



Bayesian collective learning emerges from heuristic social learning

P.M. Krafft^{a,*}, Erez Shmueli^b, Thomas L. Griffiths^c, Joshua B. Tenenbaum^d,
Alex “Sandy” Pentland^e

^a Creative Computing Institute, University of Arts London, London, England, United Kingdom

^b Department of Industrial Engineering, Tel-Aviv University, Tel-Aviv, Israel

^c Department of Psychology, Princeton University, Princeton, NJ, USA

^d Department of Brain and Cognitive Sciences, MIT, Cambridge, MA, USA

^e MIT Media Lab, Cambridge, MA, USA

ARTICLE INFO

Keywords:

Social learning
Bayesian models
Exploration-exploitation dilemma
Collective intelligence
Wisdom of crowds
Big data

ABSTRACT

Researchers across cognitive science, economics, and evolutionary biology have studied the ubiquitous phenomenon of social learning—the use of information about other people’s decisions to make your own. Decision-making with the benefit of the accumulated knowledge of a community can result in superior decisions compared to what people can achieve alone. However, groups of people face two coupled challenges in accumulating knowledge to make good decisions: (1) aggregating information and (2) addressing an informational public goods problem known as the exploration-exploitation dilemma. Here, we show how a Bayesian social sampling model can in principle simultaneously optimally aggregate information and nearly optimally solve the exploration-exploitation dilemma. The key idea we explore is that Bayesian rationality at the level of a population can be implemented through a more simplistic heuristic social learning mechanism at the individual level. This simple individual-level behavioral rule in the context of a group of decision-makers functions as a distributed algorithm that tracks a Bayesian posterior in population-level statistics. We test this model using a large-scale dataset from an online financial trading platform.

1. Introduction

There are thousands of investment opportunities listed on the world’s various stock exchanges. The options each person has for what occupations to pursue or what paths to take in life are vast. Even in decisions as mundane as where to buy a cup of coffee or where to go out to eat for dinner, a city dweller is faced with a dizzying array of options—Boston’s North End neighborhood has over 50 Italian restaurants; downtown Manhattan has hundreds of bars. Furthermore, the information available about the options in each of these cases changes over time, creating complex and ever-evolving decision-making landscapes for many of the choices we face. Yet we do not have to make these decisions alone. Entire communities of people are faced with the same sets of options in many decision-making contexts, and can communicate information about the different options available in the decisions at hand. Decision-making with the benefit of the accumulated knowledge of a community can result in superior decisions compared to what people can achieve alone (Boyd et al., 2011; Hidalgo, 2015; Hillel et al.,

2013; Mason & Watts, 2012; Rendell et al., 2010). Which of the many potential investment opportunities, career paths, or coffee shops is the best fit for a person like you? Relying on information from other people can be an effective component of how to decide.

However, decision-making in the context of a group of people presents its own challenges. Two coupled challenges that groups face in accumulating reliable knowledge to make good decisions are (1) aggregating information and (2) addressing an informational public goods problem known as the exploration-exploitation dilemma (Hills et al., 2015; March, 1991; Toyokawa et al., 2014). The problem of information aggregation is a matter of how to get information as efficiently as possible from as many people as possible who have faced the same decision. In other words, the challenge of information aggregation, at least in cases where preferences are roughly shared, is for decision-makers to pool the experiences they have had and to determine the most informed beliefs about the qualities of options available in the decision at hand. A naïve version of ideal information aggregation—directly sharing all personal preferences and experiences—is not

* Corresponding author.

E-mail address: p.krafft@arts.ac.uk (P.M. Krafft).

<https://doi.org/10.1016/j.cognition.2020.104469>

Received 23 December 2018; Received in revised form 14 September 2020; Accepted 16 September 2020

Available online 24 March 2021

0010-0277/© 2021 The Authors.

Published by Elsevier B.V. This is an open access article under the CC BY-NC-ND license

(<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

possible for people to do, but even if it were, the second challenge of the exploration-exploitation dilemma would still remain to be met. The exploration-exploitation dilemma is a matter of balancing relying on the knowledge that a population of decision-makers has accumulated with contributing to that pooled knowledge through exploration that goes beyond what is known already. If all decision-makers focus on the best-looking options at a given time according to all the available aggregated information, the group will learn little about less-explored potentially higher quality options. Both of these challenges are intrinsic computational problems that groups face in accumulating knowledge about the world. How do groups of people in shared decision-making contexts address these challenges? Are there mechanisms of human collective behavior that enable effective information aggregation and a good balance between exploration and exploitation?

We answer these question by developing a new model that synthesizes approaches from two strands of related work on modeling human social learning. We establish that a simple heuristic social decision-making procedure called social sampling is capable of achieving ideal information aggregation and a good trade-off between exploration and exploitation. To test our model, we study how people address the problems of information aggregation and exploration versus exploitation in a large, highly instrumented social system. We examine collective behavior in an online social financial trading platform. In this environment, users are able to follow and copy each other's trades, and users are therefore faced with a difficult decision of who among the many users of the platform to follow. This observational dataset allows us to study social learning in a large group regime, which is prohibitively costly in the laboratory but theoretically important since the emergent properties of our model fully appear only in large groups. A unique advantage of the environment we study among observational datasets of large groups is that explicit objective evidence of trading performance is available to both users on the site and to us as analysts. We can therefore compute normative benchmarks for ideal information aggregation and exploration versus exploitation, and check the predictions of our model by comparing how closely collective behavior accords with these normative benchmarks.

2. Background

Our work contributes to the extensive literature on social learning, which studies the use of information about other people's decisions to make your own. One key line of work on social learning has centered on what kinds of social learning behaviors and mechanisms people engage in social learning contexts, and how the various candidate behavioral models solve or fail to solve computational problems like information aggregation and exploration versus exploitation. Mathematical and computational models are commonly employed in this literature to try to answer these questions by modeling behavior and then studying the properties of those models. Two different classes of models have been especially common, heuristic social learning models and Bayesian social learning models (Acemoglu & Ozdaglar, 2011; Golub & Sadler, 2016), and prior works have also combined these classes. Our contributions rely on a new synthesis of these two approaches.

2.1. Heuristic social learning models

Heuristic social learning models describe behavior as resulting from simple hard-coded rules involving a combination of social observation and individual consideration (Laland, 2004). In cognitive science, Goldstone and colleagues have studied a range of heuristic social learning mechanisms (Wisdom et al., 2013), as well as how these different mechanisms affect task performance and allow groups to collectively solve problems, from exploring complex decision landscapes (Goldstone et al., 2013; Mason et al., 2008; Mason & Watts, 2012) to finding shortest paths (Gureckis & Goldstone, 2006). Economists and sociologists have a parallel scholarly literature on heuristic social

learning models (DeGroot, 1974; Friedkin, 2006; Friedkin & Johnsen, 2011; Gupta et al., 2006; Lazer & Friedman, 2007; March, 1991). Researchers in complex systems have also studied a range of similar models, including replicator dynamics (Henrich & Boyd, 2002), majority dynamics (Mossel et al., 2014; Tamuz & Tessler, 2015), linear opinion dynamics (Becker et al., 2017), statistical physics models (Castellano et al., 2009), and contagion models (Guille et al., 2013). The heuristic approach has been used to study both problems of information aggregation and exploration versus exploitation, although typically in separate pieces of work using separate models.

2.2. Bayesian social learning models

Bayesian social learning models relate closely to the frameworks of rational agent modeling and Bayesian cognition (Griffiths et al., 2008; Griffiths & Tenenbaum, 2006; Tenenbaum et al., 2011). The core premise of Bayesian cognition is that agents have a mental model of the world that is used for making inferences about the world. Bayesian social learning models most commonly examine how social observation can be optimally integrated into this process of rational Bayesian inference. Griffiths and others in cognitive science have studied how social learning relates to cultural accumulation (Beppu & Griffiths, 2009; Griffiths & Kalish, 2007; Kalish et al., 2007; Navarro et al., 2018; Sanborn & Griffiths, 2008; Thompson & Griffiths, 2019) and optimal use of social information (Baker et al., 2017; Miller & Steyvers, 2011; Whalen et al., 2018). Economists and sociologists have a parallel scholarly literature on Bayesian social learning models examining similar questions (Acemoglu et al., 2011; Bikhchandani et al., 1992; Chamley, 2004; Lobel et al., 2009; Mueller-Frank, 2013). Another highly related line of work is that of Pérez-Escudero and de Polavieja (Pérez-Escudero & De Polavieja, 2011) and colleagues (Arganda et al., 2012; Eguíluz et al., 2015; Pérez et al., 2016). These researchers were some of the first to specify Bayesian models of social decision-making in the context of collective animal behavior. Their model is also distinctive among Bayesian social learning models because it has been successfully empirically tested on human behavioral data. Bayesian social learning models are most often used to study information aggregation, but a similar class of rational game theoretic equilibrium-based analyses have also been used to study the exploration-exploitation trade-off (Bolton & Harris, 1999).

2.3. Boundedly rational social learning models

Researchers have also studied a class of models in between the heuristic and Bayesian approaches called boundedly rational models, which are motivated by the fact that Bayesian computation is too computationally intensive to be cognitively plausible. While Bayesian models involve optimal reasoning according to agents' veridical mental models, boundedly rational models explore relaxations of these assumptions. Most boundedly rational social learning models involve agents performing exact inference in approximate mental models of the environment (Bala & Goyal, 1998; Butts, 1998; Easley & Kleinberg, 2010; Ellison & Fudenberg, 1993; Ellison & Fudenberg, 1995; Eyster & Rabin, 2010; Feldman, Immorlica, Lucier, & Weinberg, 2014; Golub & Jackson, 2010; Goyal, 2011; Jadbabaie, Molavi, Sandroni, & Tahbaz-Salehi, 2012; Molavi, Tahbaz-Salehi, & Jadbabaie, 2018; Rahimian & Jadbabaie, 2017). Typical simplifying assumptions made in this type of boundedly rational model are that agents neglect certain dependencies between observations. Other boundedly rational models involve agents performing approximate inference, or acting probabilistically, using exact mental models of the environment (Anderson & Holt, 1997; Arganda et al., 2012; Whalen et al., 2018). As with fully Bayesian models, the focus of much of the work on boundedly rational models has been on studying information aggregation.

2.4. Synthesizing Bayesian and heuristic models with distributed algorithms

The advantage of the heuristic approach is that heuristics tend to be better descriptive models with better fidelity to data. Their key weakness is that they tend to be post hoc and underconstrained—there are few unifying principles to guide researchers towards particular forms of heuristic mechanisms that might be expected to be observed. In contrast, Bayesian social learning models are often less cognitively plausible due to the extreme complexity involved in fully rational Bayesian reasoning. The advantage of Bayesian models, though, is that they have the normative force of the axiomatic foundation of optimal statistical inference and decision theory, and therefore come with a fuller explanation for why we would expect a particular mechanism to be observed. With the exception of some theoretical work such as by [Rahimian and Jadbabaie \(2017\)](#), boundedly rational models typically keep the structure of Bayesian models while sacrificing the normative force of their axiomatic foundations, effectively specifying heuristic models using Bayesian language.

Our contribution is to use the framework of distributed algorithms to develop a much closer synthesis of the Bayesian and heuristic approaches to modeling social learning. While heuristic models and fully rational Bayesian models are at first glance inconsistent, complex distributed computations are possible from combinations of simple agents ([Chazelle, 2012](#); [Lynch, 1996](#)). Is it possible that some simple heuristic mechanism might be able to function as a distributed algorithm for ideal Bayesian inference at the group level?

The new synthesis of the heuristic and Bayesian perspectives we present combines the unique explanatory benefits of each. The key idea we explore is that Bayesian rationality at the level of a population can be implemented through a heuristic social learning mechanism at the individual level. In this formulation, the computational problems of information aggregation and exploration-versus-exploitation are solved through collective computation at the group level, while individuals behave according to a simple rule-of-thumb heuristic we call “social sampling”. Our social sampling model, which has close parallels to models of evolutionary dynamics ([Chastain et al., 2014](#); [Ewens, 2013](#); [Nowak, 2006](#)), offers a Bayesian formulation of social learning that represents population-level statistics as tracking a Bayesian posterior distribution despite more simplistic heuristic individual-level behavior. This Bayesian social sampling model shows how groups can in principle simultaneously optimally aggregate information and nearly optimally solve the exploration-exploitation dilemma through what would otherwise appear to be a simple social learning heuristic.

Our modeling effort draws upon Goldstone’s and others’ efforts to identify simple rules that can implement distributed computation ([Chazelle, 2012](#); [Goldstone & Janssen, 2005](#)), and upon the Bayesian approach of “top-down” computational modeling ([Griffiths et al., 2010](#); [Hutchins, 1995](#); [Marr, 1982](#)) in which the information processing problem is specified by analysis of the decision-making environment and an optimal solution is derived from the principles of statistics, decision-theory, and in our case, distributed algorithms. Our normative analysis explains why social sampling, among many plausible heuristics, is a uniquely suitable mechanism for groups to employ. Our approach complements other recent frameworks for reconciling heuristic and Bayesian cognition that propose certain classes of heuristic as rational under resource constraints ([Gershman et al., 2015](#); [Lieder & Griffiths, 2020](#)). We propose that *social resources* can buttress Bayesian computation in aggregate, even while individual cognition is resource-constrained.

3. Social sampling model

In order to understand how heuristic social learning behavior could lead to distributed Bayesian computation in aggregate, we first construct a model that demonstrates this effect. This model involves a large group,

i.e. a population, of agents who incorporate social and asocial sources of information in a temporally extended (repeated) shared decision task. The model makes direct predictions, which we test empirically. The goal of the model is to establish how a population of decision-makers using a simple heuristic rule might be able to address the computational problems of information aggregation and balancing exploration versus exploitation to accumulate information about a decision at hand as a population over time. The notation we use is summarized in [Table 1](#).

We assume that at each time t , a set of N agents is faced with a decision between M distinct options. Each of these options, $j \in 1, \dots, M$, has an underlying quality, $\eta_j \in (0, 1)$, and generates a directly observable asocial performance signal, $x_{jt} \in \{0, 1\}$, at each time t . This performance signal, x_{jt} , is related to the reward outcomes in the decision-making task. A decision-maker i receives a positive reward from option j on time step t if $x_{jt} = 1$ and a non-positive reward on that time step if $x_{jt} = 0$, with the probability of a positive performance signal/reward corresponding to the underlying quality of the option, $P(x_{jt} = 1 | \eta_j) = \eta_j$. We let η^* denote the underlying quality of the highest quality option, $\eta^* \geq \eta_j, \forall j$. We denote the history of performance signals for an option j up to a particular time t as $\mathbf{x}_{j,\leq t} = \{x_{j1}, x_{j2}, \dots, x_{jt}\}$, and we denote the total information that has been available about all options up to time t as $\mathbf{X}_{\leq t} = \{\mathbf{x}_{1,\leq t}, \mathbf{x}_{2,\leq t}, \dots, \mathbf{x}_{M,\leq t}\}$. In addition to the asocial information in $\mathbf{X}_{\leq t}$, decision-makers also have social information available to them at each time step. We assume that the social information decision-makers can observe is the popularity of each option in the decision at hand at each time. We let $a_{it} \in 1, \dots, M$ denote the decision of agent i at time t , i.e. the option that agent i chose to select at time t . The popularity of option j at time t is the number of decision-makers who select that option on the previous time step, $p_{jt} = \sum_{i=1}^N (a_{i,t-1} = j)$.

The social learning mechanism we study, which we call social sampling, is a variant of heuristic two-stage decision mechanisms studied by previous researchers ([Howard & Sheth, 1969](#); [Krumme et al., 2012](#); [Payne, 1976](#); [Pratt et al., 2005](#); [Seeley & Buhrman, 1999](#)). The social sampling model that we propose supposes that people first select options to consider by consulting others’ decisions, and then commit to options being considered by privately evaluating whether the options seem good according to recent information available. This first step reduces the cognitive burden of evaluating many options by allowing the decision-maker to consider only a small set of options, rather than all the options available. In contrast to prior proposed two-stage social learning models, we propose that in the second step the decision-maker performs an abbreviated Bayesian computation to assess the quality of the option being considered, which is what enables Bayesian aggregation at the group level in this model.

In the first stage of making a decision at time t , an agent i chooses option $o_{it} \in 1, \dots, M$ to consider at random with probability proportional to the current popularity of that option, $P(o_{it} = j) = \frac{p_{jt}}{\sum_{k=1}^M p_{kt}}$. In the second stage of making a decision, the agent decides whether to accept or reject the option being considered, o_{it} , based on that option’s most recent performance signal. The agent commits to making decision o_{it} with probability $P(a_{it} = j | o_{it} = j) = P(x_{jt} | \eta_j = \eta^*) = (\eta^*)^{x_{jt}} (1 - \eta^*)^{(1-x_{jt})}$. This quantity used in the second stage of decision-making is a Bayesian likelihood function giving the likelihood that option o_{it} is the highest quality option. This second stage is a heuristic use of a Bayesian quantity that is motivated by recent results in the cognitive science literature arguing that people resort to approximate Bayesian computations in many decision-making scenarios ([Gershman et al., 2015](#); [Vul et al., 2014](#)). In the case that the option o_{it} is rejected in the second stage, the agent repeats this two-stage procedure, choosing another option to consider o'_{it} with the same probability proportional to p_{jt} . The same option may be considered again or another option may be considered. This two-stage decision-making procedure is repeated until an option is accepted in the second stage. An algorithmic description of the model is given in [Fig. 1](#).

Because each option is considered according to the same process in

Table 1

Table of notation used in our social sampling model specification, analysis, and application.

Mathematical notation				
$\in, \forall, , \sum$		Standard mathematical notation for “element in”, “for all”, “such that”, and summation over a set		
$P(\cdot), P(\cdot \cdot)$		Notation for the marginal probability and conditional probabilities of observations/events		
$\mathbb{N}; \{\dots\}; (\cdot); \mathcal{P}(\cdot)$		Standard mathematical notation for the set of non-negative integers; a set of arbitrary elements; an open interval; and a power set		
Domain		Name	Social sampling model	eToro application
<i>Indices</i>				
T	$\mathbb{N}$	Time	Number of time steps in the repeated decision-making task	Total number of days of data analyzed
N	$\mathbb{N}$	Agents	Number of agents (decision-makers)	N_t is the number of following relationships on day t
M	$\mathbb{N}$	Options	Number of options available for agents to decide between	M_t is the number of traders available to follow on day t
t, i, j	$1, \dots, T;$ $1, \dots, N;$ $1, \dots, M$		Indices for time, agents, and options in the decision-making task	Indices for day, following user/follow relationship, and followed trader
<i>Decision-making environment</i>				
a_{it}	$1, \dots, M$	Decision	The decision of agent i at time t	The trader followed by the user in follow relationship i on day t
x_{jt}	$\{0, 1\}$	Performance	The outcome generated by option j at time t , which determines the reward for an agent choosing that option on that time step	An indicator variable that equals 1 if the ROI from the trades of trader j on day t is positive, and 0 otherwise
p_{jt}	$\mathbb{N}$	Popularity	The number of agents who chose option j at time $t - 1$	The number of followers trader j has at the end of day $t - 1$
<i>Model parameters</i>				
η_j	$(0, 1)$	Quality	The probability of option j generating a positive performance outcome on any time step	The estimated underlying probability of user j displaying positive performance on any day
η^*	$(0, 1)$		The probability of the highest quality option generating a positive performance outcome	The highest estimated underlying probability of positive performance among all users
o_{it}	$1, \dots, M$		The index of an option tentatively being considered by agent i at time t in an inner step of the social sampling model	The index of a trader that the user in follow relationship i is considering following during day t
<i>Model analysis</i>				
j^*	$1, \dots, M$ and $\mathcal{P}(1, \dots, M)$		In the hide-and-seek model, the index of the single	The set of indices of the traders with the highest estimated

Table 1 (continued)

Domain	Name	Social sampling model	eToro application
		highest quality option	probability of positive performance.
$x_{j \leq t}$	$\{0, 1\}^t$	Performance History	The set of performance outcomes that option j has generated on all time steps up to and including t
$\mathbf{X}_{\leq t}$	$\{0, 1\}^{(M \times t)}$	Total Information	The set of all performance histories for all options at time t
			The days up to and including t on which trader j has had positive versus nonpositive performance
			The record of all traders' current and past performances on day t

each repetition, there is a simple closed form equation that gives the overall probability that agent i makes decision j on time step t , $P(a_{it} = j)$. In each loop of the two-stage process, the joint probability of an agent considering and accepting an option j at time t is the multiplication of the probability in each of the two stages, $P(a_{it} = j, o_{it} = j) = P(o_{it} = j)P(a_{it} = j | o_{it} = j) = \frac{p_{jt}}{\sum_{k=1}^M p_{kt}} (\eta^*)^{x_{jt}} (1 - \eta^*)^{(1-x_{jt})}$. The probability that some option at all is accepted on a particular loop of the two-stage process is given by $\sum_{j=1}^M \frac{p_{jt}}{\sum_{k=1}^M p_{kt}} (\eta^*)^{x_{jt}} (1 - \eta^*)^{(1-x_{jt})}$. The conditional probability on each loop of selecting and accepting j given that some option is ultimately accepted on that iteration is then given by dividing these quantities. Since this probability is identical on each iteration of the loop, the overall probability of an agent choosing option j at time t is:

$$P(a_{it} = j) = \frac{p_{jt} (\eta^*)^{x_{jt}} (1 - \eta^*)^{1-x_{jt}}}{\sum_k p_{kt} (\eta^*)^{x_{kt}} (1 - \eta^*)^{1-x_{kt}}}$$

Both of the two stages in the social sampling model are crucial. Incorporating social information in the first stage by sampling according to popularity allows for the aggregation of information over time, while a personal assessment based on new information in the second stage allows new information to be incorporated. It's also important to note that while a Bayesian computation is being used in the second stage of the two-stage social sampling model, it is only a highly bounded one. The only information each decision-maker accesses is the most recent performance signal associated with the one option or the small set of options being considered. What we will show is that even though each individual decision-maker accesses only this small amount of information, the boundedly rational heuristic social sampling model collectively yields a fully rational Bayesian sampling scheme that leverages all the information available over time for all options.

3.1. Model analysis

Despite its simplicity and heuristic appearance, the social sampling model can achieve both excellent information aggregation and a highly efficient balance between exploration and exploitation. In order to analyze the social sampling model, we consider a simplified model of the decision-making environment. In this simplified model, known as a “hide-and-seek” problem (Shamir, 2014), there is a single best option j^* that has a probability $\eta_j^* = \eta^*$ of producing positive rewards, while all other options j' produce positive and negative rewards uniformly at random, $\eta_{j'} = 0.5$. When the number of options is large or when $\eta^* = 0.5 + \epsilon$ for small $\epsilon > 0$, this hide-and-seek setting can be thought of as a pessimistic assumption about the identifiability of the best option in the environment. In other words, this setting can be interpreted as one in which good options are rare or difficult to identify. Similar results can also be derived in more general contexts (Celis et al., 2017).

Social Sampling Procedure

```

1: for time  $t$  in  $1, \dots, T$  do
2:   for agent  $i$  in  $1, \dots, N$  do
3:     while  $a_{it} = \text{None}$  do
4:       (sample): Agent  $i$  samples  $o_{it}$  to consider with  $P(o_{it} = j) = \frac{p_{jt}}{\sum_{k=1}^M p_{kt}}$ 
5:       (accept/reject): Agent  $i$  decides  $a_{it} := o_{it}$  with probability  $(\eta^*)^{x_{jt}}(1 - \eta^*)^{(1-x_{jt})}$ 

```

Fig. 1. Algorithmic description of the social sampling model. The total number of options considered in the inner **while** loop is a geometric random variable that is finite with probability one, with mean bounded by $1/(1 - \eta^*)$ when $\eta^* > 0.5$.

3.2. Information aggregation with social sampling

We first relate the expected popularity of each option under social sampling to a Bayesian posterior distribution involving all information that has been available in the environment. In the case of the hide-and-seek environment model, the true state of the world is characterized by the identity of j^* , so the rational analysis only needs to consider whether each option is or is not option j^* .

Given the history of all rewards up to time t , $\mathbf{X}_{\leq t}$ (as defined above) the Bayesian posterior distribution over the environment parameter is

$$P(j = j^* | \mathbf{X}_{\leq t}) = \frac{(\eta^*)^{x_{jt}} (1 - \eta^*)^{1-x_{jt}} P(j = j^* | \mathbf{X}_{< t})}{\sum_k (\eta^*)^{x_{kt}} (1 - \eta^*)^{1-x_{kt}} P(k = j^* | \mathbf{X}_{< t})},$$

where we assume a uniform prior $P(j = j^*) = \frac{1}{M}$. This posterior probability bears a striking resemblance to the probability of choosing option j under social sampling. In fact, in an infinite population of decision-makers who implement social sampling, the following invariant will be maintained: $p_{jt} \propto P(j = j^* | \mathbf{X}_{< t})$; i.e., with an infinite population in a “hide-and-seek” environment the popularity of option j at time t will be proportional to the posterior probability that $j = j^*$ given all the information that has been available in the environment. Popularity can thus be precisely understood as compactly summarizing the past information about the options available to decision-makers. In other words, popularity under social sampling has an exact correspondence to a Bayesian posterior distribution that each option is best in a hide-and-seek environment.

3.3. Exploration-exploitation with social sampling

We now argue that social sampling can also be expected to achieve a good balance between exploration versus exploitation by relating the aggregate dynamics of the social sampling model to a well-known, near-optimal single-agent decision-making procedure for multiarmed bandits known as Thompson sampling (Thompson, 1933). A multiarmed bandit is a sequential decision-making task with essentially the same structure as we describe for the context of our model. The Thompson sampling algorithm is a commonly studied approach to solving multiarmed bandits that relies on establishing Bayesian posterior distributions for the underlying quality of each option available. To make a decision on a particular time step, the single-agent algorithm probabilistically samples an option with probability equal to the probability that the option is the best option available given the rewards the agent has seen. In a single-agent environment where each option has a distinct reward probability η_j and where agent i is taking an action a_{it} on each time step t and observing reward x_{it} , the Thompson sampling probability is $P(\eta_j > \eta_k \ \forall k | x_{i1}, \dots, x_{it})$. Thompson sampling works well in practice (Chapelle & Li, 2011) in both stationary and non-stationary environments, and was recently proven to achieve a near-optimal balance between exploration and exploitation in the stationary case (Agrawal & Goyal, 2012; Kaufmann et al., 2012).

To relate social sampling to Thompson sampling, we note that in the case of the hide-and-seek environment, $P(\eta_j > \eta_k \ \forall k | \mathbf{x}) = P(j = j^* | \mathbf{x})$ for any set of observations $\mathbf{x}$. Therefore, since in an infinite population of agents we have from the previous section that $P(a_{it}) = P(j = j^* | \mathbf{X}_{\leq t})$, the

probability that agent i takes action a_{it} in social sampling is equal to the Thompson sampling probability given all information available to the group and under the parametric assumption of the hide-and-seek context.

4. Modeling “follow” decisions in a financial social network

To empirically test the social sampling model, we examine collective behavior in an online social financial trading platform called eToro (Pan et al., 2012). eToro’s platform allows users to make trades on their own, predominantly in foreign exchange markets, or to choose other users on the site to follow. When one user chooses to follow another, the follower allocates a fixed amount of funds to automatically mirroring the trades that the followed user makes. eToro then proportionally executes all of the trades of the followed user on the follower’s behalf. Although in general, copying someone else’s trading could lead to market movement that affects the return of those trades, the trading on eToro is marginal enough that it is unlikely to cause such feedback effects. Therefore, the problem of choosing who to follow on eToro can be well modeled as a choice between options with exogenous reward outcomes.

4.1. eToro context

Day trading in foreign exchange markets is notoriously risky, and typically amounts to little more than gambling. eToro—as a company that mechanizes, encourages, and profits from users’ day trading—faces controversy and criticism about its intentions and practices. Many users complain about losing money because of high fees and deceptive performance statistics. However, some users systematically lose less money on eToro, and traders who follow others tend to perform better than users who make trading decisions for themselves (Pan et al., 2012)—though still often not making profit, at least losing less.

There are several decisions that are intertwined with each other on eToro: whether to put money into the platform at all, how much of your money to trade yourself, and how much to use in social trading. These decisions interact with each other, and also interact both with your own perceptions of your abilities, your judgments about the reliability and capabilities of the platform, and your impression of how well other people on the site trade.

Our study simplifies these complexities by focusing on the choices that are made within social trading behavior exclusively. We only model who decides to follow whom and we do not factor into consideration the trade-offs between investing by following others versus trading for yourself. Empirically this decision is justified by the fact that the majority of users on the site use it either for social trading (following other peoples’ trades) or for nonsocial trading exclusively or almost exclusively, generally not so much for both kinds of trading. Fig. 2 shows the distribution of the proportion of trades for each user that are nonsocial versus social. 8% of users on the site never conduct a nonsocial trade. 47% never conduct a social trade. 86% conduct 75% all of their trades as either exclusively social or exclusively nonsocial. These numbers suggest that empirically we can focus on social trading as a relatively isolated mode of behavior from nonsocial trading on the site.

This analysis choice is also informed by our theoretical motivation. Most social learning models involve discrete forced choice situations

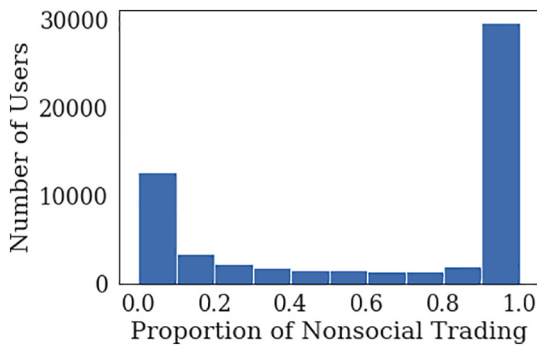


Fig. 2. A histogram plotting for each user the fraction of their trading activity that is purely nonsocial. Most users engage in either purely social trading or purely nonsocial trading on the platform we study, which justifies our analysis decision to focus only on social trading (i.e., decisions about whose trades to copy on the site).

between a fixed set of options. While we could consider nonsocial trading and social trading to both be options in a fuller decision-making model, these options are quite different in terms of how they are supported by the platform. In contrast, the choice of who to follow in social trading is a relatively clear discrete forced choice within the platform, and highly distinct in the platform design from either copying individual trades or making individual trades in the markets the site provides access to. For the remainder of the paper we therefore only consider modeling follow decisions in social trading. The decision-making problem in this case is that each user can choose to follow anyone with a public profile on the site, and traders generate new information about their performance each day through their new trades. A user who wants to follow someone must choose who to follow on each day among all the options available.

Another simplification made by our study is treating each follow relationship as a separate independent decision. While modeling the follow decisions as independent has its limitations (e.g., neglecting the constraint that a user can't follow the same trader multiple times), the actual number of traders each user tends to follow is often just one trader per day (see Fig. 4), and we only analyze the model predictions in aggregate rather than at an individual level.

While the hide-and-seek model we use to derive our analytical results does not need to apply in order for social sampling behavior to be employed, it is worth noting that the assumptions of the hide-and-seek model are also not far from what we observe in the eToro context. As shown in Fig. 3, the vast majority of users have very close to zero average return (just from trading profit/loss; not including all platform fees). Users who are looking for someone to follow can therefore reasonably expect that there are at most a handful of people on the site,

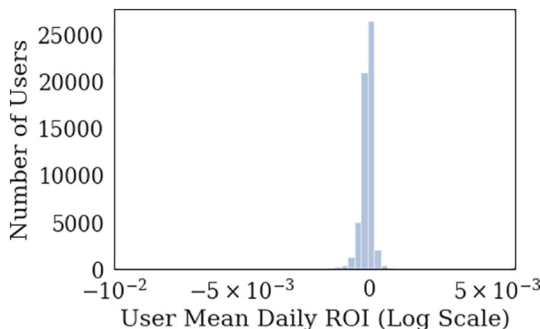


Fig. 3. The vast majority of users on the platform we study achieve close to zero mean daily return on investment (just from trading profit/loss; not including all platform fees). Users are therefore presented a difficult problem of finding good traders to follow.

if even one, that they will benefit from following. The challenge the users face is to find such people among the incredibly noisy information from trading behavior.

In order to assist in users' decisions about who to follow, eToro provides information about the trades and trading performance of each user via a search interface and public profiles. eToro's interface also reports the current popularity of each user, i.e., the number of people currently following that user. The eToro interface has many other complex facets, and has changed substantially over time. At the present time of writing, there were additional features for automatic search, such as a "Top Investors" category that uses a curated combination of search filters. While the current range of advanced curated search features did not exist to our knowledge at the time of data collection, there were many ways to find traders to follow at that time too, such as featured users and followers listed in users' profiles. We unfortunately have no way to reconstruct which users might have been highlighted in such ways, but the predominant search mechanism at the time of our data collection was a table of all the site's users that could be sorted by popularity or various user statistics. The user statistics included a measure of risk, percentage of profitable weeks, and a "gain" statistic most closely related to the performance metric we will use in our modeling. We focus on the gain statistic from that search interface because it was the performance measure that was most emphasized in the search interface and in users' profile pages at the time of data collection.

4.2. Model application

Despite all these ways to search for traders to follow on eToro at the time of our data collection, the interface does not make it easy to make good decisions. The statistics presented vary over time and must be carefully integrated to form a complete picture of each trader's performance. The statistics reported by the platform are potentially over-inflated by the platform designers to incentivize trading activity, and the traders on the platform themselves also try to manipulate their performance statistics to make themselves look good, such as by leaving trades open when those trades have lost money. The easiest way to learn if someone will make you money is to start following them and see what happens, but of course trading performance is highly stochastic. Users on eToro are therefore faced with a difficult decision problem of choosing who to follow among a large set of options of users, given unreliable, multidimensional, temporally varying performance signals, as well as social signals of popularity. We use the social sampling model to understand how well the community of users on eToro manages to address this information processing challenge.

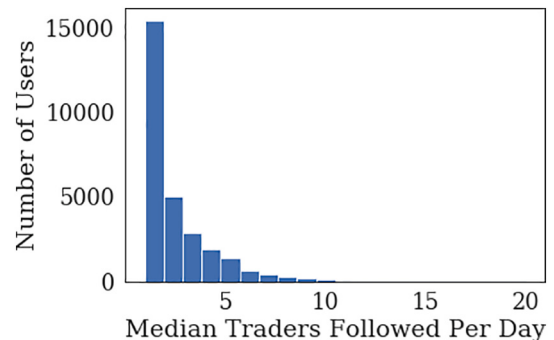


Fig. 4. A histogram of the fraction of users who followed a certain number of traders per day. For each user who followed at least one trader, we take the median number of traders that user followed across all days they followed at least one trader. 54% of the users followed a median of a single trader.

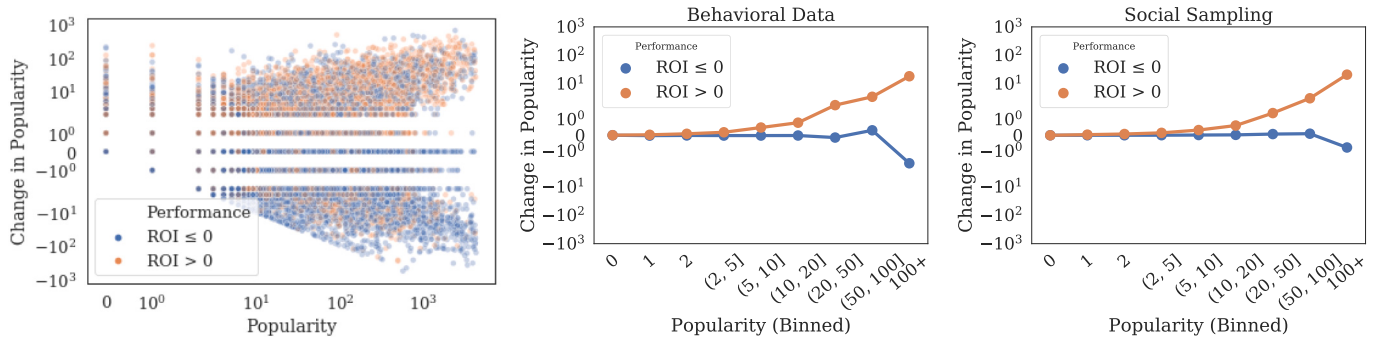


Fig. 5. Daily changes in popularity on eToro tend to be positive for those traders who are performing well and negative for those traders who are performing poorly, and the magnitude of those changes are greater as popularity increases. (Left) A scatter plot illustrating the observed relationship between daily change in popularity on eToro with past popularity and recent performance. There is one data point shown for each trader on each day. Points are colored by whether recent performance is positive or negative. (Center) A binned plot visualizing the same data to highlight the trends we observe. (Right) Predicted changes in popularity according to a fitted social sampling model.

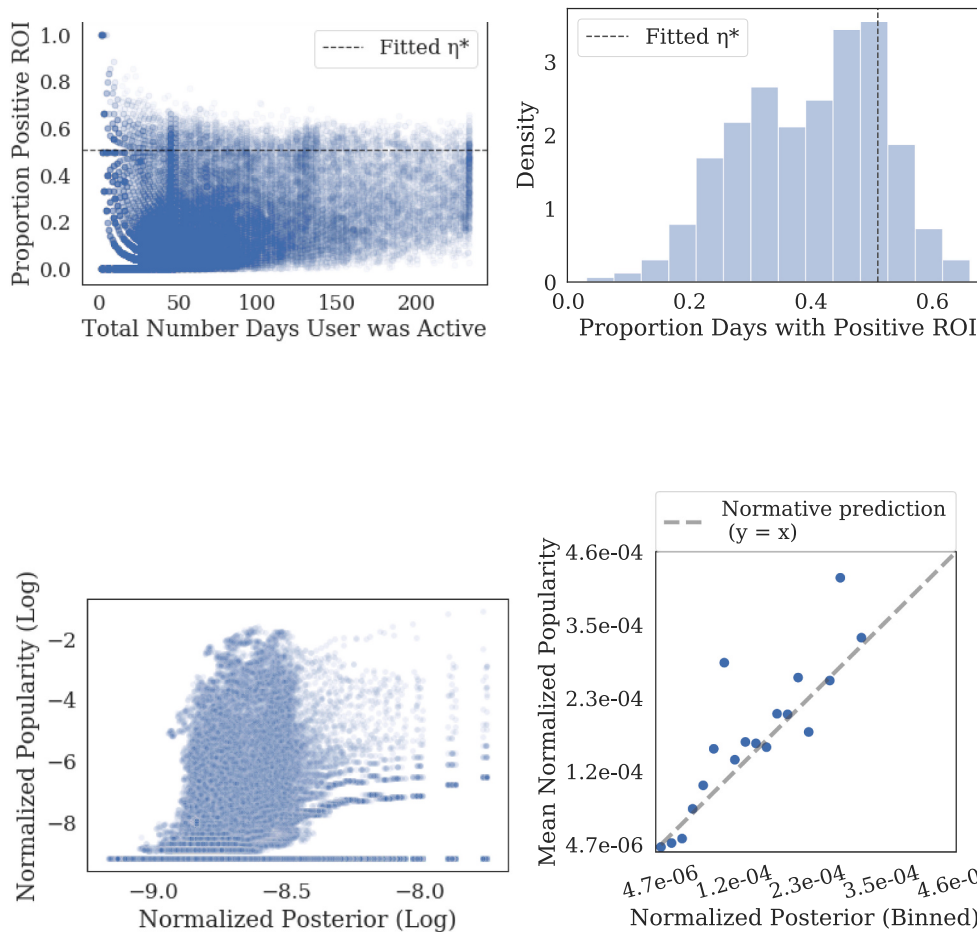


Fig. 6. Plots showing the relationship between the value of η^* fitted to follow decisions users made and the descriptive statistics of the eToro dataset. (Left) A scatter plot comparing the number of days each user was active in our dataset and the proportion of days each user achieved positive ROI. Each point in the plot is a single user. The plot shows that the vast majority of users have below 51% days with positive ROI, and many of those with higher proportions were only active on a smaller number of days. (Right) A histogram plotting the frequency with which users have certain proportions of days with positive ROI. The fitted value of $\eta^* = 0.51$ lies in a high percentile of the distribution.

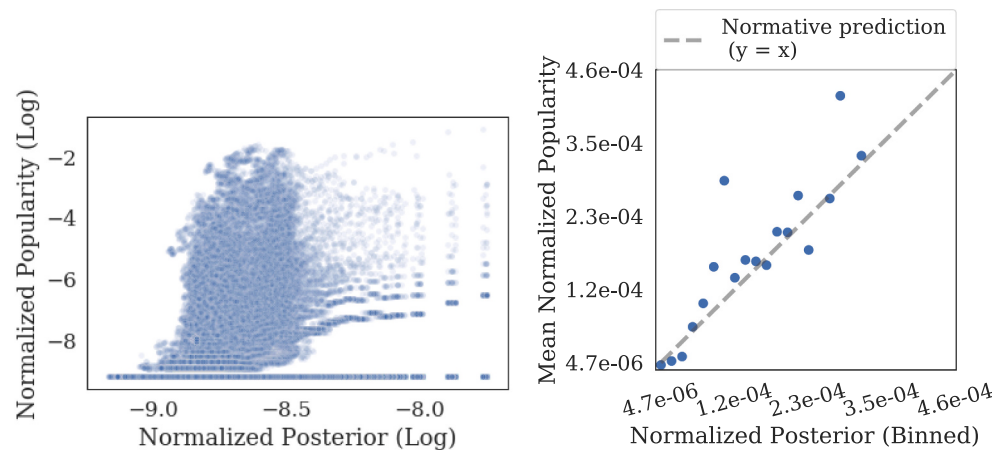


Fig. 7. Plots showing the match between normalized posterior values and popularity on eToro. (Left) Each point is one user on a particular day. The user's normalized posterior on that day is plotted against that user's normalized popularity. The plot reveals a positive relationship between the normalized posterior for each user on each day and normalized popularity. (Right) Each point represents the average popularity of all users within a range of posterior values, using an evenly spaced binning. This plot further highlights the relationship between popularity and the normative posterior distribution of the environment in our data.

5. Methods

5.1. Data

We received our data from the eToro company. The data was generated from the normal activity of users of the website etoro.com. The raw data was in the form of a list of trades conducted on the site. We processed the data to reconstruct follower relationships and aggregate performance statistics. More details on eToro and our data processing are given in the supplementary material. To keep our analysis

computationally tractable, we focused on the first year of data we received, from June 1, 2011 to June 30, 2012. Because of the way we measure active users, the actual days analyzed are July 4, 2011 to June 29, 2012. This time period included 57,455 users. We included each user on each day that user was active, giving us 3,606,903 data points to analyze. We do not endorse eToro as a company or the usage of its services.

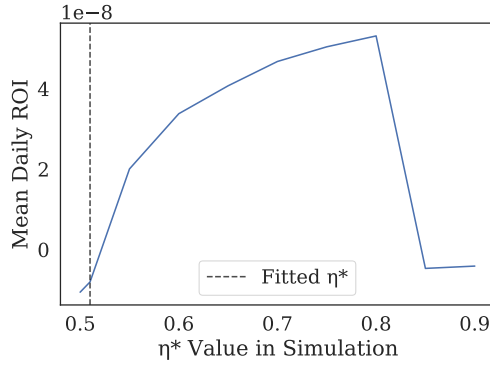


Fig. 8. Simulated mean daily ROI within a population of ideal social samplers following the traders on eToro over the time period we study, for different values of η^* . These simulations check how well the social sampling model balances exploration versus exploitation. The fitted value of η^* that achieves the best predictive accuracy of eToro follow decisions is suboptimal in terms of mean daily ROI in these simulations.

5.2. Popularity and performance signals

We analyze one year of data at a daily temporal granularity, and we model the follow decisions each user makes on each day. Although we do not have access to the specific performance statistics that were displayed to users, we summarize trading performance with return on investment (ROI) from closed trades on the most recent day, which is closely related to the “gain” performance metric presented to the site’s users. Details on how we constructed this proxy are given in our supplementary material. We set the performance signal x_{jt} associated with following a trader j on a particular day t to be positive if liquidating that person’s new trades from that day at the end of the day would yield ROI greater than zero, and set x_{jt} to zero otherwise.

Putting these pieces together with our formal model, we have that at each time t there is a set of traders $1, \dots, M_t$ that each social user on eToro could choose to follow, and each trader generates performance signals/rewards x_{jt} to followers each day. We study how the new popularity $p_{j,t+1}$ of each trader j on each day t is determined as a function of the prior popularity of that trader p_{jt} and their latest performance signal, x_{jt} . We let follows_{jt} be the number of new followers trader j gets on day t , and let unfollows_{jt} be the number of followers trader j loses on day t . In other words, $p_{j,t+1} = p_{jt} + \text{follows}_{jt} - \text{unfollows}_{jt}$. Our analysis assumes each follow decision is independent. We let $N_t = \sum_{j=1}^{M_t} p_{j,t+1}$ be the total number of follow relationships at the end of day t .

5.3. Regression analysis

In our results, we first aim to test the multiplicative interaction between popularity and performance induced by the two stages in the social sampling model. To test for the presence of this interaction, we build a regression model of the new follow decisions made on each day. The regression model we use predicts the normalized number of new followers each user gets on each day based on that user’s normalized popularity from the previous day, and that user’s performance on that day. We include fixed effects for each user and each day. This regression model is $\frac{\text{follows}_{jt}}{\sum_k \text{follows}_{kt}} = \beta_0 + \frac{p_{jt}}{\sum_k p_{kt}} \beta_{\text{pop}} + x_{jt} \beta_{\text{perf}} + \frac{p_{jt}}{\sum_k p_{kt}} x_{jt} \beta_{\text{interaction}} + \beta_j + \beta_t$. We compare the fit of this regression to a reduced model that omits the interaction term. We also compare to another baseline model that looks at aggregated performance over a 30-day period, $\frac{\text{follows}_{jt}}{\sum_k \text{follows}_{kt}} = \beta_0 + \text{ROI}_j^{(30)} \cdot \beta_{\text{perf}} + \beta_j + \beta_t$, where $\text{ROI}_j^{(30)}$ is the signed logarithm of the 30-day rolling average daily ROI of user j . This longer term performance baseline allows us to test whether new followers can be predicted just based on performance, using a longer term performance statistic. The

signed logarithm adjusts for extreme values.

5.4. Fitting the social sampling model

In addition to our regression analysis, we also directly fit the social sampling model to our data. To conduct this model fitting, and to generate model predictions, we compute the follow decisions users on eToro would have been expected to have made on each day according to the social sampling model. We examine aggregates of these decisions in the form of predicting the total number of followers each trader on eToro has on each day. We predict the number of followers each user will have on each day given the performance and popularity of that user (and of every other user) on the previous day, and given the total number of followers across all users on those days.

In the social sampling model, decision-makers make decisions independently, so the probability that a given decision-maker i chooses a specific new option j at time t is given by the decision probability $P_{jt}^{(SS)} = \frac{(\eta^*)^{x_{jt}} (1-\eta^*)^{1-x_{jt}} \cdot p_{jt}}{\sum_k (\eta^*)^{x_{kt}} (1-\eta^*)^{1-x_{kt}} \cdot p_{kt}}$. We fix the total number of follow relationships across all traders to the true value observed in the data, N_t . The expected number of followers user j gets at time t according to this model is then $p_{j,t+1} = N_t P_{jt}^{(SS)}$. We fit η^* by minimizing the mean-squared error in the log daily change in popularity, $\text{sign}(\text{follows}_{jt}^{(SS)} - \text{unfollows}_{jt}^{(SS)}) \cdot \log(|\text{follows}_{jt}^{(SS)} - \text{unfollows}_{jt}^{(SS)}| + 1)$, where $\text{follows}_{jt}^{(SS)} - \text{unfollows}_{jt}^{(SS)} = p_{j,t+1}^{(SS)} - p_{jt}$. A coarse grid search over $[0.5, 0.51, 0.55, 0.6, 0.65, 0.7, 0.75, 0.8, 0.85, 0.9]$ indicates that $\eta^* = 0.51$ presents the best fit.

5.5. Posterior comparison

A key prediction of the social sampling model is that the popularity of each trader should track their performance over time. In order to test this prediction, we compute a normative posterior distribution based on the performance histories of each trader. Rather than employing the hide-and-seek posterior we used in our model analysis, we use a slightly relaxed model that allows for the possibility that multiple traders have the highest underlying quality η^* . In this case, we let j^* be the set of traders with $\eta_j = \eta^*$, and aim to then compute the posterior $P(j \in j^* | \mathbf{X}_{\leq t})$. Without this relaxation, computing the posterior is complicated by missing data on days when traders were inactive. As in the hide-and-seek model we employed in our model analysis, we assume that traders outside the set η^* have $\eta_j = 0.5$.

The posterior here then becomes $P(j \in j^* | \mathbf{x}_{j,\leq t}) = \frac{(\eta^*)^{w_{jt}} (1-\eta^*)^{l_{jt}}}{(\eta^*)^{w_{jt}} (1-\eta^*)^{l_{jt}} + 0.5^{w_{jt}+l_{jt}}}$, where $w_{jt} = \sum_{d=1}^t \mathbf{1}_{x_{jd}=1}$ is the number of positive performance signals a trader has had over time, and $l_{jt} = \sum_{d=1}^t \mathbf{1}_{x_{jd}=0}$ is the number of negative performance signals a trader has had. We then normalize these values to obtain the probability given by Thompson sampling on this posterior: $\frac{P(j \in j^* | \mathbf{x}_{j,\leq t})}{\sum_k P(k \in j^* | \mathbf{x}_{k,\leq t})}$. We compare these Thompson sampling values to daily popularity normalized by the total number of follow relationships on each day, $\frac{p_{jt}}{\sum_k p_{kt}}$. In our empirical analysis, we use $\eta^* = 0.51$ based on our model fitting procedure.

5.6. Simulating performance

Finally, to explore the balance between exploration and exploitation achieved by the social sampling model in the eToro dataset, we simulate the behavior of an entire population of social samplers over the duration of our dataset. We look at how a population of social samplers would perform with alternative values of η^* . To do so, we retrospectively simulate social sampling using the actual profits and losses from trades on eToro. To accommodate predicting positive changes in popularity for traders with zero followers, we add a small smoothing constant to popularity, $\frac{(\eta^*)^{x_{jt}} (1-\eta^*)^{1-x_{jt}} \cdot (p_{jt} + \epsilon)}{\sum_k (\eta^*)^{x_{kt}} (1-\eta^*)^{1-x_{kt}} \cdot (p_{kt} + \epsilon)}$, where $\epsilon > 0$ is a small smoothing

parameter that ensures all users have some probability of gaining followers. We arbitrarily set $\epsilon = 0.0001$. On each day, we set the size of the population of social sampling agents to N_t . As in our model fitting, we examine simulated behavior over the range $\eta^* \in [0.5, 0.51, 0.55, 0.6, 0.65, 0.7, 0.75, 0.8, 0.85, 0.9]$.

6. Results

6.1. Evidence for social sampling

Our first analysis confirms that the dynamics of popularity on eToro are well-modeled by social sampling. These results are summarized in Fig. 5. Traders who perform well on the site tend to gain followers, and traders who perform poorly tend to lose followers. At the same time, the magnitude of these changes becomes larger as the popularity of the trader increases. People with few followers are unlikely to gain many followers, even when they perform well. People with higher popularity gain more followers when they perform well, but lose more followers when they perform poorly (the interaction coefficient, $\beta_{interaction} = 0.92$, $p < 0.0001$, is positive and statistically significant). The regression model that includes the interaction term between popularity and recent performance improves the amount of variance explained to an R^2 of 0.37 compared to the model without an interaction term, $R^2 = 0.27$, and a model that just includes a longer term performance metric, $R^2 < 0.01$. This analysis shows that neither popularity nor performance alone can explain how users decide on new traders to follow.

Simulations from a social sampling model with η^* parameter fitted to the data confirm that social sampling can replicate this pattern of changes in popularity. These results are shown in Fig. 5. We also check that the fitted value of $\eta^* = 0.51$ qualitatively matches the descriptive statistics of the dataset. Fig. 6 shows that, consistent with the interpretation of η^* in our model analysis as the plausible highest probability of positive returns, 95% of users have lower than a probability of 0.51 of achieving positive returns.

Anecdotal reports from users of eToro also corroborate the social sampling model. One website states: “Here are some tips and things we look at when selecting the Professional Investors/Traders we copy: Most people will still want to start with looking at the ‘most copied traders’... Popularity is obviously a decent ‘starting’ point for finding traders to analyze further... putting in some time and effort to analyze the additional statistics is likely to lead to better long term results.”¹ Another user explains: “[Ranking by gain] is what eToro’s [standard] ranking system is showing. And but [sic] this is not so much the way to really choose who’s a good trader because you just don’t know how long these traders have been trading until you go into details. Another good idea is to do this: What I’ll do is, you just go to ‘Copiers’, you just select the ‘Copiers’ tab and show the ones who have the most copiers. Now, this isn’t a full-on good guide either, just choosing the amount of copiers. Cuz as you see, there are other traders in this line who have been making more. Like an example, over 300%, over 300%, and they’re down the line just cuz they’ve got less copiers.”² These users are both describing how to use popularity as a kind of first-pass filter, before looking at performance information.

6.2. Information aggregation

We next directly test whether in this environment the aggregate population statistic of popularity tracks a normative Bayesian posterior distribution, i.e. that the variations in belief in the population of users reflect a rational representation of uncertainty about the decision of

whom to follow. We find that normalized popularity is positively correlated ($r = 0.03$) with the normalized posterior described in our Methods section. Looking just at users with greater than zero followers, the correlation is higher ($r = 0.17$). Taking the logarithm of each quantity to reduce the impact of outliers increases the correlation ($r = 0.06$ for all traders and $r = 0.43$ for traders with greater than zero popularity). Fig. 7 visualizes this relationship.

6.3. Exploration-exploitation

The results of our simulation analysis are shown in Fig. 8. We estimate the mean daily ROI for a population of social sampling agents that use different values of η^* in our eToro data. A value close to 0.5, like the fitted parameter value in our above analysis, may be optimal in more general contexts (Celis et al., 2017), but is conservative in our data. Values closer to 0.5 lead to slower learning, and even though the population of users on eToro maintain rational representation of uncertainty, we observe that collective learning is slower than optimal for this dataset. Therefore, even though social sampling may have the capacity to achieve a near optimal balance of exploration versus exploitation given an appropriate η value, the observed balance is suboptimal.

7. Discussion

We have examined how people address the computational problems of information aggregation and the exploration-exploitation dilemma in a large, highly instrumented social system. We proposed a social sampling model that accurately models millions of decisions performed within an online social financial trading platform. Social sampling consists of a two-step decision-making process of seeking recommendations from other people, and then privately evaluating those recommendations. We established a relationship between social sampling and a well-known, near-optimal Bayesian learning and decision-making procedure called “Thompson sampling” (Agrawal & Goyal, 2012; Kaufmann et al., 2012; Thompson, 1933). This relationship reveals that groups can in theory use a simple mechanism to dynamically aggregate information, while collectively balancing exploration and exploitation, by using a simple probabilistic decision-making mechanism that approximately implements Thompson sampling in the aggregate. We empirically validated the information aggregation property predicted by this relationship, and also explored the balance people achieve between exploration and exploitation. Our results indicate that a form of Bayesian population rationality emerges from heuristic social learning in the case we study. The balance between exploration versus exploitation we observed in this case was suboptimal, although still consistent with the social sampling model under a suboptimal parameter setting.

7.1. Connections to the wisdom of crowds

Beyond the literature on social learning we sought to inform, our paper also contributes to an ongoing debate in the literature on the wisdom of crowds around whether or in what ways social learning undermines versus promotes the wisdom of crowds (Almaatouq et al., 2020; Becker et al., 2017; Lorenz et al., 2011; Surowiecki, 2005). Many formal models of the wisdom of crowds rely on individuals in the crowd having independent pieces of information rather than information gained through social observation. The incorporation of both social information and making decisions based on your own experience is a crucial component of the social sampling model we study. In the context of eToro, all users in principle have access to all the same information. Any trader profile can be viewed by anyone. However it would be impossible for every user to view every profile to make a decision about who to follow. Independent pieces of information therefore come into play through users’ personal analyses of the performance of particular traders that those users decide to consider following. Simultaneously, paying attention to popularity facilitates ongoing aggregation of the

¹ “eToro Tips: Find Best Gurus” from SocialTradingGuru.com (<http://socialtradingguru.com/tips/etoro-tips/select-etoro-gurus>).

² “[Etoro Guide] How to Choose Good Traders or Gurus to Copy in Etoro?” YouTube (<https://www.youtube.com/watch?v=ym8Amfurzb8>, 2:42).

information from those personal assessments.

Independent judgments through personal assessment and gradually gaining experience in the environment can also play a deeper role in social sampling. In the context of the social sampling model, the way private assessments come into play is in whether to accept or reject the popular options you are considering. In the basic social sampling model decisions are resampled every day, but a more robust model with equivalent collective behavior is a win-stay lose-sample kind of dynamic (Bonawitz et al., 2014) in which individuals stick with the options they have adopted until the personal evidence they collect through paying attention to that option invalidate it as a good choice.

7.2. Potential extensions

Looking beyond the eToro context, we can think about further extensions of the social sampling model. In many cases, preferences or taste will vary greatly from one individual to the next. In such cases, overall popularity is less relevant to an individual as an informative social signal. A similar challenge is to incorporate different levels of expertise among decision-makers. In some cases, there will only be a small set of experts who can evaluate options accurately, even if many people are expressing opinions. The social sampling model we examined can be viewed as social sampling on a random network. Looking at social sampling on complex networks, or networks with properties such as homophily, is also a natural next step. Similar studies have been investigated in the broader literature on social learning (Acemoglu et al., 2011; Golub & Jackson, 2012; Lobel & Sadler, 2015). To better understand how social sampling might be related to shared belief formation, looking at social sampling in cases where individual beliefs are components of more complicated systems of beliefs is also an interesting possibility (Friedkin et al., 2016). It could also be interesting to extend our theoretical analysis to investigate optimal solutions to a related set of multi-armed bandit problems where multiple arms are pulled simultaneously. Thompson sampling algorithms have also been studied for this problem formulation (Komiyama et al., 2015). In all these potential extensions, the distinctive charge of the social sampling approach would be to constrain the extensions so that valid Bayesian learning still occurs in the aggregate. To do so may require careful attention to the computer science literature on distributed Bayesian inference (cf., (Alanyali et al., 2004; Angelino, Johnson, & Adams, 2016; Ho et al., 2016; Misra et al., 2011; Nishihara et al., 2014; Smith et al., 2013; Wang et al., 2019; Yang et al., 2018)).

7.3. Limitations

We finally address various limitations our work. Mathematical modeling always involves abstraction, and there are many details of the eToro context that are mismatched with an idealized social sampling model. First and foremost, it's almost certainly the case that our social sampling model is only an approximation of what is happening on eToro. There are undoubtedly a variety of behaviors that people display in interacting with the site, some of which conform more or less to the social sampling model, others of which are somewhat approximated by the model, and a final category of which completely diverge. All our evidence for the social sampling model is at the aggregate level, in the trends in changes in popularity. Our analysis shows that social sampling is not a bad approximation in aggregate. We focused on Marr's computational level of analysis (Marr, 1982)—of specifying a computational problem faced by users on eToro—and we studied one plausible distributed algorithm that we have some evidence for observing in terms of aggregate dynamics. Aside from assessing some anecdotal reports, we did not delve deeply into Marr's implementation level of analysis and examine how particular users' detailed interactions with the site might yield the aggregate dynamics we observe or might implement the distributed algorithm of social sampling.

There are also other ways the social sampling model itself and our

rational analysis of the eToro environment could be improved. We neglected the possibility of correlation in reward signals across time, changes in trader skill over time, traders coming or leaving the ecosystem, and the network effects of followers following followers. Along similar lines, there are several reasons to question how to interpret the fitted parameters of the social sampling model. Past performance of trading is not necessarily informative about future performance. People are risk averse: they care about both returns and the variance of returns. People are also choosing whether to follow others or make their own investment decisions. The fit of the model could surely be improved with further extensions along these lines, but these extensions would also yield a model that is more difficult to analyze. We focused on a simple model in order to isolate and test the key insights from social sampling. We are encouraged by the fact that our current social sampling model can replicate the average dynamics of popularity in its predictions.

There are also several questions about the generalizability of our findings. eToro is an online sociotechnical system with an interface designed by software developers to facilitate following behavior. A major threat to generalizability is the contingency of our results on the design choices of this system. For instance, affording users the ability to sort traders by popularity surely encourages attention to that feature, and position bias in that returned list may affect how social sampling is implemented (Lerman & Hogg, 2014). However, the ways in which information is presented on the website do not guarantee the outcomes of our results. There is always an interaction between the structure of an environment and agent behavior in the environment in determining collective behavior. In the case we study, the users could very plausibly rely exclusively on performance statistics, which are prominently featured and easily searchable, rather than relying on social information at all, for instance. Even if the results were determined by the interface, though, our model would still be a contribution to understanding how that particular interface design yields good properties in terms of information aggregation and exploration versus exploitation in the population of users on the site. One of the strengths of normative models of the sort we pursue is exactly their usefulness for design questions of that sort. This topic deserves further research and offers a compelling connection to the literature on platform design from the computer science communities of human-computer interaction and computer-supported cooperative work. Recent research in these fields has begun to investigate how interface design can impact the extent to which individuals conform to Bayesian reasoning in interacting with data visualizations (Kim et al., 2019; Krafft & Spiro, 2019). Future research synthesizing that line of work with our own could study how interface design moderates the extent to which online communities are able to effectively co-produce knowledge and accumulate information through something like distributed Bayesian computation. It is possible that design aspects such as featured traders might promote knowledge production relative to this normative ideal, or it is possible that such design features undermine knowledge production. What we have focused on showing in our work is that the framework of distributed Bayesian computation is at least a useful lens for studying these questions, and that the behavior on eToro shows a surprising degree of conformity to a normative standard along these lines—whether that is due to innate human behavioral mechanisms, encouragement from the interface, or a combination of both.

In weighing our evidence for social sampling against the limitations of our analysis, we are further encouraged by the similarity between the social sampling model and the many other two-stage decision-making models from the existing literature (Howard & Sheth, 1969; Krumme et al., 2012; Payne, 1976; Pratt et al., 2005; Seeley & Buhrman, 1999). The specific form of social sampling that differentiates it from other two-stage models was motivated more by theoretical considerations and constraints from the literature on cognitive science than by fitting to the eToro context. Social sampling is an intuitive heuristic that could easily be implemented in a variety of contexts. eToro was uniquely suitable as

a test of social sampling because of its ecological validity and the existence of the explicit objective information signals needed to compute a normative posterior.

Credit authorship contribution statement

PK led in conceptualizing the study, designing the model, designing and conducting the analysis, and writing the paper. ES assisted with study conceptualization, data processing, and writing. TG, JT, and AP assisted with study conceptualization, model design, and writing.

Declaration of competing interest

The authors declare no competing interests.

Acknowledgements

Special thanks to Julia Zheng for contributing to an early version of this work, Wei Pan for discussion about the data, Nicolás Della Penna for advice on statistical tests, Yaniv Altshuler for providing the dataset, Guy Zyskind for assisting with data curation, and Jonathan Huggins for discussions about the mathematics of our model. This research was partially sponsored by the Army Research Laboratory Cooperative Agreement Number W911NF-09-2-0053, the United States Defense Advanced Research Projects Agency (DARPA) Cooperative Agreement D17AC00004, and a National Science Foundation Graduate Research Fellowship Grant No. 1122374. Views and conclusions in this document are those of the authors and should not be interpreted as representing the policies, either expressed or implied, of the sponsors.

Appendix A. Supplementary material

Supplementary code, data, and text to this article can be found online at <https://doi.org/10.1016/j.cognition.2020.104469>. Our supplementary text includes descriptions of our released data, further details of the online platform our data is from, further details of our data processing, and further descriptive statistics of the data. Our supplementary code and data allows for reproducing our results and figures.

References

- Acemoglu, D., Dahleh, M. A., Lobel, I., & Ozdaglar, A. (2011). Bayesian learning in social networks. *The Review of Economic Studies*, 78(4), 1201–1236.
- Acemoglu, D., & Ozdaglar, A. (2011). Opinion dynamics and learning in social networks. *Dynamic Games and Applications*, 1(1), 3–49.
- Agrawal, S., & Goyal, N. (2012). Analysis of Thompson sampling for the multi-armed bandit problem. In *JMLR: Workshop and Conference Proceedings: 23. 25th annual conference on learning theory* (pp. 39.1–39.26).
- Alanyali, M., Venkatesh, S., Savas, O., & Aeron, S. (2004). Distributed Bayesian hypothesis testing in sensor networks. In *6. Proceedings of the 2004 American control conference* (pp. 5369–5374). IEEE.
- Almaatouq, A., Noriega-Campero, A., Abdulrahman, A., Krafft, P. M., Moussaid, M., & Pentland, A. (2020). Adaptive social networks promote the wisdom of crowds. *Proceedings of the National Academy of Sciences*, 117(21), 11379–11386.
- Anderson, L. R., & Holt, C. A. (1997). Information cascades in the laboratory. *The American Economic Review*, 87(5), 847–862.
- Angelino, E., Johnson, M. J., & Adams, R. P. (2016). Patterns of scalable Bayesian inference. *Foundations and Trends® in Machine Learning*, 9(2–3), 119–247.
- Arganda, S., Pérez-Escudero, A., & de Polavieja, G. G. (2012). A common rule for decision making in animal collectives across species. *Proceedings of the National Academy of Sciences*, 109(50), 20508–20513.
- Baker, C. L., Jara-Ettinger, J., Saxe, R., & Tenenbaum, J. B. (2017). Rational quantitative attribution of beliefs, desires and percepts in human mentalizing. *Nature Human Behaviour*, 1(4), 0064.
- Bala, V., & Goyal, S. (1998). Learning from neighbours. *The Review of Economic Studies*, 65(3), 595–621.
- Becker, J., Brackbill, D., & Centola, D. (2017). Network dynamics of social influence in the wisdom of crowds. *Proceedings of the National Academy of Sciences*, 114(26), E5070–E5076.
- Beppu, A., & Griffiths, T. (2009). Iterated learning and the cultural ratchet. In *31. Proceedings of the annual meeting of the cognitive science society*.
- Bikhchandani, S., Hirshleifer, D., & Welch, I. (1992). A theory of fads, fashion, custom, and cultural change as informational cascades. *Journal of Political Economy*, 100(5), 992–1026.
- Bolton, P., & Harris, C. (1999). Strategic experimentation. *Econometrica*, 67(2), 349–374.
- Bonawitz, E., Denison, S., Gopnik, A., & Griffiths, T. L. (2014). Win-stay, lose-sample: A simple sequential algorithm for approximating bayesian inference. *Cognitive Psychology*, 74, 35–65.
- Boyd, R., Richerson, P. J., & Henrich, J. (2011). The cultural niche: Why social learning is essential for human adaptation. *Proceedings of the National Academy of Sciences*, 108, 10918–10925.
- Butts, C. (1998). A Bayesian model of panic in belief. *Computational & Mathematical Organization Theory*, 4(4), 373–404.
- Castellano, C., Fortunato, S., & Loreto, V. (2009). Statistical physics of social dynamics. *Reviews of Modern Physics*, 81(2), 591.
- Celis, E., Krafft, P. M., & Vishnoi, N. (2017). A distributed learning dynamics in social groups. In *ACM symposium on principles of distributed computing (PODC)*.
- Chamley, C. (2004). *Rational herds: Economic models of social learning*. Cambridge University Press.
- Chapelle, O., & Li, L. (2011). An empirical evaluation of Thompson sampling. In *Advances in Neural Information Processing Systems*, 2249–2257.
- Chastain, E., Livnat, A., Papadimitriou, C., & Vazirani, U. (2014). Algorithms, games, and evolution. *Proceedings of the National Academy of Sciences*, 111(29), 10620–10623.
- Chazelle, B. (2012). Natural algorithms and influence systems. *Communications of the ACM*, 55(12), 101–110.
- DeGroot, M. H. (1974). Reaching a consensus. *Journal of the American Statistical Association*, 69(345), 118–121.
- Easley, D., & Kleinberg, J. (2010). *Networks, crowds, and markets: Reasoning about a highly connected world*. Cambridge University Press.
- Egürlüz, V. M., Masuda, N., & Fernández-Gracia, J. (2015). Bayesian decision making in human collectives with binary choices. *PLoS One*, 10(4), Article e0121332.
- Ellison, G., & Fudenberg, D. (1993). Rules of thumb for social learning. *Journal of Political Economy*, 101(4), 612–643.
- Ellison, G., & Fudenberg, D. (1995). Word-of-mouth communication and social learning. *The Quarterly Journal of Economics*, 110(1), 93–125.
- Ewens, W. J. (2013). *Population genetics*. Springer Science & Business Media.
- Eyster, E., & Rabin, M. (2010). Naive herding in rich-information settings. *American Economic Journal: Microeconomics*, 2(4), 221–243.
- Feldman, M., Immorlica, N., Lucier, B., & Weinberg, S. M. (2014). Reaching consensus via non-Bayesian asynchronous learning in social networks. In *28. LIPICs-Leibniz International Proceedings in Informatics*. Schloss Dagstuhl-Leibniz-Zentrum fuer Informatik.
- Friedkin, N. E. (2006). *A structural theory of social influence*. 13. Cambridge University Press.
- Friedkin, N. E., & Johnsen, E. C. (2011). *Social influence network theory: A sociological examination of small group dynamics*. 33. Cambridge University Press.
- Friedkin, N. E., Proskurnikov, A. V., Tempo, R., & Parsegov, S. E. (2016). Network science on belief system dynamics under logic constraints. *Science*, 354(6310), 321–326.
- Gershman, S. J., Horvitz, E. J., & Tenenbaum, J. B. (2015). Computational rationality: A converging paradigm for intelligence in brains, minds, and machines. *Science*, 349(6245), 273–278.
- Goldstone, R. L., & Janssen, M. A. (2005). Computational models of collective behavior. *Trends in Cognitive Sciences*, 9(9), 424–430.
- Goldstone, R. L., Wisdom, T. N., Roberts, M. E., & Frey, S. (2013). Learning along with others. In *58. Psychology of learning and motivation* (pp. 1–45). Elsevier.
- Golub, B., & Sadler, E. (2016). Learning in social networks. In *The Oxford handbook of the economics of networks*.
- Golub, B., & Jackson, M. O. (2010). Naive learning in social networks and the wisdom of crowds. *American Economic Journal: Microeconomics*, 112–149.
- Golub, B., & Jackson, M. O. (2012). How homophily affects the speed of learning and best-response dynamics. *The Quarterly Journal of Economics*, 127(3), 1287–1338.
- Goyal, S. (2011). Learning in networks. In *1. Handbook of social economics* (pp. 679–727). Elsevier.
- Griffiths, T. L., & Kalish, M. L. (2007). Language evolution by iterated learning with Bayesian agents. *Cognitive Science*, 31, 441–480.
- Griffiths, T. L., Chater, N., Kemp, C., Perfors, A., & Tenenbaum, J. B. (2010). Probabilistic models of cognition: Exploring representations and inductive biases. *Trends in Cognitive Sciences*, 14(8), 357–364.
- Griffiths, T. L., Kemp, C., & Tenenbaum, J. B. (2008). Bayesian models of cognition. In *The Cambridge handbook of computational psychology*. Cambridge University Press.
- Griffiths, T. L., & Tenenbaum, J. B. (2006). Optimal predictions in everyday cognition. *Psychological Science*, 17(9), 767–773.
- Guille, A., Hacid, H., Favre, C., & Zighed, D. A. (2013). Information diffusion in online social networks: A survey. *ACM Sigmod Record*, 42(2), 17–28.
- Gupta, A. K., Smith, K. G., & Shalvey, C. E. (2006). The interplay between exploration and exploitation. *Academy of Management Journal*, 49(4), 693–706.
- Gureckis, T. M., & Goldstone, R. L. (2006). Thinking in groups. *Pragmatics & Cognition*, 14(2), 293–311.
- Henrich, J., & Boyd, R. (2002). On modeling cognition and culture: Why cultural evolution does not require replication of representations. *Journal of Cognition and Culture*, 2(2), 87–112.
- Hidalgo, C. (2015). *Why information grows: The evolution of order, from atoms to economies*. Basic Books.
- Hillel, E., Karnin, Z. S., Koren, T., Lempel, R., & Somekh, O. (2013). Distributed exploration in multi-armed bandits. In *Advances in neural information processing systems* (pp. 854–862).

- Hills, T. T., Todd, P. M., Lazer, D., Redish, A. D., Couzin, I. D., & Cognitive Search Research Group. (2015). Exploration versus exploitation in space, mind, and society. *Trends in Cognitive Sciences*, 19(1), 46–54.
- Ho, M. K., Littman, M., MacGlashan, J., Cushman, F., & Austerweil, J. L. (2016). Showing versus doing: Teaching by demonstration. In , 3027–3035. *Advances in Neural Information Processing Systems*.
- Howard, J. A., & Sheth, J. N. (1969). *The theory of buyer behavior*. MIT Press.
- Hutchins, E. (1995). *Cognition in the wild*. MIT Press.
- Jadbabaie, A., Molavi, P., Sandroni, A., & Tahbaz-Salehi, A. (2012). Non-Bayesian social learning. *Games and Economic Behavior*, 76(1), 210–225.
- Kalish, M. L., Griffiths, T. L., & Lewandowsky, S. (2007). Iterated learning: Intergenerational knowledge transmission reveals inductive biases. *Psychonomic Bulletin & Review*, 14(2), 288–294.
- Kaufmann, E., Korda, N., & Munos, R. (2012). Thompson sampling: An asymptotically optimal finite-time analysis. In *Algorithmic learning theory* (pp. 199–213). Springer.
- Kim, Y.-S., Walls, L. A., Krafft, P. M., & Hullman, J. (2019). A Bayesian cognition approach to improve data visualization. In , 1–14. *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*.
- Komiyama, J., Honda, J., & Nakagawa, H. (2015). Optimal regret analysis of Thompson sampling in stochastic multi-armed bandit problem with multiple plays. In *International Conference on Machine Learning*, 1152–1161.
- Krafft, P. M., & Spiro, E. S. (2019). Keeping rumors in proportion: Managing uncertainty in rumor systems. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems, CHI '19* (pp. 1–11). New York, NY, USA: Association for Computing Machinery.
- Krumme, C., Cebrian, M., Pickard, G., & Pentland, A. (2012). Quantifying social influence in an online cultural market. *PLoS ONE*, 7(5), Article e33785.
- Laland, K. N. (2004). Social learning strategies. *Animal Learning & Behavior*, 32(1), 4–14.
- Lazer, D., & Friedman, A. (2007). The network structure of exploration and exploitation. *Administrative Science Quarterly*, 52(4), 667–694.
- Lerman, K., & Hogg, T. (2014). Leveraging position bias to improve peer recommendation. *PLoS ONE*, 9(6), Article e98914.
- Lieder, F., & Griffiths, T. L. (2020). Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behavioral and Brain Sciences*, 43.
- Lobel, I., Acemoglu, D., Dahleh, M., & Ozdaglar, A. (2009). Rate of convergence of learning in social networks. *Institute of Electrical and Electronics Engineers*.
- Lobel, I., & Sadler, E. (2015). Preferences, homophily, and social learning. *Operations Research*, 64(3), 564–584.
- Lorenz, J., Rauhut, H., Schweitzer, F., & Helbing, D. (2011). How social influence can undermine the wisdom of crowd effect. *Proceedings of the National Academy of Sciences*, 108(22), 9020–9025.
- Lynch, N. A. (1996). *Distributed algorithms*. Morgan Kaufmann.
- March, J. G. (1991). Exploration and exploitation in organizational learning. *Organization Science*, 2(1), 71–87.
- Marr, D. (1982). *Vision: A computational investigation into the human representation and processing of visual information*. Freeman and Company: W. H.
- Mason, W. A., Jones, A., & Goldstone, R. L. (2008). Propagation of innovations in networked groups. *Journal of Experimental Psychology: General*, 137(3), 422.
- Mason, W., & Watts, D. J. (2012). Collaborative learning in networks. *Proceedings of the National Academy of Sciences*, 109(3), 764–769.
- Miller, B., & Steyvers, M. (2011). The wisdom of crowds with communication. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 33.
- Misra, V., Goyal, V. K., & Varshney, L. R. (2011). Distributed scalar quantization for computing: High-resolution analysis and extensions. *IEEE Transactions on Information Theory*, 57(8), 5298–5325.
- Molavi, P., Tahbaz-Salehi, A., & Jadbabaie, A. (2018). A theory of non-Bayesian social learning. *Econometrica*, 86(2), 445–490.
- Mossel, E., Neeman, J., & Tamuz, O. (2014). Majority dynamics and aggregation of information in social networks. *Autonomous Agents and Multi-Agent Systems*, 28(3), 408–429.
- Mueller-Frank, M. (2013). A general framework for rational learning in social networks. *Theoretical Economics*, 8(1), 1–40.
- Navarro, D. J., Perfors, A., Kary, A., Brown, S. D., & Donkin, C. (2018). When extremists win: Cultural transmission via iterated learning when populations are heterogeneous. *Cognitive Science*, 42(7), 2108–2149.
- Nishihara, R., Murray, I., & Adams, R. P. (2014). Parallel MCMC with generalized elliptical slice sampling. *The Journal of Machine Learning Research*, 15(1), 2087–2112.
- Nowak, M. A. (2006). *Evolutionary dynamics: Exploring the equations of life*. Harvard University Press.
- Pan, W., Althuler, Y., & Pentland, A. (2012). Decoding social influence and the wisdom of the crowd in financial trading network. In *International Conference on Social Computing*, 203–209.
- Payne, J. W. (1976). Task complexity and contingent processing in decision making: An information search and protocol analysis. *Organizational Behavior and Human Performance*, 16(2), 366–387.
- Pérez, T., Zamora, J., & Eguíluz, V. M. (2016). Collective intelligence: Aggregation of information from neighbors in a guessing game. *PLoS One*, 11(4), Article e0153586.
- Pérez-Escudero, A., & De Polavieja, G. G. (2011). Collective animal behavior from Bayesian estimation and probability matching. *PLoS Computational Biology*, 7(11), Article e1002282.
- Pratt, S. C., Sumpter, D. J. T., Mallon, E. B., & Franks, N. R. (2005). An agent-based model of collective nest choice by the ant *Temnothorax albigipennis*. *Animal Behaviour*, 70(5), 1023–1036.
- Rahimian, M. A., & Jadbabaie, A. (2017). Bayesian learning without recall. *IEEE Transactions on Signal and Information Processing over Networks*, 3(3), 592–606.
- Rendell, L., Boyd, R., Cownden, D., Enquist, M., Eriksson, K., Feldman, M. W., ... Laland, K. N. (2010). Why copy others? Insights from the social learning strategies tournament. *Science*, 328(5975), 208–213.
- Sanborn, A., & Griffiths, T. L. (2008). Markov chain Monte Carlo with people. In *Advances in Neural Information Processing Systems* (pp. 1265–1272).
- Seeley, T. D., & Buhrman, S. C. (1999). Group decision making in swarms of honey bees. *Behavioral Ecology and Sociobiology*, 45(1), 19–31.
- Shamir, O. (2014). Fundamental limits of online and distributed algorithms for statistical learning and estimation. In *Advances in Neural Information Processing Systems*, 163–171.
- Smith, N. J., Goodman, N., & Frank, M. (2013). Learning and using language via recursive pragmatic reasoning about other agents. *Advances in Neural Information Processing Systems*, 3039–3047.
- Surowiecki, J. (2005). *The wisdom of crowds*. Anchor.
- Tamuz, O., & Tessler, R. J. (2015). Majority dynamics and the retention of information. *Israel Journal of Mathematics*, 206(1), 483–507.
- Tenenbaum, J. B., Kemp, C., Griffiths, T. L., & Goodman, N. D. (2011). How to grow a mind: statistics, structure, and abstraction. *Science*, 331(6022), 1279–1285.
- Thompson, B., & Griffiths, T. L. (2019). Inductive biases constrain cumulative cultural evolution. *Proceedings of the 41st Annual Conference of the Cognitive Science Society*, 41, 1111–1117.
- Thompson, W. R. (1933). On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika*, 285–294.
- Toyokawa, W., Kim, H.-r., & Kameda, T. (2014). Human collective intelligence under dual exploration-exploitation dilemmas. *PLoS One*, 9(4), Article e95789.
- Vul, E., Goodman, N., Griffiths, T. L., & Tenenbaum, J. B. (2014). One and done? Optimal decisions from very few samples. *Cognitive Science*, 38(4), 599–637.
- Wang, P., Paranamana, P., & Shafto, P. (2019). Generalizing the theory of cooperative inference. In *The 22nd International Conference on Artificial Intelligence and Statistics* (pp. 1841–1850).
- Whalen, A., Griffiths, T. L., & Buchsbaum, D. (2018). Sensitivity to shared information in social learning. *Cognitive Science*, 42(1), 168–187.
- Wisdom, T. N., Song, X., & Goldstone, R. L. (2013). Social learning strategies in networked groups. *Cognitive Science*, 37(8), 1383–1425.
- Yang, S. C.-H., Yu, Y., Givchi, A., Wang, P., Vong, W. K., & Shafto, P. (2018). Optimal cooperative inference. In *International Conference on Artificial Intelligence and Statistics* (pp. 376–385).