number of states in the underlying automata grows. Our work (i) supports this literature by showing that this inference problem is indeed empirically difficult for decision makers and (ii) extends this literature by showing that it is difficult even for the simplest (two- and three-state) automata because of sources of complexity unrelated to the size of the automaton.^{6,7}

Finally, our work relates to several strands of literature in psychology. First, there is a long literature, beginning with Byrne and Johnson-Laird (1989), examining whether people form mental models when trying to solve propositional and spatial reasoning

⁶Our paper also relates less directly to work that studies the complexity of other objects than automata. For instance, Murawski and Bossaerts (2016) shows that metrics of the computational complexity of instances of optimization problems (knapsack problems) predicts failures to solve them. Abeler and Jäger (2015) show that increasing the complexity of taxes reduces responsiveness to their marginal incentives. Jin, Luca and Martin (2021) show that people use complexity instrumentally to shroud information.

⁷Our paper also relates to a literature that models bounded rationality by postulating costs associated with metrics of decision problems. For instance, an important literature models the effects of attentional costs, often measured by mutual information, on attention costs (e.g. Sims (2003), Caplin (2016), Dean and Neligh (2019)). Likewise, Gabaix (2014) assumes that sparser models of decision problems are less costly to use and examines the implication for standard optimization problems in economics. We do not fit or test these models directly but instead examine the predictive power of some of the key underlying metrics they are built around (e.g. mutual information, sparsity).

problems.⁸ By comparison, the DGP-inference problem we study is unambiguously related to mental models (it would be difficult to perform in our task without forming a model of some sort), making it a particularly valuable setting for studying what makes forming such models difficult. Second, our finding that articulation of models is more difficult than extraction of models parallels a large literature in psychology that makes a distinction between “explicit” and “implicit” learning. Most of this literature studies very different settings from ours, uses very different methods, and asks fundamentally different questions. Most closely related to our work is a small literature on “biconditional grammars” (Mathews and Buss (1989) and Shanks, Johnstone and Staggs (1997)) which attempts to distinguish between implicit and explicit learning.⁹ Subjects are given a DGP similar to our “instruct” DGP, and asked to identify a rule that relates one sequence of letters to another (“x” goes with “T”, etc.). To identify implicit learning, these papers vary whether or not subjects know such a rule is being applied. Because our primary goal is to understand what makes a DGP difficult to model, we instead vary the DGP, holding fixed subjects’ knowledge that there is a DGP. To identify explicit learning, these experiments also ask subjects to articulate the rule, but we do so in an incentive compatible way, arguably leading to more credible data.

The remainder of this paper is organized as follows. In Section 2, we describe the problem we are interested in, the class of algorithms we use to generate models in the experiment and preview a number of metrics of complexity that form hypotheses for our experiment. In Section 3, we discuss and motivate our experimental design and, in Sections 4 and 5, our empirical results. We conclude in Section 6 with a discussion.

⁸In propositional reasoning, subjects are faced with a series of logical propositions and attempt to derive the logical conclusion. Spatial reasoning problems are similar except that the statements describe locations of objects (e.g. “C” is in front of “x”).

⁹The two most popular tasks are artificial grammar tasks (Reber (1967)) and serial reaction time tasks. In the former, subjects observe test “sentences” constructed from a Markovian process and then attempt to determine whether or not subsequent sentences are consistent with the grammar of not. In the latter, they respond to a visual stimulus as fast as possible (see Jimenez, Vaquero and Lupianez (2006) for a recent contribution). In both tasks, subjects are not told beforehand that a set of rules is generating the patterns they see so that any ability to do better than chance in subsequent guesses is thought to occur through implicit learning.

2 Conceptual Background and Questions

Consider a DM who observes a “dataset”, D^T , consisting of (i) a variable called “inputs,” $u_t \in U$ observed at dates, $t = 1 \dots T$ and (iii) a second variable called “outputs,” $v_t \in V$ observed at dates, $t = 0 \dots T$. The dataset is formed by a “data generating process” (DGP) that produces an initial output, v_0 , and subsequent outputs, v_t , at each date t based on the sequence of inputs prior to t , $u_1 \dots u_t$.

We deliberately focus on what is arguably the simplest class of DGPs: DGPs in which (i) inputs and outputs are discrete, (ii) the mapping between inputs and outputs are deterministic and (iii) the direction of causality between inputs and outputs runs weakly in a single direction. As we discuss below, we focus on this class of DGPs in order to limit the number of forces influencing the DM’s inference problem – a choice which allows us to draw sharper conclusions on the drivers of the complexity of inference in a particularly simple domain. Nonetheless, our questions and methods can and should be extended to richer DGPs that include stochasticity (e.g. by studying Markov chains) or unclear causality (by studying directed acyclic graphs) in future work.

DGP’s in this class are describable as finite state machines (or finite automata) – among the simplest languages for modeling and describing algorithms in computer science. Although in what follows we will use automaton descriptions to describe the DGPs we study, we emphasize that these DGPs could equally have been described by a great many alternative formalizations of algorithms (some quite rich and sophisticated). We use automata as a descriptive language here largely for expositional purposes and to help to taxonomize and organize the DGPs we select for our design. Later, (in Section 2.3) we will compare how well features of automata descriptions (like states and transitions) compare to a number of alternative formal ways of describing these same DGPs in predicting complexity responses.

Formally, an automaton, M , is a four-tuple, (S, s^0, f, τ) , where S is a finite set of states, $s^0 \in S$ is the initial state, $f : s \rightarrow V$ is an output function that specifies the output $v \in V$ in each state, and $\tau : S \times U \rightarrow S$ is a transition function that maps

the current state and input, $u \in U$, into the subsequent state.¹⁰ Table 1 contains a sample of eight DGPs, diagrammed using automata descriptions. In these diagrams, (i) each circle represents a state, $s \in S$, (ii) each arc a transition, (iii) letters next to arcs show inputs $u \in U$ and (iv) letters inside circles represent outputs, $v \in V$, in each state. A freestanding arrow pointing to one of the states indicates the initial state, s^0 . For instance, automaton I2 (the 2-state “Instruct” DGP) specifies that (i) output “x” appears at date 0 and (ii) afterwards output “x” always appears in the period following input “b” while “y” always follows input “a”.

2.1 Examples

Anticipating our experiment, the examples in Table 1 are restricted to DGPs that are non-trivial for our purposes – ones describable by “fully connected” automata with no source states (states which are entered but never exited) and no sink states (those which, once entered, are never left).¹¹ Likewise we restrict attention to DGPs with the *simplest* input and output alphabets, each containing only two values (e.g. $U = \{a, b\}$, and $V = \{x, y\}$).

The left column in Table 1 shows the four canonical DGPs that are describable by 2-state automata and have only two inputs and outputs (modulo swaps of labels). A2 describes an “autonomous” process that is unresponsive to inputs, switching back and forth between x and y . I2 describes an *instruct* process: each input value uniquely determines the subsequent output value, creating a direct mapping (an instruction) between the previous input and current output. S2 describes a *switch* process: input b leaves the output unchanged, while input a toggles the output, forcing the output to change from its previous value. Finally, H2 describes a *hybrid* process in which one input (b) performs an instruct function while the other (a) performs a switch function.

¹⁰Technically an automaton with such a description is a Moore machine whose outputs only depend on the state. Mealy machines allow outputs to also depend on the current input. There is no loss of generality in focusing on Moore machines as every Mealy machine also has a Moore machine description.

¹¹Formally, any automaton can be described as a directed edge graph. We restrict attention to automata for which the associated directed graph is fully connected: every state is reachable from every other.

These four DGPs summarize common processes that arise in inference problems. Autonomous process (A2) are patterns in outputs that are causally distinct from any other variables – simple regularities like the change from day to night or cyclical fads in consumer tastes. Instruct processes (I2) are simple causal relationships that reliably dictate the conditions under which outcomes occur – like the predictable way a person responds to incentives or a chemical is required to generate a reaction. Switch processes (S2) are processes that do not directly cause outcomes, but instead cause changes in outcomes – like the way investors may flip between value and growth stocks when past performance is bad or foraging animals may switch food sources whenever food is scarce. Hybrid (H2) processes blend instruct and switch processes.

In the right hand column of Table 1 we show the most direct extensions of these four basic processes to richer DGPs – richer in the sense that they require 3-states to describe with automata instead of 2-states. These rules maintain the non-triviality (no sinks or sources) and simplicity (a 2-value input and output library) but require an additional state grafted to each of the four basic 2-state processes.^{12,13}

2.2 Forming Models

Our interest is in understanding how a DM who observes a dataset $D^T \equiv \{u_t\}_{t=1}^T \cup \{v_t\}_{t=0}^T$ generated by an exogenous input string $\{u_t\}_{t=1}^T$ and DGP M forms a “model” of the DGP. We mean three distinct things by “form a model”:

Extraction: Suppose a DM after observing D^T observes a further sequence of inputs $\{u_t\}_{t=1+T}^{2T}$ and must guess or forecast the subsequent output sequence $\{v_t\}_{t=1+T}^{2T}$. We say a DM has “extracted” a model, M , of the DGP if she can correctly list the further output sequence that M produces in response to the continuation of the input string. This ability to reliably “imitate” a DGP is evidence of implicit learning or implicit

¹²An alternative way of extending our four 2-state DGPs to three states would be to add another output value (“z”). Anticipating our experiment, we opted against this to avoid simultaneously changing two things when moving from 2-state to 3-state DGPs (the size of the output alphabet and the number of states).

¹³Note that unlike DGPs describable as 2-state automata which have only four basic types (among non-trivial, 2 input/output DGPs), 3-state DGPs have 356. We therefore selected 3-state algorithms that seem most similar to the four 2-state exemplars in the left column of Table 1.

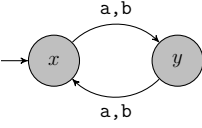
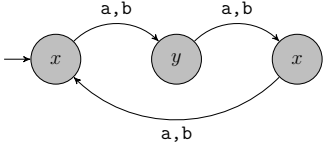
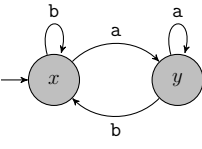
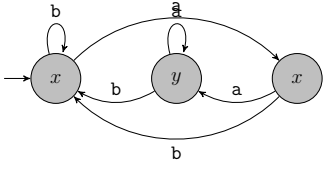
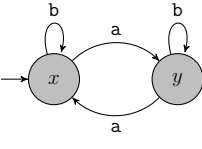
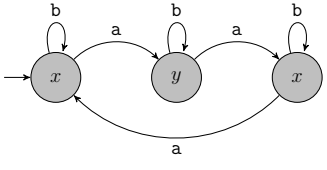
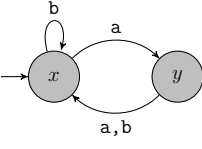
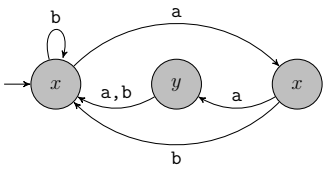
DGP	2-State DGP	3-State Extension
Autonomous <i>A fixed pattern, independent of inputs.</i>	 A2	 A3
Instruct <i>A direct mapping of previous input to output.</i>	 I2	 I3
Switch <i>A direction to move back or forth over a sequence of outputs.</i>	 S2	 S3
Hybrid <i>One input instructs, the other prescribes a Switch</i>	 H2	 H3

Table 1: Examples of data generating processes (DGPs) described using finite automata. Notes: We include the 4 types of non-trivial (no sink or source states), simple (2-input and 2-output) 2-state automata and examples of a 3-state extension of each of these. These are the DGPs we assign to subjects in the experiment.

formation of a model.

Articulation: Suppose a DM after observing D^T is asked to directly describe the DGP M (e.g. using words). We say a DM has “articulated” a data-consistent model of the DGP if she can correctly describe M . In order to do this, the DM has to be conscious of M . Our distinction between “articulation” and “extraction” is not natural in economic theory, but it mirrors a long-made distinction between “explicit

learning” (learning of something one is aware of) and “implicit learning” in psychology (Reber (1967)).

Selection: Suppose a DM observes a D^T which can be rationalized by any of a set of models/DGPs, $M \in \mathcal{M}$, in the sense that the output sequence $\{v_t\}_{t=0}^T$ could be generated by the input sequence $\{u_t\}_{t=1}^T$, $\forall M \in \mathcal{M}$. We say that the DM has “selected” a model if she has successfully extracted or articulated an $M \in \mathcal{M}$ (even if it is ex post not the “true” model of the DGP). Here, too, the DM has “formed a model” but has had to choose between several candidates. One of our interests in the experiment described below is to understand how a DM chooses between or searches over the multiple rationalizing models in $\mathcal{M}$.

Our aim is to study and compare each of these three aspects of model formation.

2.3 Complexity

We are interested in understanding (i) what makes it more difficult to form a model of one DGP than another (i.e. what drive rates of “extraction” and “articulation” across DGPs) and (ii) what makes one model more likely to be formed when more than one model is consistent with the dataset (i.e. what drives “selection”).

We are particularly interested in the possibility that either or both might be driven by the “complexity” of the inference task itself. By this, we mean that there is an ordering across DGPs $\succsim_c$ such that if $M' \succsim_c M$, M is easier for any DM to extract or articulate than M' , and there is a (possibly different) ordering, $\succsim_{c'}$, such that M is more likely to be selected by any DM than M' when $M, M' \in \mathcal{M}$.

The existence of a consistent complexity ordering is hardly obvious – many other possibilities exist. For instance, it may be that individual instead DMs have idiosyncratic and heterogeneous talents or proclivities for extracting/articulating models of some types of DGPs over others. We call this alternative possibility “affinity” (in the sense that individual DMs may have different individual affinities for identifying some DGPs over others). Even if, in the aggregate, rates of extraction or articulation differ across DGPs, this may be driven not by a consistent complexity ordering over DGPs, but instead by different affinities being more or less common in the population.

Likewise, it is not obvious, *ex ante*, that selection should be governed by a systematic preferential ordering over models. While it is possible that DMs are systematically drawn to some types of models over others, it is also possible that something more closely resembling random search governs selection. In this case, what matters for selection might be random, with DMs more likely to successfully infer models whenever more models are available. And even if there is a systematic preferential ordering over models governing selection, it is possible that it is a different ordering from the complexity ordering driving the difficulty of extraction and selection.

Finally, if there is such a thing as a consistent complexity ordering governing model formation, we are interested in understanding what generates that complexity and predicts it. Is there a metric on DGPs (or datasets) that organizes the difficulty of forming models? Does this same metric describe how people choose between equally justifiable models when more than one is available? There are many *ex ante* notions of complexity (some related to automata descriptions and some completely unrelated) we might consider and in Section 5.1 we consider a number of them, drawing notions from information theory, computer science and economics.

3 Experimental Design

3.1 The Experimental Task

Each session of the experiment consists of eight *tasks*. Each task consists of three Parts, described below.

3.1.1 Part 1: Observing the Dataset

In Part 1 of each task, we show subjects twelve *inputs* (letters, “a” or “b”) one at a time on their screens. The inputs are drawn independently with equal probability of each letter. After each input appears, a downward arrow and an output (a letter, “x” or “y”) appears below it. The outputs are determined by one of the DGPs from Table 1 (varied across tasks) based on the history of inputs. We inform the subject

that the input sequence is entirely exogenous, but that outputs are determined by a rule (possibly linked to the prior input sequence), implemented by the computer. Subjects are told nothing about the rule except that (i) it does not choose outputs randomly and (ii) outputs may or may not depend on the preceding history of inputs under the rule. We specify (ii) to avoid deception since two of our rules (A2 and A3) do not depend on prior inputs, and we did not tell subjects that outputs can depend on past outputs because it is technically incorrect - outputs depend on the state which does not map one-to-one to an output in the 3-state rules.

We call the set of input/output pairs for Part 1 the task’s “dataset.”

3.1.2 Part 2: Extracting a Model of the DGP

In Part 2 of each task, we show subjects a 13th input and ask them to guess which output comes next in the sequence. They submit their guess by typing a letter (“x”, “y”) which shows up on the screen after it is typed. Immediately after typing her guess, a 14th input appears and the subject must guess the next output in the sequence. Inputs continue to appear on the subject’s screen until she has observed 12 inputs in Part 2 (24 inputs in total including those from Part 1) and submitted 12 guesses.

If the task is selected for payment, the subject is paid \$0.35 for every output she guesses correctly. Importantly, the subject *is not shown* the correct sequence (or given any feedback on the correctness of her guesses) until Part 3 of the task has concluded, long after she has made all of her Part 2 guesses.

We call the subject’s ability to repeat the Part 1 DGP in Part 2, “extraction”, and we say a subject has successfully extracted a model if she chooses a sequence of guesses that is consistent with a DGP that (i) can rationalize the Part 1 dataset and (ii) is describable by the class of automata with less than four states and no source or sink states. (In Section 4.4, we show that criterion (ii) is not binding in our data: subjects essentially never describe models that satisfy (i) but violate (ii).)

3.1.3 Part 3: Articulating a Model of the DGP

In Part 3 of each task, we induce the subject to describe (“articulate”) the model she has formed (if any) to explain the dataset of inputs and outputs in Part 1, by asking her to describe the rule she thinks the computer used to generate the Part 1 sequence of outputs. The subject types a description of the rule in a text box on her screen.

The subject is told that this description will (with 10% likelihood) be given to a future participant who will face a guessing task similar to Part 2, with outputs determined by the same rule but in response to a different sequence of inputs. That future participant will not have access to a Part 1 dataset and will therefore rely exclusively on the subject’s description of the rule to guide her guesses. The future participant therefore will only be able to correctly guess the sequence of outputs if the subject correctly and clearly articulates a model of the data generating process.

To incentivize the subject to articulate the DGP/rule she inferred from Part 1, we pay her \$0.35 for each output the future participant manages to guess correctly after reading the subject’s description of the rule.¹⁴ In this we follow previous experimental work incentivizing advice (e.g. Schotter and Sopher (2003)).

3.2 Tasks and Treatments

We collected data in two treatments that differ only in the specific datasets we assigned to subjects.

3.2.1 The Unique Treatment

In the first treatment – called “Unique” – we assign each subject eight Tasks, each consisting of a dataset (input/output sequence) that can only be rationalized by one, unique, DGP that is discredable by a fully connected automaton of fewer than four

¹⁴We took pains to assure subjects that the future participant would be sophisticated enough to understand a well-articulated rule in order to further incentivize effort. Specifically we told them that the participant would be a graduate student in a quantitative field. In practice, we gave each description selected for payment to a graduate student whose guesses were then used to determine payment.

states. We have two aims with this treatment. The first is to study the degree to which subjects can extract and articulate the unique model that describes the data generating process and how these rates of extraction and articulation vary with characteristics of the generating rule. The second is to study the degree to which subjects’ ability to extract and articulate a model varies with characteristics of the dataset presented to them.

In Unique, we studied the eight data generating processes listed in Table 1 and discussed in Section 2.1. For each DGP we created four distinct datasets and assigned each to 1/4 of subjects. Thus in the Unique treatment we vary DGPs within-subject and, for each, vary datasets between-subject.

3.2.2 The Multiple Treatment

In the second treatment – called “Multiple” – we inverted the design of the Unique treatment. We assigned subjects the same eight DGPs as in Unique, but gave all subjects Part 1 input/output sequences that could be rationalized by *multiple* 2- and 3-state models/DGPs. In Multiple, therefore, subjects could extract any one of a number of simple models, each of which is consistent with Part 1 evidence (but only one of which is ex post correct in Part 2). The number of such models capable of rationalizing the dataset ranged between 2 and 20 depending on the task. Our goal with the Multiple treatment is to study how subjects choose among models when multiple models can rationalize the dataset – a model-formation task we call “selection.” To collect enough data to credibly answer this question, we gave all subjects in Multiple the same dataset for each DGP.

3.3 Implementation

We ran the experiment in June of 2021 on the online research panel Prolific using custom Javascript programmed by the authors and deployed using Qualtrics.¹⁵ Each session consisted of detailed instructions (reproduced in Online Appendix C) explain-

¹⁵We restricted participation to residents of the U.S. and those that had not previously participated (only). The average age of participants was 31.7 and 46.6 percent were male.

ing the task followed by a set of comprehension questions that subjects were required to correctly answer before moving on to the experiment. Subjects experienced the eight DGPs in a random order across eight tasks. In the Unique treatment, each subject was, with equal likelihood, assigned one of four distinct datasets in each DGP. Experimental sessions took on average 31 minutes. Subjects on average earned \$5.63 including the showup fee, resulting in an hourly rate of \$10.90 which greatly exceeds Prolific’s minimum required rate of \$6.50. We targeted 100 subjects per treatment and in the end collected data for 99 subjects in Unique and 101 subjects in Multiple.

3.4 Understanding the Design

We designed this experiment with several inferential goals in mind.

First, we wanted to directly measure how difficult it is to form models of data generating processes and we designed the Unique treatment with this measurement in mind. By providing subjects with a setting in which there is only one simple model to explain the data we can directly study the mapping between the true data generating process and the difficulty of forming a model of it.

Second, we wanted to separately evaluate the difficulty of inferring a model implicitly (to guide future behavior) and explicitly (to formally describe it). For this reason, we included both Part 2 (in which subjects only have to reproduce and imitate the pattern of the data generating process) and Part 3 (in which subjects have to be conscious enough of their model of the DGP to articulate it in words) in each task in the design. Including both allows us to directly contrast what we call “extraction” of a model with “articulation” of the same model, within subject. Furthermore, because subjects are free to describe any model they like in Part 3, we can see whether or not it is restrictive to test for extraction within the class of models describable by automata with less than four states and no sink or source states.

Third, we wanted to study what formalizable features of DGPs – what notions of complexity from literatures in computer science, information theory, algorithmic information theory and economics – make them more and less difficult (“complex”) for humans to model. To achieve this we attempted to vary DGPs in a way that

plausibly varied difficulty, given recent work in the literature. We achieved this by (i) deploying DGPs that are describable by the canonical four types of 2-state automata and (ii) the most obvious three-state-describable extensions of these four canonical types. This allowed us to vary states in the automaton description (an important driver of implementation complexity in past work, see Oprea (2020)) and also generated significant variation in a number of other complexity metrics that are unrelated to automaton descriptions.

Fourth, although we wanted to vary difficulty, we deliberately limited ourselves to the simplest class of model formation problems available on several margins. For instance, in order to focus on the complexity of extracting patterns, we studied purely deterministic problems in which explicit statistical reasoning is not required. Our methods can easily be extended to study how model formation and its complexity is impacted by stochasticity by, for example, applying them to more general Markov chains. Likewise, we simplified our problems significantly by making the direction of causality clear. But our methods can be easily adapted to study Directed Acyclic Graphs (DAGs), which are typically used to describe causal relationships. We limited ourselves to DGPs with two possible inputs and two possible outputs. But, again, our exact methods can be extended to richer datasets.

Fifth, we wanted to ensure that our results weren't overfitted to idiosyncracies in the datasets we assigned to subjects in Part 1. Varying the dataset also allowed us to study the performance of complexity notions that depend on the details of the dataset rather than the underlying DGP. For this purpose, we need sufficient data for each dataset. Striking a balance between these two goals, we generated four datasets for each DGP in the Unique treatment and assigned them evenly across subjects.

Sixth, we wanted to study not only what makes a given DGP difficult to model but also what models people are drawn to when more than one can rationalize the data. In particular, we wanted to study the degree to which these two model-formation tasks – extraction and selection – are related to one another. To accomplish this we introduced a second treatment (“Multiple”) with the same true DGPs as the Unique treatment, but with Part 1 datasets that could be rationalized with more than one simple model/DPG.

Finally, we parameterized our tasks so that the Part 1 dataset and the Part 2 guessing task each included 12 inputs with corresponding outputs. 12 inputs/outputs in the Part 1 dataset struck a balance between being long enough to be able to find datasets that ensure unique models for the Unique treatment, but not so long as to preclude datasets consistent with multiple models for the Multiple treatment. 12 inputs/outputs in the Part 2 dataset is roughly the smallest number we need to structurally distinguish between models that are consistent with the Part 1 dataset in the Multiple treatment, and was chosen over longer datasets to avoid decision fatigue among subjects. The input sequences were chosen randomly, subject to two constraints: (i) in Part 1, they were consistent with only one model (Unique) or more than one model (Multiple), and (ii) in Part 2, they uniquely identified the model (Unique and Multiple).

4 Results

4.1 Extraction

We first examine the rate at which subjects successfully “extract” a model from the dataset they observe. We say a subject “extracts a model” if she makes choices in Part 2 that are consistent with a DGP that can rationalize the dataset provided in Part 1. To operationalize this, we say a DGP (and by extension a model) is “data-consistent” if it generates the outputs in the Part 1 dataset, given the sequence of inputs in that dataset. We say a subject has successfully extracted a “data-consistent model” if a binomial test can reject, at the one-percent level,¹⁶ the hypothesis that the subject made choices in Part 2 that are random relative to a data-consistent model with fewer than four states.¹⁷

We will say a DGP is relatively “complex” if extracting a model of it proves difficult, on average (though we will revisit the appropriateness of this terminology in Section 4.2).

¹⁶In order to pass this threshold for extraction, the subject had to correctly forecast 11 out of the 12 Part 2 outputs.

¹⁷We show in Section 4.4 that this restriction is not binding in our data.

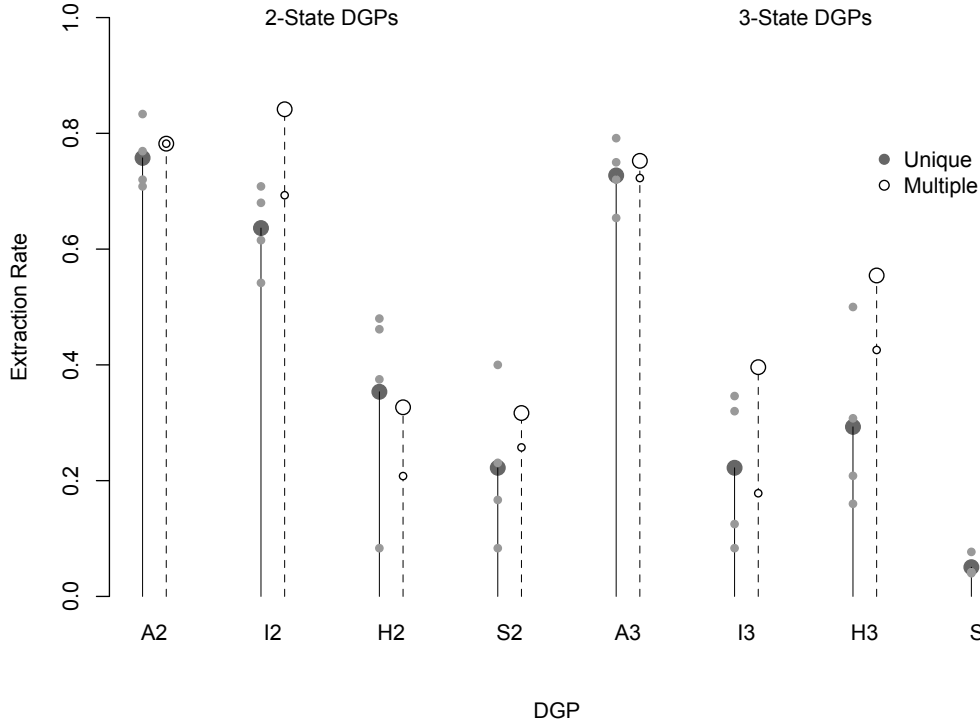


Figure 1: Mean extraction rates by DGP and treatment. *Notes: Large dark dots give the rate at which subjects extracted a model in the Unique treatment. Small dark dots break this rate down by dataset. Large hollow dots give the rate at which subjects extracted any data-consistent model in the Multiple treatment. Small hollow dots give the rate which subjects in the Multiple treatment extracted a model of the (ex post) correct DGP.*

We focus first on the Unique treatment. Figure 1 plots the extraction rate from the Unique treatment in dark gray for each task (indexed by the true DGP for that task). Small dark dots show the extraction rates for each of the four datasets assigned across subjects in the Unique treatment (in order to visualize variation in extraction rates across datasets for each DGP).

The results show that, overall, subjects successfully extract a model less than half (41%) of the time. However, this global mean conceals substantial heterogeneity across DGPs. Extraction rates vary from very high (over 75% for autonomous processes) to almost zero (as low as 5% for switch processes). This heterogeneity is highly significant – we can reject the hypothesis that extraction rates are identical across

DGPs at the 0.0001 level (Chi-squared test).¹⁸

Result 1 *Subjects successfully extract models 41% of the time, and this rate varies across DGPs between 76% (for the “simplest” DGP) and 5% (for the most “complex”).*

We defer detailed discussion of what characteristics of DGPs and datasets drive this heterogeneity across tasks, but note here one striking fact. The number of states in the DGP (which prior work suggests has a first order impact on the costs of *implementing* a rule) seems to have only modest influence over the difficulty of *inferring* a DGP. Subjects extract three-state DGPs 17 percentage points less often than two state DGPs (32% vs. 49%). But this is small compared to the variation in extraction rates within-state. Indeed, the difference between extraction rates for the most- and least-complex models within-state is 60 percentage points, four times larger than the rate difference between states. Clearly complexity of inference is, in large part, driven by factors other than state counts.

Next, we look at how variation across datasets influences extraction rates. Recall, in the Unique treatment that we assigned four different datasets across subjects, each generated by the same DGP but with different input sequences. This variation allows us to ask whether variation in the dataset itself generates variation in complexity, holding the DGP constant. The small gray dots in Figure 1 visualize this variation (for each generating model there are four distinct datasets assigned to subjects, producing four small dots).

Under many DGPs, we see little variation in extraction across datasets. For instance, extraction rates are similar (small gray dots are close together) across datasets under the autonomous DGPs (A2 and A3), the 2-state instruct DGP (I2) and the three state switch DGP (S3). To formally evaluate this, we ran Chi-square models on each DGP-type separately. Doing so, we can reject the null hypothesis that extraction rates are invariant to the dataset at the five percent level only for the hybrid models (H2 and H3); we can reject the same hypothesis at only the ten percent level for I3

¹⁸There is, by contrast, no evidence of order effects in extraction rates. The correlation between task number and extraction rate in the Unique data is -0.026 ($p = 0.459$).

and S2.

Thus, although there is some suggestive evidence that characteristics of the dataset may sometimes influence the complexity of extraction, this effect is very small compared to the influence of the underlying DGP itself. Overall, the mean difference in extraction rates between the least and most complex datasets is 22 percentage points which is less than a third the size of the mean difference between rates for the least and most complex DGPs (70 percentage points). We conclude that:

Result 2 *Variation in datasets has a much smaller influence over the complexity of extraction than variation in the underlying DGP.*

In Section 5 we consider the relative influence of DGPs and datasets on model formation in more depth, comparing formal measures of complexity that depend on the dataset and measures that depend only on the DGP.

4.2 Complexity

So far we have interpreted a low extraction rate as an indication that a DGP is itself “complex” (e.g. difficult or costly). The assumption underlying this interpretation is (i) that there is a latent cost or difficulty ordering across the DGPs themselves and that (ii) subjects differ in their willingness to bear cognitive costs or in their ability to solve problems so that (iii) extraction rates differ on average across tasks.

But, as we discuss in Section 2.3, this “complexity interpretation” may be wrong. For instance, an appealing alternative possibility is that (i) individual subjects have idiosyncratic affinities for inferring some DGPs over others, but that (ii) the rates at which subjects have such affinity varies across DGPs. Call this the “affinity interpretation” of the variation we observe in extraction rates across DGPs. Under this interpretation, subjects specialize – some subjects may, for instance, be good at detecting switch processes while others may be good at detecting instruct processes.

To distinguish between these interpretations (and avoid committing the “ecological fallacy” in interpreting the data),¹⁹ we look for evidence in the Unique treatment of

¹⁹This is the fallacy of supposing that patterns that arise in the aggregate are driven by the same

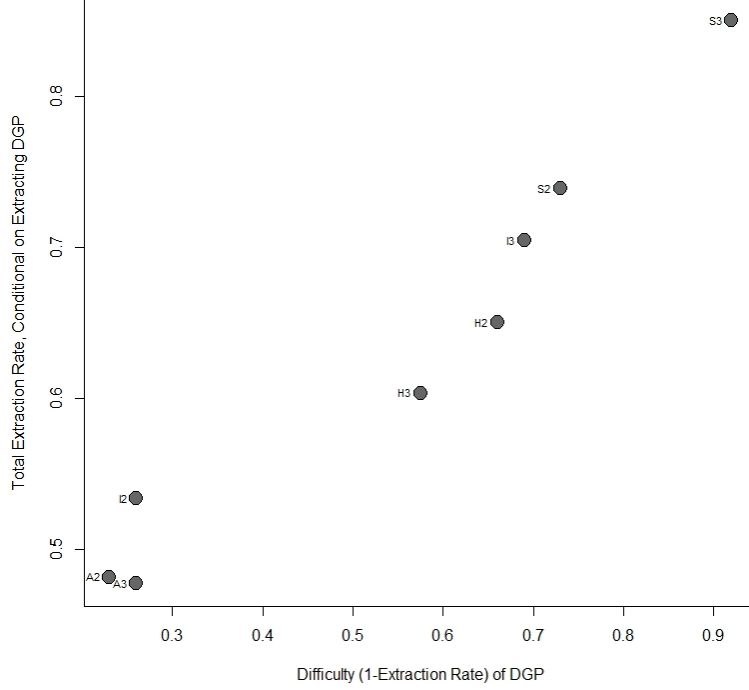


Figure 2: Total rate of extraction conditional on extracting a model of each individual DGP *Notes: The x-axis plots the difficulty (one minus the extraction rate) of extracting each DGP, calculated across subjects. The y-axis plots, for each DGP, the average extraction rate (over all DGPs) for the subset of subjects who successfully extracted that DGP.*

an ordering that is consistent with the complexity but not the affinity interpretation of the extraction data. Under the complexity interpretation, having extracted a relatively “difficult” DGP (a DGP with a low overall extraction rate) will predict extraction of on-average easier DGPs, while having extracted an “easy” DGP (a DGP with a high extraction rate) should be much less predictive of overall extraction rates. In Figure 2, we plot on the x-axis the average difficulty (across subjects) of extracting each of our eight DGPs (one minus the extraction rate pictured in Figure 1), and on the y-axis, for each DGP, the average rate of extraction (over all DGPs) for the subset of subjects who extracted that DGP.

Under a complexity interpretation we expect to observe a strong upward sloping relationship, while under the affinity interpretation we needn’t see any relationship at all. The data clearly shows a strong upwards relationship, indicating a consistent ordering and therefore supporting a complexity interpretation of the extraction data.

patterns at the individual level.

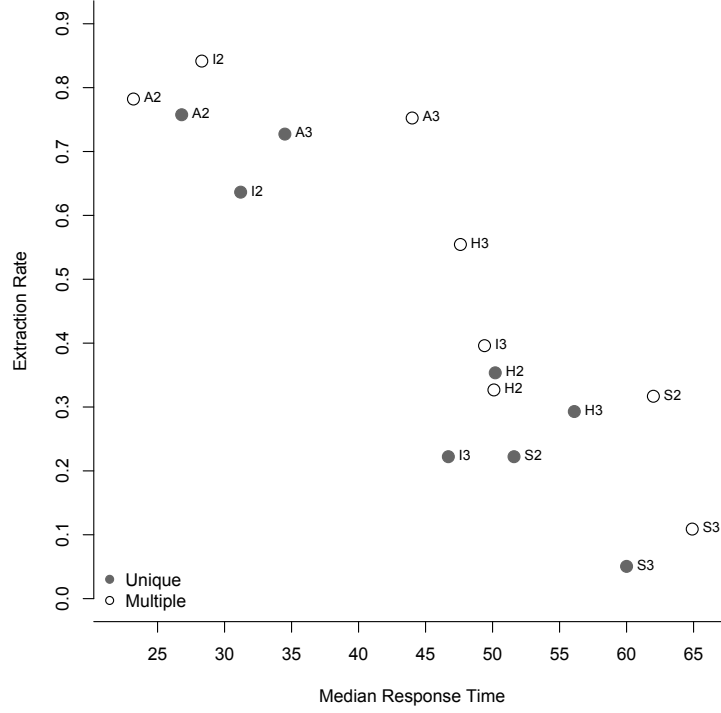


Figure 3: Extraction rates vs. median response time across DGPs.

Result 3 *There is a consistency in extraction behavior across subjects that suggests a complexity ordering across DGPs.*

Figure 3 uses response times across DGPs as additional evidence in favor of our interpretation of the difficulty of extraction as an outgrowth of DGP “complexity.” Response times are widely used as measures of the complexity of decision tasks in economics and psychology (e.g. Rubinstein (2007), Gill and Prowse (2018), Frydman and Jin (2021)). Figure 3 plots the extraction rate for each DGP in each treatment (y-axis) against the median time subjects spend making their full set of decisions in Part 2 of the experiment (whether they successfully extracted or not). The plot reveals a strong linear relationship and an extremely tight correlation ($\rho = -0.88$) between extraction rates and response time across DGPs. Subjects spend over twice as long on DGPs with the smallest extraction rate as they do on DGPs with the

largest extraction rate.²⁰

Result 4 *Subjects spend more time attempting to form models DGPs that they have more difficulty successfully modeling.*

4.3 Selection

Our design allows us to study not only extraction, but also “selection”: how people choose between data-consistent models when more than one model is capable of rationalizing a dataset. In the Unique treatment, we assigned subjects datasets that can be rationalized by only one model (among those describable as fully connected automata with three or fewer states), while in the Multiple treatment we deliberately assigned subjects datasets (generated by the same DGPs as in Unique) that can be rationalized by *multiple* models in this class. Comparing behavior across the two treatments allows us to answer some first questions about the nature of selection.

Figure 1 gives us some first evidence on the impact of selection on subjects’ inference. For each DGP, we plot (using large hollow dots) the rate at which subjects in the Multiple treatment extracted *any* (of the many) models that are data-consistent (that can rationalize the dataset given in Stage 1 in the Multiple treatment). Comparing the hollow dots (extraction rates from Multiple) to the dark gray dots (from Unique), we find that subjects sometimes (but not always) extract at a higher rate when more than one model is available. Having multiple models available to rationalize a dataset overall increases subjects’ ability to extract *some* consistent model from that dataset (Chi-square test, $p < 0.001$). Conducting Chi-square tests at the level of the true DGP, we find that subjects extract significantly more often in Multiple than Unique in three tasks: H3, I2 and I3 (we cannot reject the null at the five-percent level for the other DGPs).

Result 5 *Subjects extract data-consistent models at a higher rate when multiple models are consistent with the data.*

²⁰The direct relationship we observe between response time and mistakes rate is predicted by recently developed “sequential sampling models” of decision making (e.g. Fudenberg, Strack and Strzalecki (2018)).

Figure 1 also plots small hollow dots showing the rate at which subjects in Multiple extracted the, ex post, *true* DGP.²¹ It is clear that the higher rate of extraction in Multiple is *not* driven by higher rates of extraction of the true DGP. Focusing only on extraction of the (ex post) true model, the rates of extraction are identical in the Unique and Multiple cases (41% in both cases). The increase is therefore due to the availability of models other than the, ex post, correct one.

These patterns suggest that the increased rate of model extraction in Multiple occurs due to something akin to search – when more models are available to rationalize the dataset, extraction is more likely to be successful, mirroring comparative statics from any costly search model. This raises the question of whether this process resembles random, undirected search or whether, rather, it is more akin to directed search. In other words, is the increased extraction rate in Multiple a statistical effect of there being more options to stumble upon, or is it driven by the appearance of specific additional models that are easier or more appealing to extract than the true one?

We can test for evidence of undirected search because we deliberately varied the number of rationalizing models available across tasks. For instance for our S2 task, we gave subjects a dataset rationalizable by only 2 models, while for I2 subjects were given one rationalizable by 20. Most tasks featured datasets rationalizable by 3-6 models. Under the hypothesis of random search there should be a strong relationship between the number of available models and the degree to which extraction improved relative to the same DGP under Unique. We find no such relationship. The Spearman correlation coefficient of 0.34 is not significantly different from zero ($p = 0.41$). The Multiple task with the most data-consistent models available (I2) featured only the third largest improvement over Unique and the task with the largest improvement (H3) had only four data-consistent models. We conclude:

Result 6 *The number of models available to rationalize the dataset does not predict improvements in extraction rates in the Multiple treatment. This suggests that characteristics of the specific set of data-consistent models available drives these im-*

²¹Note that failing to extract the true DGP in the Multiple treatment is *not* irrational. Extracting any data-consistent model based on the Part 1 dataset is perfectly rational given the information available to subjects. However, ex post, failing to extract the true DGP will cause the subject to make mistakes in Part 2.

provements, and therefore that subjects engage in some form of directed search.

In Section 5, we structurally estimate the models subjects extract in the Multiple treatment and examine in greater depth what characteristics of models govern selection and this directed search process.

4.4 Articulation

Our measure of successful model formation so far has been based exclusively what we call “extraction”: the rate at which subjects act “as if” they use a model in Part 2 that rationalizes the data observed in Part 1. Extraction requires only *implicit* learning, relying on subjects’ ability to imitate the pattern of the Part 1 data generating process. Do subjects also have an *explicit* understanding of what they have learned, in the sense that they can articulate the model explicitly? This needn’t be true – it is possible that subjects can form a mental model of the true DGP without having the ability to describe what it is they have learned. Indeed, philosophers of science such as Karl Polanyi (Polanyi (2009)) have emphasized that a great part of our knowledge is “tacit knowledge” – knowledge that the learner is unable to describe or articulate in words or symbols and perhaps even has acquired unconsciously. Likewise, psychologists (Reber (1967)) have long emphasized and studied a distinction between “implicit” learning (producing knowledge agents can’t describe and may be unaware of) and “explicit” learning (producing knowledge agents are fully aware of).

Part 3 of each of our tasks allows us to evaluate the rates at which subjects are fully conscious of what they have learned in Part 1 and applied in Part 2. Recall that, at the end of each task, we ask subjects to describe in words the rule the computer used to produce the outputs in Part 1, and we incentivized them to give an accurate and complete description. After data collection, we created a dataset that did not contain the true model or the subject’s choices and, line by line, translated the rule the subject wrote into automata descriptions. We also recorded when subjects admitted they were unable to describe or articulate the rule, when they described processes outside of the true class of DGPs, or when they described overly-elaborate DGPs that could not rationalize the data. Finally, we matched this data back to the

original dataset and coded a subject as having successfully articulated the model if the automaton formed from her words fully matched a data-consistent automaton.

Subjects’ articulation efforts are often successful. For instance, one subject, in trying to describe her model of I3 wrote, “A double bb results in a y. An a or single b is an x” and another trying to articulate S3 wrote “a output stays consistent with the previous output letter. b output goes in the pattern of yxxyxxyxxyxxyx...” Often, however, subjects have difficulty articulating models they’ve successfully extracted. Sometimes subjects admit (often with frustration) that they can’t describe the model they’ve just successfully extracted, like this subject attempting to articulate H3: “I’m sorry, you’re on your own with this one. There’s gonna be more x’s than y’s. I don’t understand the pattern very well.” Sometimes subjects who successfully extracted express doubt that there actually is a coherent DGP (“I’m not sure there is a rule.” or “no idea, seems random”). And sometimes subjects articulate incomplete models, overly elaborate models, or models containing errors. For example, after successfully extracting, one subject attempted to articulate H3 by writing “the letter b input will always produce an x output. The letter a input will produce a y output randomly, but only if there are two or more a inputs in a row. If there is only one a input in a row (for example: abab), then the output for a is x. The y output is random, but can only be triggered if there are two or more a inputs in a row (for example: baab).”

Our main question is how rates of articulation (Part 3) compare to rates of extraction (Part 2) – a comparison that tells us how much of what subjects infer in forming models is implicit. In Figure 4 we plot the mean extraction rate for each model on the x-axis and the corresponding articulation rate on the y-axis. The 45-degree line is shown as a dashed line.²² Virtually all DGPs appear below the 45-degree line, indicating that subjects are generically worse at articulating models than at extracting them. Overall, subjects articulate at 65% of the rate they extract – a rate that falls to just over 50% for 3-state DGPs. There is significant variation in this proportion. For some DGPs, subjects articulate at rates as low as 28% of the rate at which they extract. This gap between extraction and articulation is highly significant

²²In order to filter out subjects who were undermotivated to articulate in Part 3, we remove from the dataset subjects who failed to *ever* articulate a model correctly (the results are similar if we include the entire dataset instead).

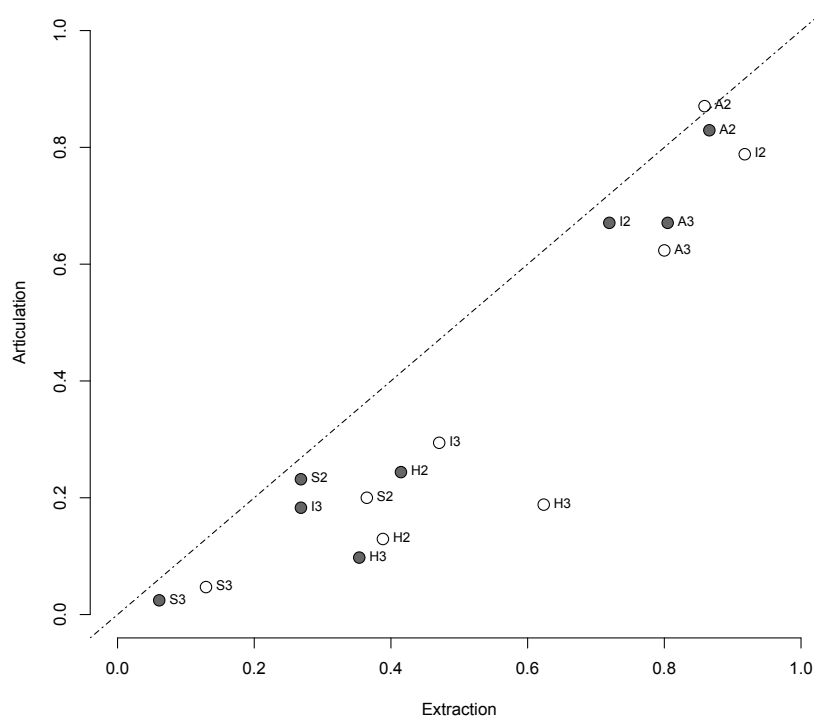


Figure 4: Rates of articulation and extraction for each DGP and treatment. *Notes: The x-axis plots the extraction rate for each DGP/treatment and the y-axis plots the corresponding rate of articulation of a data-consistent model of each DGP.*

($p < 0.001$, paired Wilcoxon test on rates of extraction and articulation).

Thus, overall, and for some DGPs in particular, much of what subjects learn from inference in Part 1 is coded as tacit or implicit knowledge – knowledge that subjects cannot articulate and may even not be explicitly aware of.

Result 7 *Subjects are often unable to explicitly articulate the models they extract. This suggests that a great deal of model formation occurs via implicit rather than explicit learning.*

In Online Appendix A.2, we provide additional details on articulation. Most importantly, this data directly supports and confirms our operationalization of extraction. Recall, in determining whether subjects extract a model, we restrict attention to models/DGPs with fewer than four states. Our analysis of the articulation data suggests

that few subjects (17%) actually form models more complex than this, and almost all of these overly-elaborate models are incomplete and fail to rationalize the Part 1 dataset. Thus, our restriction to 2- and 3-state models in operationalizing extraction is not binding.

5 What Makes a Model Complex?

Our results illustrate significant variation in the rates of extraction, selection and articulation of models across DGPs and (to a much lesser extent) datasets. In this section, we search for a principled explanation for this variation by examining a number of distinct notions of complexity, drawn from computer science, information theory, algorithmic information theory, and economics. Our aim in this search is (i) to provide some insight into why model formation problems are complex for humans and (ii) to examine whether there is a unity in the notions of complexity that drive the somewhat different cognitive acts of extraction, articulation and selection.

5.1 Complexity Notions

We collected a total of 14 notions/measures of complexity that plausibly predict the complexity of forming models and group them into three broad categories, described in the following three subsections. Additional details for Partition Complexity and Sparsity Complexity, as well as the algorithms used to calculate them, are provided in Online Appendix B. Notably, although some of these complexity metrics (s-complexity, t-complexity) are related to the automaton descriptions of DGPs provided in Table 1, most are built on orthogonal descriptions that are unrelated to automata.

5.1.1 Implementation Notions

First, the complexity of forming a model may be linked to characteristics of the DGP itself, and the costs of implementing it. In particular, there are several notions

advanced in economics and computer science of what Oprea (2020) calls “implementation complexity” – the resource burden of implementing the DGP’s automaton for a human or a computational device. We consider three salient measures in this class, which we collectively call **Implementation Measures**:

States (s-complexity): The number of states in the automaton describing the DGP. This is the most popular measure of implementation complexity in the automata literature in economics (e.g. Rubinstein (1986)) and was found to be an important driver of complexity costs in Oprea (2020).

Transitions (t-complexity): The number of transitions in the automaton describing the DGP, hypothesized to generate complexity costs by Banks and Sundaram (1990) and verified as a significant source of complexity costs in Oprea (2020). This is the number of arcs in the automaton diagrams shown in Figure 1.²³

Machine Complexity: The size of the DGP as measured by the minimum number of two-input gates (NOR, NAND, etc.) that would be required to implement it physically in hardware. To the degree human brains work similarly to physical silicon chips, machine complexity describes the cognitive effort required to implement a correct model of the DGP for prediction.

5.1.2 String/sequence Notions

Second, at the other extreme, are measures focused entirely on characteristics of the dataset (rather than the DGP). In particular, these **String/sequence Measures** focus on characteristics of the input or output strings that constitute the dataset subjects observe in Part 1:

Entropy: The Shannon entropy of the input string or output string in the dataset. This is a measure of the amount of information or surprise in a string and is calculated as $-\sum_{u \in U} p(u) \log(p(u))$ (for an input string), where $p(u)$ is the probability of u (approximated using by the relative frequency in the string). We separately consider the entropy of inputs (**Entropy In**) and the entropy of outputs (**Entropy Out**) as

²³Note that our transition complexity measure is different from the one used in Oprea (2020), reflecting our very distinct question. Oprea (2020) focused on a measure of transitions in excess of states to separately identify the contribution of states and transitions to implementation costs.

two different measures.

Kolmogorov Complexity: A measure of the computational resources needed to create a string (e.g. the input or output string), it is measured as the length of the shortest computer program that can create the string. Intuitively, it characterizes a string or sequence as being easier to draw inferences from if it features stronger regularity. It is the core complexity notion in algorithmic information theory, but an important set of theorems (see Li and Vitanyi (2008)) show that it is uncomputable. However, computable approximations to Kolmogorov Complexity exist and we include these as complexity measures. One approximation is *Lempel-Ziv complexity* (Lempel and Ziv (1976)) which we calculate separately for inputs and outputs (**LZ In** and **LZ Out**). Another is *approximate entropy* (Pincus, Gladstone and Ehrenkranz (1991)) (**Approximate In** and **Approximate Out**).^{24,25}

5.1.3 Relational Notions

Finally, we consider measures that focus on both the dataset and the DGP. These **Relational Metrics** are descriptions of the relationship between input and output strings generated, on average, by the DGP:

Mutual Information: A statistical measure of the symmetric informativeness of one random variable about another that is widely used as a measure in rational inattention models Sims (2003)). We approximate the mutual information between the entire set of inputs, $u^T \equiv \{u_t\}_{t=1}^T \in U^T$ and the current output, v_T , by simulating the DGP for 5000 periods to obtain an approximation of the joint distribution, $P(v_t = v, u^T = \tilde{u})$ with $\tilde{u} \in U^T$. Mutual information is then calculated as $I(u^T, v_T) = \sum_{v,u} P(v_t = v, u^T = u) \log \frac{P(v_t=v, u^T=u)}{P(v_t=v)P(u^T=u)}$, where T is the size of the dataset a decision-maker has to make inferences from.

Sensitivity: A deterministic measure of the sensitivity of outputs to the history of

²⁴Approximate entropy was developed for measuring regularity or irregularity in heartbeats. Lempel-Ziv complexity was developed by computer scientists as a practical alternative to Kolmogorov complexity that forms the basis for Lempel-Ziv lossless compression.

²⁵More promising might be conditional Kolmogorov complexity which is measured as the length of the shortest computer program that can create the output string *from the given input string*. Unfortunately this too is uncomputable and we are not aware of any methods for approximating it.

inputs in DGPs, proposed by Lipman and Srivastava (1990). Specifically, the measure counts the number of alterations one can make to the history of inputs that would change the output at the end of that history, summed over all histories.²⁶ Our main measure calculates this at the DGP level, e.g. over all possible histories. We will call this notion simply **sensitivity**. We also calculate a version fitted to the specific Part 1 dataset subjects observed, which we call **string sensitivity**.

Partition Complexity: A measure of the information processing of the dataset required to correctly implement the DGPs output sequence, adapted from a framework for modeling bounded rationality suggested by Lipman (1995). Our adaptation counts the number of unique elements in the coarsest partition of the set of subhistories in the dataset that a decision maker must use to correctly forecast the outputs required by the DGP. The idea is that models that require an agent to distinguish between more distinct patterns in the dataset (more unique partition elements) are more difficult to identify and model. For instance, a model of DGP I2 only requires the agent to partition the dataset into two kinds of elements to forecast outputs: (i) output x occurs if the input is a and (ii) output y occurs if the input is b. But DGP S2 (which has the same number of states and transitions) requires the subject to distinguish between four patterns: (i) output x occurs if the input is a and previous output is x; (ii) output y occurs if the input is a and the previous output is y; (iii) output x occurs if the input is b and the previous output is y; (iv) output y occurs if the input is b and the previous output is x. Thus S2 is more “partition complex” than I2.

Sparsity Complexity: A measure of the sparsity of a DGP – the amount of the dataset that can be safely ignored when applying a correct model of the DGP to predict outputs. The measure is an adaptation of ideas from Gabaix (2014) which postulates that sparser models are less difficult or costly to apply. To operationalize, we say that a model/DGP is “less complex in the sense of sparsity” if fewer elements (inputs or outputs) in the dataset need to be attended to in order to predict outputs. For instance S2 is more “sparsity complex” than I2 because an agent must attend to both the input and previous output in S2, but only the input in I2. This measure is

²⁶We weight all histories equally and only consider alterations in which one input changes. Lipman and Srivastava (1990) consider a more general model.

closely related to partition complexity and, like partition complexity, is a measure of the amount of information in a dataset that must be processed to apply the model.²⁷

5.2 The Complexity of Extraction and Articulation

In Figure 5 we plot, for each complexity notion, the absolute value of the estimated correlation between the complexity metric²⁸ and (i) extraction (the y-axis) and (ii) articulation (the x-axis), using data from the Unique treatment.²⁹ To complement, in Online Appendix A, we show scatterplots of mean extraction rates against each of these 14 measures.

We make three observations.

First, there is an extremely strong relationship between the values of complexity metrics in explaining each of extraction and articulation. High correlation of a metric with one measure of model formation tends to be accompanied by high correlation with the other.

Second, String/sequence measures like entropy and approximations to Kolmogorov Complexity (LZ In/Out and Approx In/Out) tend to be particularly bad predictors

²⁷One additional notion we hoped (but were unable) to use is “perceptrons,” which have been successfully used in modeling bounded rationality in economics (e.g. Rubinstein (1993)). Perceptrons are binary classifiers that produce an output classification by calculating the weighted sum of a set of inputs and comparing it to a threshold. The number of inputs (a perceptron’s order) can be taken to be a measure of its complexity (Rubinstein (1993)). Thus, a natural measure of the complexity of a model is the minimal order of the perceptron that can implement it. However, in attempting to adapt it we ran into a well known problem of perceptrons - they can only implement linearly separable functions. If the function corresponding to a model is not linearly separable, no perceptron exists. For example, the output of S2 corresponds to the XOR of the previous input and output, which is the quintessential example of a function that is not implementable by a perceptron. This non-implementability prevents us from being able to use perceptrons to construct a complexity measure.

²⁸Because it isn’t clear how to cardinality interpret many of these complexity notions, we focus entirely on their ordinal explanatory value by normalizing all of our notions to their ordinal ranks. For instance, although our design varies the number of states between 2 and 3 we normalize them to vary between 1 and 2, by rank.

²⁹For this calculation we use Goodman and Kruskal’s gamma, which, like Spearman correlation and Kendall’s tau, measures the rank correlation between two variables. The main advantage of Goodman and Kruskal’s gamma is that it allows for more general relationships in the data than Spearman and does not penalize ties like Kendall’s tau (and thus does not penalize the metrics that only have a few values such as States).

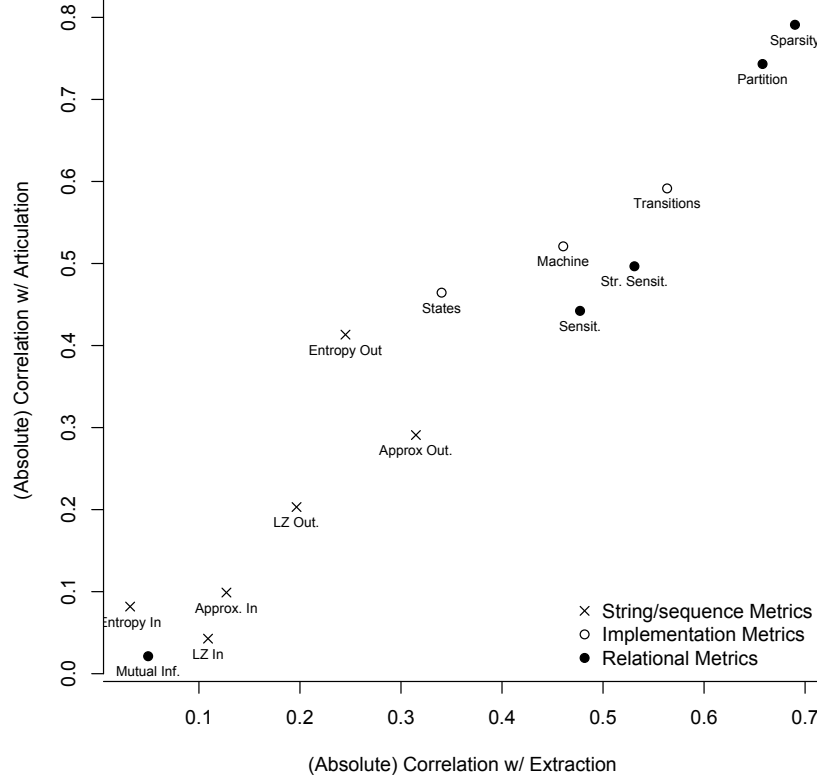


Figure 5: Absolute value of the correlation between complexity metrics and articulation/extraction. *Notes: A separate dot is presented for each complexity notion/metric. The absolute value of Goodman/Kruskal gamma correlation is plotted for extraction on the x-axis and for articulation on the y-axis.*

of model formation. This mirrors and reinforces the observation from our analysis of extraction that DGPs are a much stronger driver of variation than are datasets. Indeed, Implementation metrics (States, Computational Complexity and Transitions) that are rooted in the DGP are all better than the String/sequence measures.³⁰

³⁰To give string-based metrics their best shot, we also calculated correlation coefficients within-DGP, effectively using them only to explain the small and somewhat inconsistent variation we observe across datasets. Thus we calculated our correlation coefficients between our string-based complexity metrics (Entropy In, Entropy Out, Lz In, LZ Out, Approx. In, Approx. Out and String Sensitivity) metrics and extraction/articulation *within* DGP for each DGP. Although for some DGPs correlations are quite strong (rising to as high as 0.79 for some DGP/metric combinations), we find little systematic evidence that string-based metrics do much better within- than between-DGP. Averaging within-correlations across DGPs, we find they are never higher than 0.46 for any metric. We conclude that the variation we observe across datasets is either due to sampling error or to perceptual difficulties presented by some datasets that are ill-captured by our complexity metrics.

Third, there is significant variability in explanatory power across Relational measures. The overall worst measure (Mutual Information) is a relational measure, but so are the very best (Partition Complexity and Sparsity Complexity).

Most importantly, we find that two closely related metrics that describe the *information processing* required to apply a model for prediction – Partition and Sparsity – do the best job of explaining what makes model formation complex, with impressive correlations of 2/3 or more.

Result 8 *Measures of the amount of information processing of the dataset required strongly predict subjects’ abilities to extract and articulate the model.*

5.3 Complexity and Selection

Next, we use these same complexity notions to return to the question of what drives selection of models when multiple models are available. Recall from Section 4.3 that (i) subjects tend to extract at a greater rate when multiple models can rationalize the dataset, but (ii) the sheer number of available models is a poor predictor of this improvement in extraction across tasks. Our conclusion from this analysis is that selection is likely driven by something akin to directed search, with subjects preferentially selecting some models over others based on some characteristics of the underlying DGP.

In this section we consider the possibility that something like a behavioral analogue of Occam’s Razor applies – that subjects preferentially select models that are simpler under some notion of simplicity. To do this, we estimate a structural model described in Online Appendix A.3 on Part 2 choice data to determine which of the available data-consistent models subjects choose in the Multiple treatment. Then, we test the hypothesis that subjects are drawn to simplicity (and examine what sort of simplicity they are drawn to) by examining how frequently subjects select models of the lowest complexity under each of our complexity notions.

We estimate subjects’ models from Part 2 data using a standard random utility framework, specifically a finite mixture model adapted from the Strategy Frequency

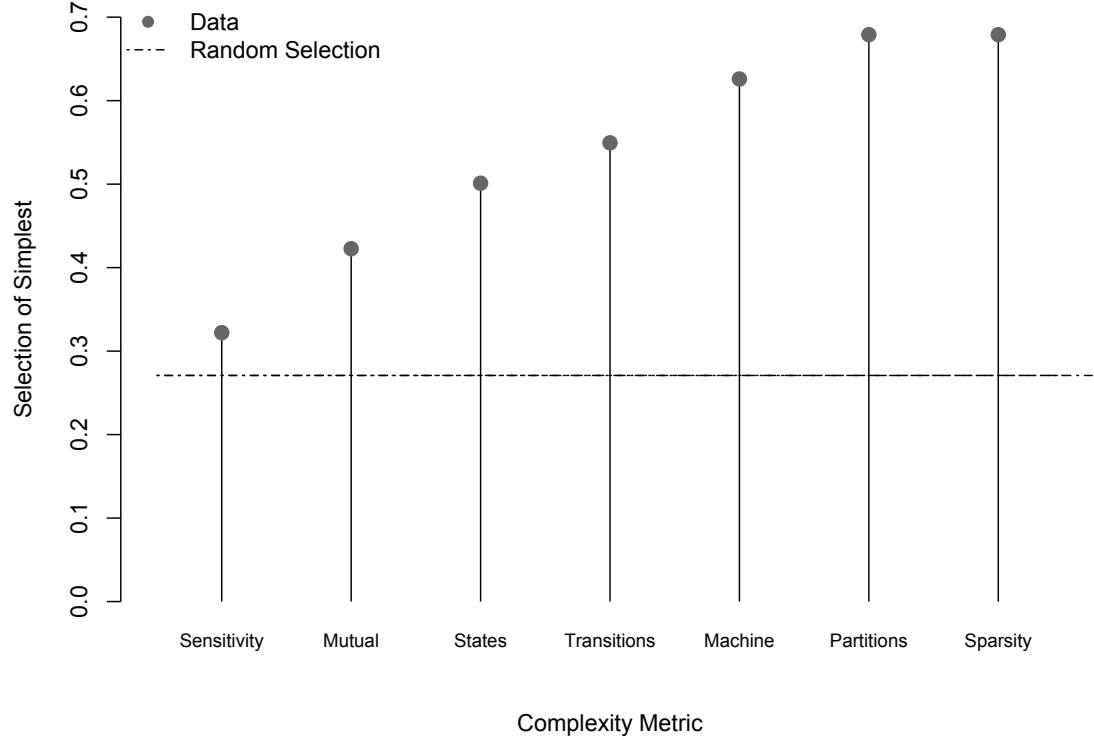


Figure 6: Proportion of subjects who select the simplest available model under various notions of simplicity. *Notes: Each data point corresponds to a different notion of complexity. The y-axis shows the proportion of subjects who selected the most commonly selected of the simplest available model under each notion.*

Estimation Method (SFEM) developed by Dal Bó and Fréchette (2011) to estimate strategies used in repeated games. This procedure give us individual level estimates of the weight each subject places on each of the models (of fewer than four states) available to rationalize the datasets observed in Part 1. We use these subject/task level estimates to summarize which models subjects are drawn to in the following way. For each task in the Multiple treatment, we identify the *simplest* (least complex) available data-consistent model *under each complexity notion*. Because this exercise concerns models and not datasets, we use only complexity notions that relate purely to characteristics of the DGP: Sensitivity, Mutual Information, States, Transitions, Machine Complexity, Partition Complexity and Sparsity Complexity.

Next, we calculate the fraction of subjects who clearly selected this simplest model

under each metric. For this, we consider only subject/task combinations in which there is clear evidence that the subject extracted a particular model: subject-tasks in which the subject put a weight greater than 0.95 on one of the models according to the structural estimates. These cases make up 61% of the data. Our question is whether and under what complexity notions these subjects put greater than random weight on the simplest available model.

Figure 6 plots the rate at which subjects selected the simplest model as solid lines with dots for each of our complexity notions. A dotted horizontal line shows the rate at which subjects would select the simplest model if they were choosing randomly.

The results show that under *any* complexity notion, subjects select the simplest model more often than random. Moreover, as with extraction and articulation, selection is most strongly explained by the information processing notions (Partition Complexity and Sparsity Complexity). Subjects choose the simplest model under these metrics nearly 2/3 of the time (three times more often than would be expected from random selection), suggesting that subjects are not only particularly good at extracting models with low information processing requirements, they are also strongly attracted to (or good at discovering) these same types of models. These findings therefore suggest a notable unity in the determinants of model formation.

Result 9 *Subjects have a “bias” towards simplicity in selecting models when multiple models are consistent with data. As with extraction and articulation, the strongest predictors of selection are information processing notions of complexity (Partition and Sparsity Complexity).*

Finally, we emphasize that the preceding analysis is deliberately conservative and should be interpreted as a lower bound estimate of the weight subjects attach to simple models in selection. In calculating this weight, we are sometimes confronted with the fact that there is no unique simplest model – often multiple available models are equally simple. In these cases, we said a subject picked the simplest model only if she picked the model most often picked among the set of simplest models. A less conservative approach would have been to add up the weight across all models in the simplest set – an approach that mechanically overweights complexity notions that

make fewer distinctions across models (that have fewer distinct values). Using this alternative method of classifying, subjects select the Partition-simplest and Sparsity-simplest model over 90% of the time.

6 Discussion

In the last half century, economists have learned a great deal about the difficulties people face in making inferences. Most of the hundreds of empirical investigations on inference in economics center on the difficulties people face in extracting simple probabilities from stochastic reservoirs of information (see e.g. Benjamin (2019)). But in economic life, many of our most important inference tasks involve extracting, not simple probabilities, but rather *algorithms* from data. The data generating processes we have to understand in order to forecast and choose, after all, often have rule-like structures ranging from very simple cyclical regularities to elaborate causal relationships. Forming models of such processes to inform beliefs and guide behavior is difficult, not due to stochasticity (as in simple probabilistic inference) but rather to the sheer computational difficulty of extracting patterns from data and representing these patterns in form suitable for prediction.

Our experiment takes some first steps towards understanding how humans form models of algorithms and what makes this type of inference complex (costly or difficult). Because these are first steps, we radically simplify the problem on multiple dimensions. In order to isolate the effects of computational difficulty we study perfectly deterministic data generating processes. Likewise, we study processes in which the direction of causality is clear to decision makers, *ex ante*. Finally we study the simplest possible examples from this already simple class of problems: DGPs describable by two and three state finite automata with datasets consisting of binary inputs and outputs. All of these simplifications can be relaxed by applying our methods in future work.

Despite these simplifications, we find evidence that forming accurate models of algorithms is extremely difficult for decision makers. Although our data generating processes (DGPs) are rudimentary, deterministic and causally clear, subjects are able

to form accurate models of them less than half of the time. This difficulty varies dramatically across DGPs, with subjects accurately forming models of some DGPs upward of 80% of the time and of others less than 10% of the time.

Importantly, we find significant structure to this complexity and are able to make progress in identifying its source. We find evidence of a consistent (across subjects) complexity ordering across DGPs: subjects who successfully form on-average difficult models typically are also able to form on-average easier models. Searching over a wide array of formal complexity notions, we find that this complexity ordering is best organized by a pair of closely related metrics of the information processing required to form a model of the DGP. One of these notions (“Partition Complexity”), adapted from Lipman (1995), counts the number of distinct patterns a modeler has to identify in a dataset in order to make future predictions. Closely related is “Sparsity Complexity”, adapted from Gabaix (2014), which counts the number of distinct input/output elements in the dataset a modeler has to attend to in order to predict. Both of these measures of information processing are highly correlated with rates of extraction across modeling tasks.

These findings put some important context on the often remarked-upon fact that decision makers seem eager to borrow and re-use models from other people and contexts rather than form new models of their own. For instance, human beings seem eager to borrow models from their social environment, adopting narratives and explanatory schema provided by family members, peers, mentors and political and religious affiliates. This social and cultural contagion of models is understandable in light of the difficulty people seem to have in forming even simple models in our data on their own. Likewise decision makers often seem to use thin analogies between circumstances as bases for re-using mental models from very different contexts, even when these second-hand models aren’t perfectly appropriate fits. It is plausible that one of the roots of the widespread use of heuristics in decision-making observed in behavioral economics and psychology, is the sheer cost and difficulty people face when attempting to form new models of novel problems. The fact that this complexity has significant structure suggests we may be able to model such phenomena and predict the contexts in which social contagion and heuristic re-use of models will tend to arise.

In addition to providing evidence on how well subjects “extract” models to guide behavior, our experiment also allows us to examine how well they can consciously “articulate” what they have learned. We find that subjects are often able to extract models that they are not able to afterwards accurately describe. This is evidence that at least some of the learning involved in modeling complex phenomena is implicit and perhaps even unconscious. Nonetheless, we find that the same formal notions of complexity – ones that measure the information processing required to recognize an algorithm – best predict articulation just as with extraction.

Finally, our design allows us to understand not only what makes an algorithm difficult to accurately model, but also what types of models people are preferentially drawn to when explaining data. When more than one model is available to rationalize data in our experiment, we find a kind of behavioral analogue to Occam’s Razor: subjects tend to preferentially select simple models to explain data. What’s more, this model selection is governed by the same complexity notions – Partition and Sparsity Complexity – that govern the difficulty of extracting and articulating models, with subjects preferentially selecting models that require less information processing. The fact that selection is “biased” in this way may be important for understanding bounded rationality and its persistence. A growing theoretical literature on misspecified models in economics shows that mistaken models are often resistant to learning, leading to persistent and even reinforcing mistakes in behavior and beliefs. Our findings point to structure in what types of models people may tend to get stuck in and therefore may be useful for building future models of bounded rationality.

Our results reveal a fundamental unity in the complexity of forming models. Extraction, articulation and selection of models are all best organized by the same metrics of complexity, rooted in the information processing required to form and deploy a model. Future work should test the limits of this unity. Does the same type of complexity organize the formation of models of DGPs complicated by stochasticity (e.g. non-degenerate Markov chains) or confounded by causal uncertainty (e.g. DAGs)? Conducting experiments that extend our methods to these richer settings seems an important next step in this agenda.

More generally, our findings suggest that the difficulty of finding a solution to a prob-

lem (e.g. inferring a correct model) is directly linked to the information processing required to later deploy that solution in decision-making (e.g. to *use* a model to make predictions). Is this true more generally of the wide range of judgemental and optimization failures documented in behavioral and experimental economics? Answering this question will require continued development of models that characterize the information processing required of rational behavior, and experiments testing the organizing power of resulting metrics in explaining deviations from optimal choice in a much wider range of settings. Exploring these questions both theoretically and empirically seems a promising path towards learning how to build robust, predictive models of boundedly rational decision makers.

References

- Abeler, Johannes, and Simon Jäger.** 2015. “Complex Tax Incentives.” *American Economic Journal: Economic Policy*, 7(3): 1–28.
- Abreu, Dilip, and Ariel Rubinstein.** 1988. “The Structure of Nash Equilibrium in Repeated Games with Finite Automata.” *Econometrica*, 56(6): 1259–1281.
- Aragones, Enriqueta, Itzhak Gilboa, Andrew Postlewaite, and David Schmeidler.** 2005. “Fact-Free Learning.” *American Economic Review*, 95(5): 1355–1368.
- Banks, Jeffrey S, and Rangarajan K Sundaram.** 1990. “Repeated Games, Finite Automata, and Complexity.” *Games and Economic Behavior*, 2(2): 97–117.
- Benjamin, Daniel J.** 2019. “Errors in probabilistic reasoning and judgment biases.” *Handbook of Behavioral Economics: Applications and Foundations 1, 2*: 69–186.
- Ben-Porath, Elchanan.** 1990. “The Complexity of Computing a Best Response Automaton in Repeated Games with Mixed Strategies.” *Games and Economic Behavior*, 2(1): 1–12.
- Bohren, Aislinn, and Daniel Hauser.** 2021. “Learning with Heterogeneous Misspecified Models: Characterization and Robustness.”

- Byrne, Ruth, and P.N. Johnson-Laird.** 1989. “Spatial Reasoning.” *Journal of Memory and Language*, 28: 564–575.
- Caplin, Andrew.** 2016. “Measuring and Modeling Attention.” *Annual Review of Economics*, 8: 379–403.
- Chatterjee, Kalyan, and Hamid Sabourian.** 2009. “Game Theory and Strategic Complexity.” *Encyclopedia of Complexity and Systems Science*, 4098–4114.
- Dal Bó, Pedro, and Guillaume R Fréchette.** 2011. “The Evolution of Cooperation in Infinitely Repeated Games: Experimental Evidence.” *American Economic Review*, 101(1): 411–429.
- Dean, Mark, and Nathaniel Neligh.** 2019. “Experimental Tests of Rational Inattention.”
- Eliaz, Kfir, and Ran Spiegler.** 2020. “A Model of Competing Narratives.” *American Economic Review*, 110(12): 3786–3816.
- Enke, Benjamin.** 2020. “What You See Is All There Is.” *Quarterly Journal of Economics*, 135(3): 1363–1398.
- Esponda, Ignacio, and Demian Pouzo.** 2016. “Berk–Nash Equilibrium: A Framework for Modeling Agents With Misspecified Models.” *Econometrica*, 84(3): 1093–1130.
- Esponda, Ignacio, Emanuel Vespa, and Sevgi Yuksel.** 2021. “Mental Models and Learning: The Case of Base-Rate Neglect.”
- Frydman, Cary, and Lawrence Jin.** 2021. “Efficient Coding and Risky Choice.” *The Quarterly Journal of Economics*.
- Fudenberg, Drew, Gleb Romanyuk, and Philipp Strack.** 2017. “Active Learning with a Misspecified Prior.” *Theoretical Economics*, 12(3): 1155–1189.
- Fudenberg, Drew, Philipp Strack, and Tomasz Strzalecki.** 2018. “Speed, Accuracy, and the Optimal Timing of Choices.” *American Economic Review*, 108(12): 3651–3684.

- Gabaix, Xavier.** 2014. “A Sparsity-Based Model of Bounded Rationality.” *The Quarterly Journal of Economics*, 129(4): 1661–1710.
- Gagnon-Bartsch, Tristan, Matthew Rabin, and Joshua Schwarzstein.** 2021. “Channeled Attention and Stable Errors.”
- Gilboa, Itzhak.** 1988. “The Complexity of Computing Best-Response Automata in Repeated Games.” *Journal of Economic Theory*, 45(2): 342–352.
- Gill, David, and Victoria L Prowse.** 2018. “Strategic Complexity and the Value of Thinking.”
- Graeber, Thomas.** 2021. “Inattentive Inference.”
- Handel, Benjamin, and Joshua Schwartzstein.** 2018. “Frictions or Mental Gaps: What’s Behind the Information We (Don’t) Use and When Do We Care?” *The Journal of Economic Perspectives*, 32(1): 155–178.
- Hanna, Rema, Sendhil Mullainathan, and Joshua Schwartzstein.** 2014. “Learning Through Noticing: Theory and Evidence from a Field Experiment.” *The Quarterly Journal of Economics*, 129(3): 1311–1353.
- Heidhues, Paul, Botond Koszegi, and Philipp Strack.** 2018. “Unrealistic Expectations and Misguided Learning.” *Econometrica*, 86(4): 1159–1214.
- Jimenez, Luis, Joaquin Vaquero, and Juan Lupianez.** 2006. “Qualitative Differences Between Implicit and Explicit Sequence Learning.” *Journal of Experimental Psychology Learning, Memory, and Cognition*, 32(3): 475–490.
- Jin, Giner, Michael Luca, and Daniel Martin.** 2021. “Complex Disclosures.” *Management Science*, forthcoming.
- Kalai, Ehud, and William Stanford.** 1988. “Finite Rationality and Interpersonal Complexity in Repeated Games.” *Econometrica*, 56(2): 397–410.
- Lempel, Abraham, and Jacob Ziv.** 1976. “On the Complexity of Finite Sequences.” *IEEE Transactions on Information Theory*, 22(1): 75–81.

- Li, Ming, and Paul M.B. Vitanyi.** 2008. *An Introduction to Kolmogorov Complexity and its Applications*. . 3 ed., Springer Publishing Company, Incorporated.
- Lipman, Barton L.** 1995. “Information Processing and Bounded Rationality: A Survey.” *Canadian Journal of Economics*, 28(1): 42–67.
- Lipman, Barton L, and Sanjay Srivastava.** 1990. “Informational Requirements and Strategic Complexity in Repeated Games.” *Games and Economic Behavior*, 2(3): 273–290.
- Mathews, Robert, and Ray Buss.** 1989. “Role of Implicit and Explicit Processes in Learning from Examples: A Synergistic Effect.” *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 15(6): 1083–1100.
- Murawski, Carsten, and Peter Bossaerts.** 2016. “How Humans Solve Complex Problems: The Case of the Knapsack Problem.” *Scientific reports*, 6: 34851.
- Oprea, Ryan D.** 2020. “What Makes a Rule Complex.” *American Economic Review*, 110(12): 3913–3951.
- Pincus, S.M., I.M. Gladstone, and R.A. Ehrenkranz.** 1991. “A Regularity Statistic for Medical Data Analysis.” *Journal of Clinical Monitoring and Computing*, 7(4): 335–345.
- Polanyi, M.** 2009. *The Tacit Dimension*. University of Chicago Press.
- Reber, Arthur.** 1967. “Implicit Learning of Artificial Grammars.” *Journal of Verbal Learning and Verbal Behavior*, 6: 855–863.
- Rubinstein, Ariel.** 1986. “Finite Automata Play the Repeated Prisoner’s Dilemma.” *Journal of Economic Theory*, 39(1): 83–96.
- Rubinstein, Ariel.** 1993. “On Price Recognition and Computational Complexity in a Monopolistic Model.” *Journal of Political Economy*, 101(3): 473–484.
- Rubinstein, Ariel.** 2007. *Instinctive and Cognitive Reasoning: A Study of Response Times*. Vol. 117.

- Schotter, Andrew, and Barry Sopher.** 2003. “Social Learning and Coordination Conventions in Intergenerational Games: An Experimental Study.” *Journal of Political Economy*, 111(3): 498–529.
- Schwartzstein, Joshua, and Adi Sunderam.** 2021. “Using Models to Persuade.” *American Economic Review*, 111(1): 276–323.
- Shanks, David, Theresa Johnstone, and Leo Staggs.** 1997. “Abstraction Processes in Artificial Grammar Learning.” *The Quarterly Journal of Experimental Psychology*, 50A(1): 216–252.
- Sims, Christopher A.** 2003. “Implications of Rational Inattention.” *Journal of Monetary Economics*, 50(3): 665–690.
- Spiegler, Ran.** 2016. “Bayesian Networks and Boundedly Rational Expectations.” *Quarterly Journal of Economics*, 131(3): 1243–1290.

Online Appendices

A Online Appendix A: Additional Data Analysis

A.1 Extraction and Complexity Notions

Figure A.1 provides a detailed view of the relationship between each of the complexity notions considered in the paper and extraction rates. It includes a separate panel for each complexity notion, plotting the rank value of the complexity metric on the x-axis and the mean extraction rate for each DGP on the y-axis. A dashed line plots a LOESS fit to the raw data. A corresponding figure for articulation rates is provided in Figure A.2, revealing similar patterns.

A.2 Articulation Details

To better understand failures to articulate and to check whether or not subjects are extracting more complex models than we consider when declaring that a subject “extracted” a model, in Figure A.3 we categorize and plot the types of responses subjects give in Part 3 of the experiment. Overall, we find that most subjects make serious efforts to articulate: nearly 85% of subjects give at least partly coherent advice on at least one task. Thus, the vast majority subjects are willing and capable of articulating models in principle. The left-most category in the Figure (“Any Articulation”) shows that subjects at least partially articulate the true model just under 40% of the time. The next most common response is the subject saying explicitly that she does not know or cannot articulate what the true model is (23%). The remaining subjects provide descriptions of the model that are not in the form of a rule (“Not a Rule”, 4%), too ambiguously described to form an automaton description (“Ambiguous”, 15%), present an overly-complicated description of the model that fails to rationalize the Part 1 dataset (“Too Complex”, 17%) or give a heuristic prescription to simply choose the same action throughout the task (“Just X/Y”).

Hollow dots above each of these categories reveal the rates at which subjects who

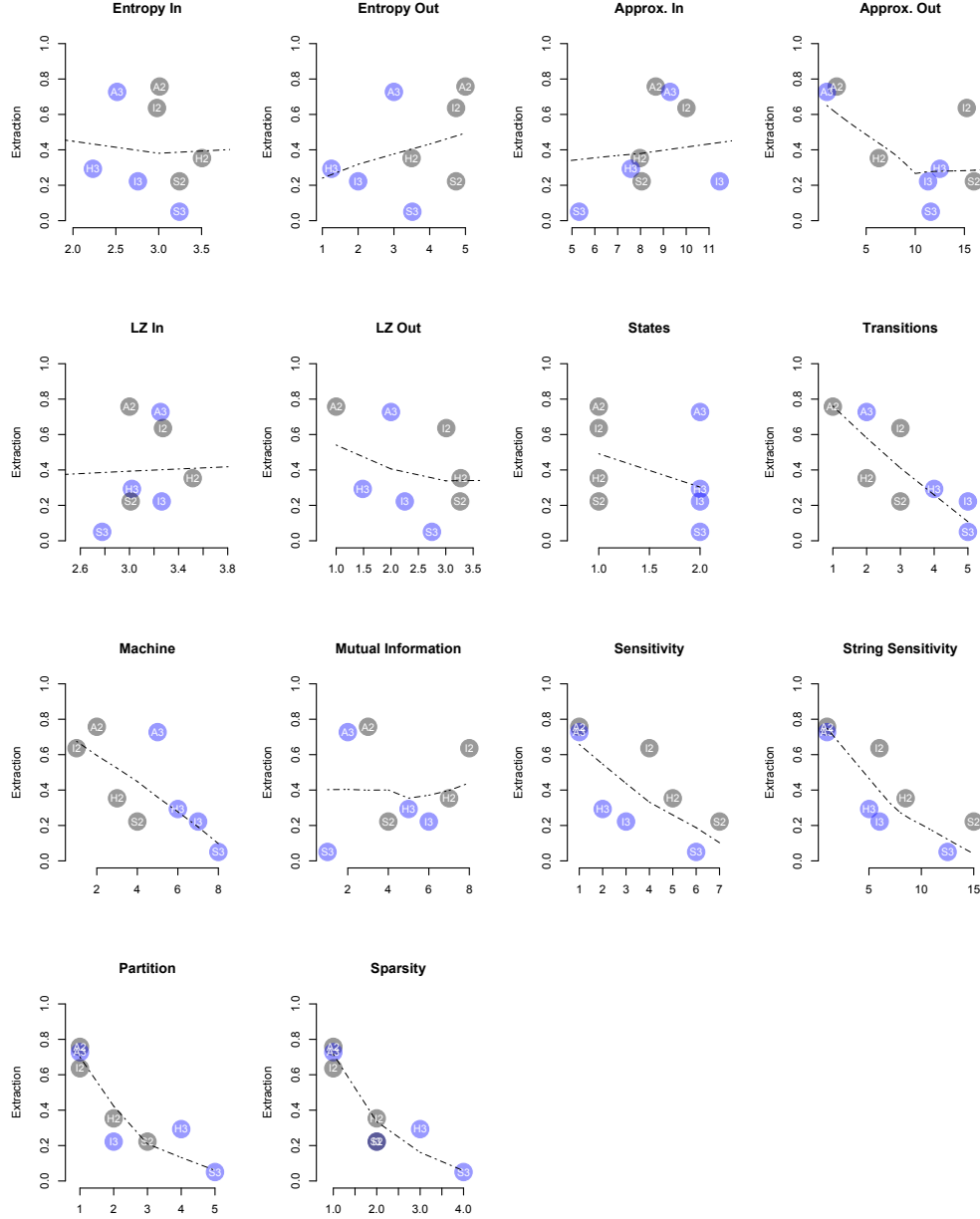


Figure A.1: Average extraction rate as a function of complexity metrics. *Notes: Each panel corresponds to one of our 14 complexity metrics. Each pictures the relationship between the rank ordering of the complexity metric (x-axis) and the mean rate of extraction (y-axis) across DGPs. Note that for the Dataset/string measures these averages aggregate over multiple datasets, producing fractional values for the rank ordering.*

provided each description type managed to extract the true model in their choices in Part 2. Subjects who managed to articulate were by far the most likely to extract. But there is also significant extraction for subjects who failed to articulate the true

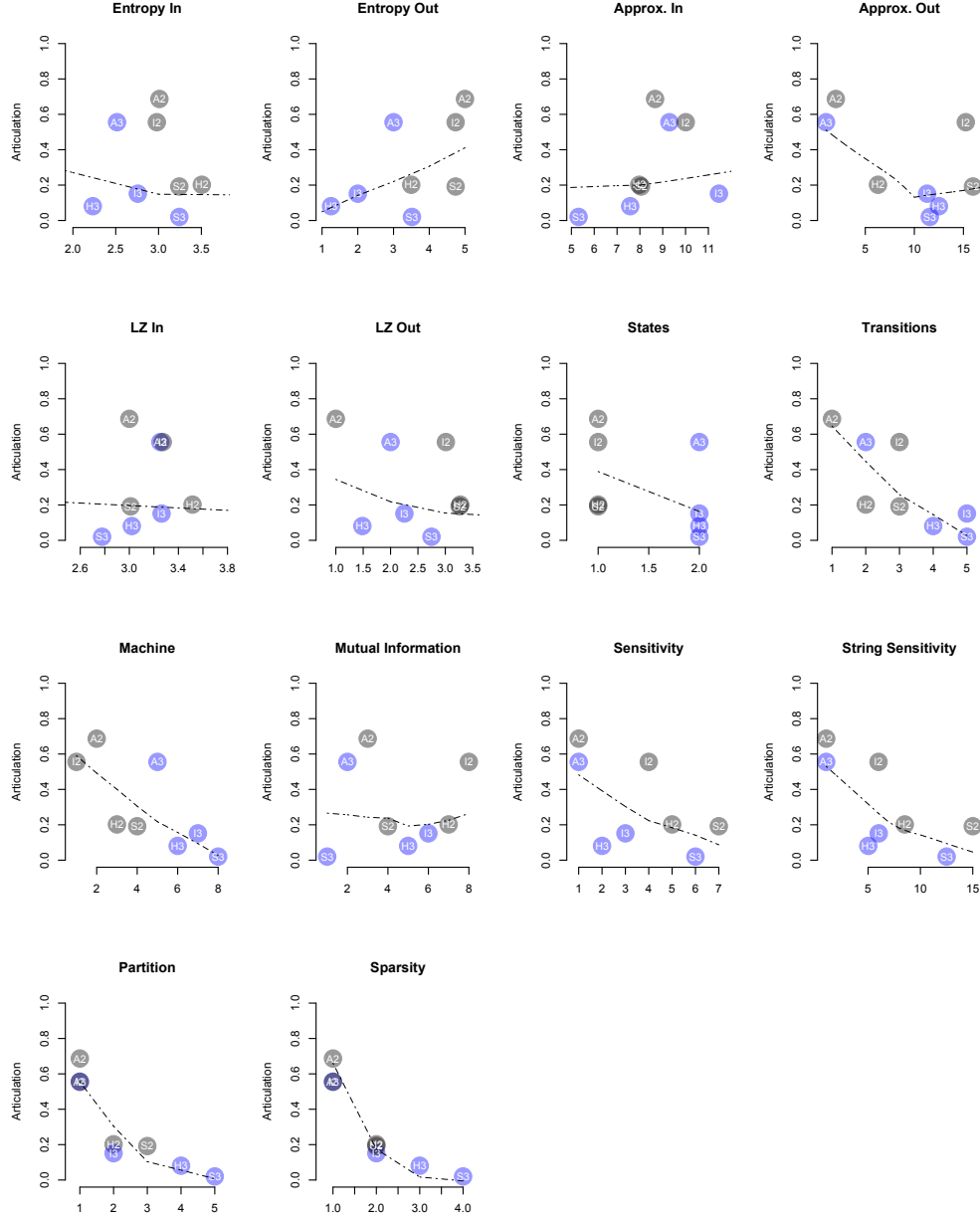


Figure A.2: Average articulation rate as a function of complexity metrics. *Notes: Each panel corresponds to one of our 14 complexity metrics. Each pictures the relationship between the rank ordering of the complexity metric (x-axis) and the mean rate of articulation (y-axis) across DGPs. Note that for the Dataset/string measures these averages aggregate over multiple datasets, producing fractional values for the rank ordering.*

model. Once again, this suggests that many subjects who are unable or unwilling to explicitly articulate the true model, nonetheless were able to implicitly understand the model well enough to use it to guide their actions.

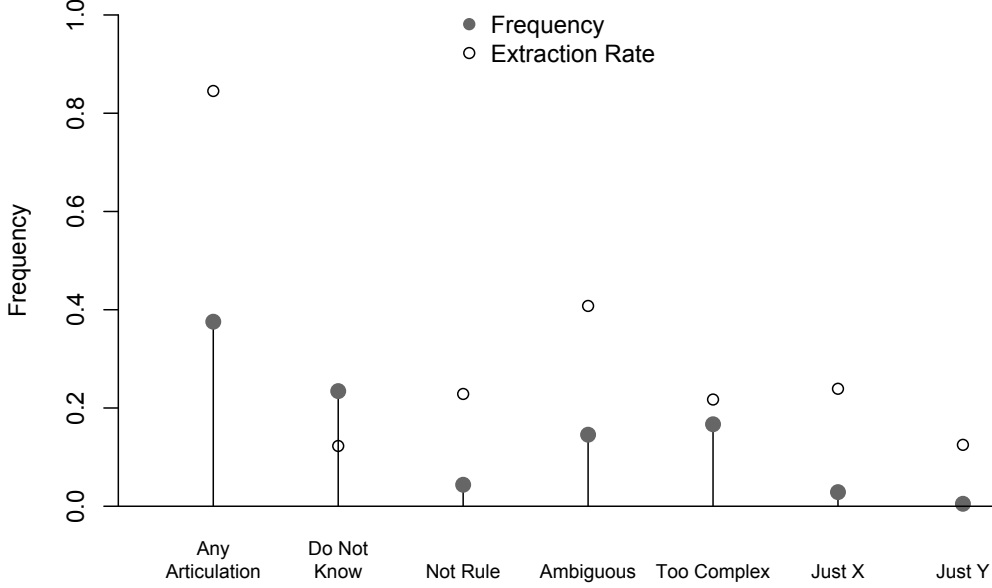


Figure A.3: Breakdown of types of responses given in Part 3. *Notes: Dark dots show the frequency of each type of response. Hollow dots show the extraction rate for subjects who responded in each way.*

A.3 Structural Model

In order to understand what types of models subjects tend to extract when multiple are available, we estimate a structural model of model-selection adapted from structural models of strategy choice from the experimental repeated games literature, specifically the Strategy Frequency Estimation Method (SFEM) developed by Dal Bó and Fréchet (2011). For each task, we estimate the proportion of subjects extracting each model. As in SFEM, we allow for random deviations between the actual output of the model and a subject’s guess: $y_{ist}(m^k) = I \{m_{ist}(m^k) + \gamma \varepsilon_{ist} \geq 0\}$, where $I()$ is the indicator function. y_{ist} indicates the observed guess by subject i , on task s , in period t , where guesses of “x” are coded as 0, and guesses of “y” are coded as 1. $m_{ist}(m^k)$ is the actual output for model, m^k , coded as -1 for “x” and 1 for “y”. ε_{ist} is an error term that is assumed independent across subjects, tasks, and periods, and γ_{is} parameterizes the probability of a mistake. We assume that the error term has

a logistic functional form such that it results in the standard logistic form for the likelihood that subject i matches model m^k on task s :

$$p_{is}(m^k) = \prod_T \left(\frac{1}{1 + e^{-\frac{m_{ist}(m^k)}{\gamma_{is}}}} \right)^{y_{ist}} \left(\frac{1}{1 + e^{\frac{m_{ist}(m^k)}{\gamma_{is}}}} \right)^{1-y_{ist}}$$

where T is the number of periods (guesses). Allowing each subject to match a different model on each task, the overall log-likelihood we estimate is given by

$$\mathcal{L} = \sum_S \sum_I \ln \left(\sum_{K_s} \tau_s(m^k) p_{is}(m^k) \right) \quad (1)$$

where S is the number of tasks, I is the number of subjects, and K_s is the number of consistent machines on task s . $\tau(m^k)$ are the main parameters of interest, representing the fraction of subjects that best match each model, m^k , in task s . We estimate the mixture model, (1), with the Expectation-Maximization (EM) algorithm, which provides better convergence properties than maximum likelihood.³¹ We allow the “noise” parameter, γ_{is} , to vary by task to reflect the fact that some tasks are more difficult than others. We ran the estimation both with a homogeneous (across subjects) γ_s for each task and with heterogeneous γ_{is} . A likelihood ratio test rejects the homogeneous version in favor of the heterogeneous version so we report results for the heterogeneous version, but the results are very similar across the two specifications.

B Online Appendix B: Details on Complexity Notions

B.1 Partition Complexity

Lipman (1995) describes a general model of bounded rationality based on limited abilities to process information. A decision maker that is trying to learn a state

³¹We estimate with 50 different starting points and choose the maximum likelihood across runs, but each run gives very similar estimates.

of the world observes a series of inputs, but may not fully process the information contained in those inputs. In particular, Lipman (1995) describes a class of partitional models in which the decision maker partitions the state of the world into several elements, where a finer partition corresponds to more fully processing information. Our partitions complexity measure builds upon this general idea. Intuitively, models which require more of the information in the dataset to be processed (i.e. a finer partition) may be more difficult to infer.

In our setting, both the state of the world and the “input” that the decision maker processes can be taken to be the observed dataset: the ordered set of the inputs and outputs up to time T : $h^T \equiv \{v_0, u_1, v_1, u_2, \dots, v_{T-1}, u_T\}$, where $v_t \in V$ denotes an output of the model and $u_t \in U$ an input.³² The information contained in the dataset is processed to determine the subsequent output, v_T . In what follows, we define a model-based measure of complexity, allowing $T \rightarrow \infty$.³³

We assume decision makers process the dataset by looking for *patterns*: recurring combinations of inputs and past outputs that lead to the same output. Formally, we define a *subdataset*, $h^t \equiv \{v_0, u_1, v_1, u_2, \dots, v_{t-1}, u_t\}$ with $t \in \{1, 2, \dots, T\}$. Let a generic element of the subdataset be denoted $e_t \in \{u_t, v_t\}$. We denote a subset of h^t by S . A pattern, denoted $S \rightarrow v$ is then defined to be a mapping from some S to a constant v , such that $v_t = v$ for all of the subdatasets in the dataset, except for perhaps the subdatasets given by $t = 1, \dots, t'$ for some finite t' . Examples of simple patterns are $u_t = a \rightarrow x$, $v_{t-3} = x \rightarrow x$, and $u_t = a, v_{t-1} = x \rightarrow x$.

A set of patterns forms a *partition* of the set of subdatasets if each subdataset is associated with one and only one pattern in the set. We denote such a set of patterns, $R^* \in \mathbf{R}$ where $\mathbf{R}$ is the set of all possible sets of patterns. Our *partitions complexity* measure is then taken to be the minimum number of partition elements in any $R^* \in \mathbf{R}$: models which require more patterns to describe the dataset require more information processing, and are therefore more complex. For example, I2 can be described by two

³²This construction is similar to the description in Lipman (1995) of the history of a repeated game as both the state of the world and the input. In his formulation, a strategy partitions the history and more complex strategies use finer partitions.

³³A model-based measure is somewhat easier to work with than a measure based on a dataset of finite-length because it avoids having to deal with exceptions that may occur at the beginning of the dataset.

patterns, $u_t = a \rightarrow x$ and $u_t = b \rightarrow y$. A3 can also be described by two patterns, $v_{t-3} = x \rightarrow x$ and $v_{t-3} = y \rightarrow y$. S2 requires four patterns, each of which maps a value of v_{t-1} and a value of u_t to an output.

It is important to note that in constructing a partition of the set of subdatasets through patterns as defined, we are not allowing for any partition. For example, one partition that is ruled out is “all of the subdatasets for which $v_t = x$ and all of the subdatasets for which $v_t = y$ ”. Although a possible partition, this partition is not particularly useful in thinking about the complexity of a model because it leads to a complexity measure of two for all models. Put another way, to construct this partition, a decision maker must infer the model, which, if allowable, can tell us nothing about how difficult a model is to infer.

In limiting to patterns as defined, it may not always be possible to construct a partition with a finite number of patterns. It turns out though that for the eight models we study, finite partitions can be constructed for seven of the them.³⁴ Importantly, given that only one of models requires an infinite number of partitions, our partitions complexity measure can still be used to rank the eight machines. Thus, although we could allow for more complex patterns, simpler patterns are sufficient for our purposes.³⁵ In addition, allowing for more complex patterns would significantly increase the algorithmic complexity of finding the partition with the minimum number of elements. In fact, even with relatively simple patterns, it is a non-trivial problem.

B.2 Sparsity Complexity

Gabaix (2014) provides a general framework for studying standard optimization problems in economics when decision makers must pay costs to attend to inputs to the

³⁴S3 is the exception, requiring a countably infinite number of simple patterns. It requires one to look for the pattern, “if the input is b and the output was x when the third most recent input of b occurred, the next output will be x” which, when restricted to simple patterns, requires a partition with a countably infinite set of simple patterns due to the fact that an arbitrary number of a inputs can occur between b inputs.

³⁵In the selection treatment analysis, several of the models consistent with the datasets also require a countably infinite number of simple patterns so that we cannot rank all consistent models. However, for each of the eight datasets, at least one model can be characterized by a finite number of patterns. As we are only interested in whether or not the simplest model is most often selected, this suffices.

optimization process. In his framework, the decision maker first decides which inputs to attend to by comparing the attention cost to an approximation of the value of the input. He shows that decision makers will optimally consider only a sparse set of inputs, ignoring those with relatively little value. A natural complexity measure inspired by this framework, then, is the number of inputs a decision maker must attend to (i.e. must avoid ignoring) in order to behave optimally. In our setting, the complexity of a model is the number of input/output elements in the dataset the DM must examine in order to correctly apply the model correctly. In our operationalization, each element of the dataset is a potential input that might be processed, and our measure of complexity is the number of elements in the dataset that must be attended to in order to correctly predict the output. For example, A2 requires the DM to attend to v_{t-2} only to guess the next output and is therefore one of the simplest possible models. S2 is more complex, requiring the DM to attend to both u_t and v_{t-1} in order to predict. As with Partitions Complexity, S3, is the only one of our DGPs out of the eight that we study for which there is no finite set of elements that will always predict the output and we code it as maximally complex under this metric.

B.3 Algorithms for Determining Complexity Measures

We discuss the algorithm used to find partition complexity. Sparsity complexity is determined at the same time because once each partition that can describe the model is identified, both the number of patterns in the partition and the number of dataset elements used in constructing the patterns are known. Partition complexity is the minimum number of patterns in any partition and sparsity complexity is the minimum number of elements in any partition.

We take a brute force approach to finding the partition with the minimum number of elements. For a given number, τ , time periods, we form all the possible partitions with the subset of elements, $\{v_{T-\tau-1}, u_{T-\tau}, \dots, v_{T-1}, u_T\}$. If a partition cannot be found, we set partition complexity to infinity and otherwise set it to the number of patterns in the smallest partition. We cannot search over all partitions as $\tau \rightarrow \infty$ and therefore cannot guarantee we have found the partition of minimum size. However,

$\tau = 3$ appears to be sufficient in the sense that the complexity measures do not change for a sample of larger values. Intuitively, the minimum number of patterns should use only elements from the most recent three periods for automata with at most three states, but given all the ways in which partitions can be formed, we have been unable to prove this formally. Proposition 1 does show though that if a partition cannot be formed using elements from the two most recent periods, then adding a finite number of prior elements does not help. So, we can be certain that the cases in which we set the partition complexity to infinity do indeed require an infinite number of patterns.

Proposition 1: If the subset of elements, $v_{t-2}, u_{t-1}, \dots, u_t$, does not produce a partition with a finite number of elements, then incorporating an additional finite set of prior elements does not either.

Proof: For any one or two state automaton, v_{t-1} and u_t , always produce a partition with at most four elements because the output, v_{t-1} , directly maps to the current state of the automaton which, together with the input, determines the next state and output. For three state automata in which two states map to the same output, v_{t-1} and u_t are no longer sufficient. Without loss of generality, assume a single state corresponds to the output, y , and two states, x_1 and x_2 , correspond to x . For v_t not to be determined by some subset of past elements, it must be the case that the state of the automaton is indeterminate after the subset because if the state is uniquely determined, then so is the output, v_t .

For the subset of $v_{t-2}, u_{t-1}, v_{t-1}$, and u_t , not to determine the automaton state, it must be the case that either x_1 and x_2 both transition back to themselves on the same input, or that each transitions to the other on the same input. To see this, suppose to the contrary that neither of these conditions holds. Then for each input, u , either one of the x states transitions to the other and the other transitions back to itself, or one or both of the x states transitions to y . In the first case, v_{t-1} and $u_t = u$ uniquely determines the state, and in the second case, the subset of $v_{t-2}, u_{t-1}, v_{t-1}$, and u_t uniquely determines the state.

Finally, if x_1 and x_2 transition back to themselves on the same input, u , then no finite number of past outputs and inputs necessarily determines the state because u could repeat indefinitely. Similarly, if each of x_1 and x_2 transitions to the other on the same

input, u , then no finite number of past outputs and inputs necessarily determine the state. Therefore, no partition with a finite number of elements exists. $\square$

C Online Appendix C: Instructions to Subjects

The same instructions were provided to subjects in both treatments. Each set of bullet points corresponds to a different page in the instructions and bullets were revealed to subjects one-by-one. Subjects were required to correctly answer comprehension questions, reproduced below, before proceeding to the experiment.

Instructions

- We will start by providing you with **INSTRUCTIONS** for the study.
- We will ask you **COMPREHENSION QUESTIONS** to check that you understand the instructions. You should be able to answer all of these questions correctly.
- Please read and follow the instructions closely and carefully.
- If you **COMPLETE** the main parts of the study, you will receive a **GUARANTEED PAYMENT** of **\$2.50**.
- In addition, your **CHOICES** in the **DECISIONS** portion of the study will result in **PERFORMANCE-BASED EARNINGS**. You will experience several **TASKS** worth **REAL MONEY**. Your points from **ONE OF THE TASKS** will be randomly selected by computer and will be converted into an additional payment.

Inputs and Outputs

- The experiment will consist of several separate **TASKS**. Each Task will consist of a series of **PERIODS**. The screen will look like this:

Part 1: Observe how the computer produces **outputs** for this task.

Inputs: **abaabaabaaaa**
 ↓↓↓↓↓↓↓↓↓↓↓↓
Outputs: **xyxyxyxyxyxy**

- In every Period we will show you an INPUT (a letter in red) followed by an OUTPUT (a letter in blue) as in the screen above.
- The Periods are shown on your screen from left to right – each horizontal location is a different period. In order to help you remember that outputs come **AFTER** inputs in each period. we draw a little arrow ↓ from inputs to outputs in each period.
- The computer uses a **RULE** to decide on outputs each period. That rule may depend on the current and/or past inputs or it **may not depend** on the inputs at all. You will **not know**. Most importantly, the rule is **not random**.
- Although outputs **may** depend on inputs, inputs never depend on outputs (the inputs were selected before the experiment began).
- **Comprehension question: Which of the following is true?**
 1. Inputs come before outputs in each period.
 2. Outputs come before inputs each period.
 3. Inputs and outputs happen at the same time each period.
- **Comprehension question: The rule the computer uses to determine outputs:**
 1. is random.
 2. is a rule that always depends on inputs.
 3. is a rule that never depends on inputs.
 4. is a rule that may or may not depend on inputs.

Making Guesses

- In Part 1 of each Task, we will show you a set of inputs and outputs (one at a time) for 12 periods to help you get a sense of how outputs are determined. Then, in Part 2 of the task we will have a series of 12 additional periods in which we will show you inputs and have you try to **GUESS** the correct output.

•

Part 2: Now **guess** the remaining **outputs** by typing your **guesses**.

Inputs: abaabaabaaaaabbbaaa
 ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓
Outputs: xyxyxyxyxyxyx
Guesses: xyxyx

The screen will look like the image above. In the first 12 periods you make no guess but just observe the process, as it automatically happens. In the last 12 periods, you make a guess each period after each new input appears.

- To make your guess, just type (on your keyboard) the letter of the output that you think comes next (outputs can only be “x” or “y”). The more periods in which you correctly guess the outputs, the higher your **BONUS**.
- **IMPORTANT.** The **rule** the computer uses to determine outputs will stay the same throughout the Task. But it **may change from task to task**. So don’t assume that the way outputs are determined is the same for every Task!
- **IMPORTANT.** The **rule** the computer uses to determine outputs does **not** depend on your guesses. The set of inputs and outputs was determined ahead of time and **does not respond at all** to what you do in the experiment.

- After the last Task, the computer will **randomly select one Task** and pay you a bonus of \$0.35 for each output you correctly guessed in that task.
- **Comprehension question: My guesses influence the inputs and outputs in the experiment**
 1. False
 2. True
 3. It depends
- **Comprehension question: The rule determining outputs**
 1. changes throughout each Task and across Tasks.
 2. is the same throughout each Task, but changes from Task to Task.
 3. is the same throughout each Task and across Tasks.

Giving Advice, Describing the Rule

- At the end of each Task, when you are done making your guesses, we will have you **give advice** to another participant who will face a similar Task in the future.
- The future participant **will not** experience Part 1 of the Task and will see a **different set of inputs** than you did in Part 2.
- However, the **rule** the computer uses to determine inputs will be exactly the same as in your Task. That means, you can only help the future participant by trying your best to **describe the rule** you think the computer used during the Task.
- In Part 3 of the Task we will therefore give you a box to type the rule you think the computer used to determine outputs

- The better job you do in describing the rule, the more money you can earn. With 10% likelihood, we will show what you typed to a future participant. This will be a sophisticated participant – a graduate student in a quantitative field. **You will earn \$0.35 for every letter that participant gets right.** That means, the better job you do of accurately describing the rule, the **more money** you will earn on average in the experiment.
- **Comprehension question: I earn money in Part 3 of the Task by describing to a future participant**
 1. what my guesses were.
 2. the rule the computer uses in the Task.
 3. how people are paid in the experiment.
- **Comprehension question: If I do a good job of describing the rule**
 1. it will have no impact on anything so I shouldn't bother.
 2. it may help another participant earn more money.
 3. it may may help another participant earn more money, AND the more they earn the greater my bonus will be.