

LEARNING FROM NEIGHBORS ABOUT A CHANGING STATE

KRISHNA DASARATHA, BENJAMIN GOLUB, AND NIR HAK

ABSTRACT. Agents learn about a changing state using private signals and past actions of neighbors in a network. Bayesian learning in equilibrium yields a DeGroot-style learning dynamic, where agents use social information simply by averaging neighbors’ recent estimates, with time-invariant weights. We examine when a community can aggregate information well, responding quickly to recent changes. A key sufficient condition for good aggregation is that each individual’s neighbors have sufficiently different types of private information. In contrast, when signals are homogeneous, aggregation is suboptimal on any network. Behavioral variations of the model demonstrate that achieving good aggregation requires a sophisticated response to correlations in neighbors’ actions. Finally, we find that an agent’s social influence is much more sensitive to the precision of her private signal than in the DeGroot benchmark.

(KD) YALE UNIVERSITY. EMAIL: KRISHNADASARATHA@GMAIL.COM

(BG) NORTHWESTERN UNIVERSITY. EMAIL: BENJAMIN.GOLUB@NORTHWESTERN.EDU

(NH) UBER TECHNOLOGIES, INC. EMAIL: NIRHAK@GMAIL.COM

Date: April 29, 2022.

We gratefully acknowledge funding from Pershing Square Fund for Research on the Foundations of Human Behavior at Harvard, the Rockefeller Foundation (Dasaratha), and the National Science Foundation under grants SES-1847860 and SES-1629446 (Golub), as well as the hospitality of the economics department at the University of Pennsylvania during a key phase of this work. Hershdeep Chopra, Yixi Jiang, Joey Feffer, and Rithvik Rao provided excellent research assistance. For valuable comments we are grateful to (in random order) Erik Madsen, Alireza Tahbaz-Salehi, Drew Fudenberg, Matt Elliott, Nageeb Ali, Tomasz Strzalecki, Alex Frankel, Jeroen Swinkels, Margaret Meyer, Philipp Strack, Erik Eyster, Michael Powell, Michael Ostrovsky, Leeat Yariv, Andrea Galeotti, Eric Maskin, Elliot Lipnowski, Jeff Ely, Eddie Dekel, Annie Liang, Iosif Pinelis, Ozan Candogan, Ariel Rubinstein, Bob Wilson, Omer Tamuz, Jeff Zwiebel, Matthew O. Jackson, Matthew Rabin, and Andrzej Skrzypacz. We also thank Evan Sadler for important conversations early in the project, David Hirshleifer for detailed comments on a draft, and Rohan Dutta and Kevin He for helpful discussions. The paper was greatly improved by advice from four anonymous referees and the editor, Adam Szeidl. All errors are our own.

1. INTRODUCTION

People learn from others about conditions relevant to decisions they have to make. For instance, students who are about to start their careers learn from the behavior of recent graduates. In many cases, the conditions—for example, the market returns to different specializations—are changing. Thus, welfare depends not on learning a static “state of the world,” but rather on staying up to date with a changing state. The phenomenon of adaptation and responsiveness to new information is central in many economic applications, including in economic development and the study of organizations. When is a group of agents successful, collectively, in adapting efficiently to a changing environment? The answers lie partly in the structure of the social networks that shape agents’ social learning opportunities. Our model is designed to analyze how a group’s adaptability is shaped by the properties of such networks, the inflows of information into society, and the interplay of the two.

We consider overlapping generations of agents who are interested in tracking an unobserved *state* that evolves over time. The state is an AR(1) process: somewhat persistent, but with constant innovations to learn about. Each agent, before making a decision, engages in social learning: she learns the actions of some members of prior generations, which reveal their estimates of the state.¹ The social learning opportunities are embedded in a network, in that one’s network position determines the *neighborhood* of peers whom one observes. Neighborhoods reflect geographic, cultural, organizational, or other kinds of proximity. In addition to social information, agents also receive private signals about the current state, with distributions that may also vary with network position; in particular, some agents may receive more precise information about the state than others.

We give some examples. When a university student begins searching for jobs, she becomes interested in various dimensions of the relevant labor market (e.g., typical wages for someone like her), which naturally vary over time. She uses her own private research (a private signal) but also learns from the choices of others (e.g., recent graduates) who have recently faced a similar problem. The people she learns from depend on her academic specialization, dormitory, extracurricular activities, and so forth: she will predominantly observe predecessors who are “nearby” in these ways.² Similarly, when a new cohort of professionals enters a firm (e.g., a management consultancy or law practice) they learn about the business environment from their seniors. Who works with whom, and therefore who learns from whom, is shaped by the structure of the organization. Beyond heterogeneity in network position, agents differ in the precision of the private signals they can access from outside the network: for example, students in quantitative majors may be better placed to analyze compensation trends.

¹In using an overlapping-generations model we follow a tradition in social learning that includes, for example, the models of Banerjee and Fudenberg (2004) and Wolitzky (2018).

²Sethi and Yildiz (2016) argue that, even without explicit communication costs or constraints, people can end up listening only to some others due to the investments needed to understand sources.

Our first contribution is to develop a network learning model suited to examples such as these. In the model, the state and the population are refreshed over time. The environment in which agents learn is dynamic, but its distribution over time is stationary. This makes equilibria and welfare simple in ways that facilitate the analysis. Indeed, our setup removes a time-dependence inherent in many models of learning, where society accumulates information over time about a fixed state: Those models typically imply rational updating rules that depend on the time elapsed since the learning process started. In terms of outcomes, they often focus on an eventual *rate of learning* about a fixed state.³ In contrast, our model features stationary equilibria with time-invariant learning rules. The outcomes we focus on are learning quality and social influence in such equilibria.

We begin by characterizing equilibria in this model. Bayesians update estimates by taking linear combinations of neighbors’ estimates and their own private information. The weights are endogenously determined because, when each agent extracts information from neighbors’ estimates, the information content of those estimates depends on the neighbors’ learning rules. We characterize these weights and the distributions of behavior in a stationary equilibrium. This characterization enables the various comparative statics and estimation exercises—both under the Bayesian equilibrium benchmark and behavioral alternatives. Equilibria can be numerically computed quickly in networks of thousands of nodes, which makes the model practically useful for structural exercises.

Turning from the model’s features to substantive findings, our second contribution is to analyze classic questions about learning in networks within our framework. We begin by analyzing the steady-state quality of aggregation of social information, and how it depends on signal endowments and network structure. At every point in time, agents use neighbors’ actions to form an estimate of the most recent state before the current period. Our measure of aggregation quality is the accuracy of these estimates. The main finding is that, in large Bayesian populations, an essentially optimal benchmark is achieved in an equilibrium, as long as each individual has access to a set of neighbors that is sufficiently diverse, in the sense of having different signal distributions from each other. A key mechanism behind the value of diverse signal endowments is that it leads to diversity of neighbors’ strategies. This avoids “collinearity problems” in agents’ sources of information, which helps them to construct better statistical estimators of the most recent state and better filter out confounds.

If signal endowments are *not* diverse, then our good-aggregation result does not hold. Indeed, equilibrium learning can be bounded far away from efficiency, even though each agent has access to an unbounded number of observations, each containing independent information. Thus, Bayesian agents who understand the environment perfectly are not guaranteed to be able to aggregate information well. We first make this point in highly symmetric networks,

³See, for instance, Molavi, Tahbaz-Salehi, and Jadbabaie (2018) and Harel, Mossel, Strack, and Tamuz (2021).

where we can show the failure of aggregation is quite severe. We also identify conditions under which this has severe welfare consequences, making agents worse off by an unbounded amount relative to a world with diverse signals. Beyond these highly symmetric networks, it is natural to ask whether diversity of network positions (as opposed to signals) can substitute for diversity of information endowments by making neighbors' equilibrium strategies sufficiently diverse. We show that the answer is no: network asymmetry is a poor substitute for asymmetric signal distributions. In large networks, it is impossible in equilibrium to achieve accuracies of aggregation of the same order as in our positive result under signal diversity.

To achieve good learning when it is possible, agents must respond in a sophisticated way to the correlations in their neighbors' estimates. Thus, the second contrast we emphasize is between Bayesians who are correctly specified about others' behavior, and agents who are too unsophisticated about correlations among their social observations to remove confounds, as in some canonical behavioral learning models (Eyster and Rabin, 2010).⁴ We identify a class of such models in which information aggregation is essentially guaranteed to fall short of good aggregation benchmarks for all agents. The deficiencies of naive learning rules are different from and more severe than those in similar problems with an unchanging state, where naive heuristics can aggregate information very well.⁵

Having discussed the implications of our model for questions of asymptotic learning, we address a second classic question—namely, which agents are influential. We define a notion of steady-state social influence—how an idiosyncratic change in an individual's information affects others' average behavior. This is analogous to exercises familiar from the standard DeGroot model of network learning (where the weights agents place on others are given exogenously). The endogenous determination of weights makes a big difference for how the structure of the environment affects social influence. Relative to the DeGroot model benchmark studied in DeMarzo, Vayanos, and Zweibel (2003), an agent's social influence is much more sensitive to the quality of her private information. On the other hand, just as in the standard benchmark, an agent's influence is approximately proportional to her degree.

Our closing discussion makes two main points. First, some of our theoretical aggregation results use large random graphs. We perform a numerical exercise to show that the main message about information aggregation—diversity of signal types helps learning—remains valid when we calculate equilibria on graphs reflecting real social networks. Second, our analysis generalizes readily to richer models of multidimensional states and signals. As one application of such a generalization, we consider a manager who wishes to facilitate better learning in an organization, and ask what distributions of expertise such a designer would

⁴See also Bala and Goyal (1998), a seminal model of boundedly rational learning rules in networks.

⁵In analogous fixed-state environments where individuals have sufficiently many observations, if everyone uses certain simple and stationary DeGroot-style heuristics (requiring no sophistication about correlations between neighbors' behavior), they can learn the state quite precisely (Golub and Jackson, 2010; Jadbabaie, Molavi, Sandroni, and Tahbaz-Salehi, 2012). A changing state makes such imitative heuristics quite inefficient.

prefer. Our results provide a distinctive rationale for *informational specialization* as a design feature that facilitates good information aggregation.

An Example. We now present a simple example that illustrates our dynamic model, highlights obstacles to learning that distinctively arise in a dynamic environment, and gives a sense of some of the main forces that play a role in our results on the quality of learning.

Consider a particular environment, with a single perfectly informed source S ; many media outlets $M_1, \dots, M_n$ with access to the source as well as some independent private information; and the general public. The public consists of many individuals who learn only from the media outlets. We are interested in how well each member of the public could learn by following many media outlets. More precisely, we consider the example shown in Figure 1.1 and think of P as a generic member of the large public.

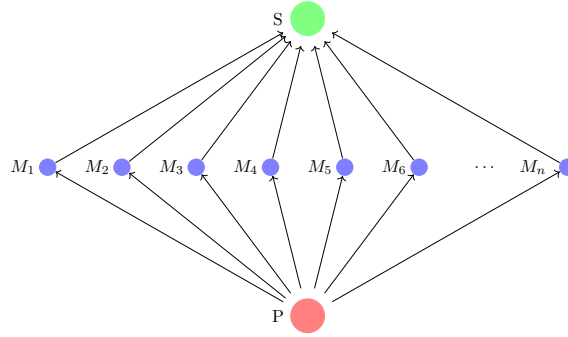


FIGURE 1.1. The network used in the “value of diversity” example

The state θ_t follows a Gaussian random walk: $\theta_t = \theta_{t-1} + \nu_t$, where the innovations ν_t are standard normal. Each period, the source learns the state θ_t and takes an action (which can be thought of simply as making an announcement) that reveals it. The media outlets observe the source’s action from the previous period, which is θ_{t-1} . At each time period, they also receive noisy private signals, $s_{M_i,t} = \theta_t + \eta_{M_i,t}$ with normally distributed, independent, mean-zero errors $\eta_{M_i,t}$. They then announce their posterior means of θ_t , which we denote by $a_{M_i,t}$. The member of the public, in a given period t , makes an estimate based on the observations $a_{M_1,t-1}, \dots, a_{M_n,t-1}$ of media outlets’ actions in the previous period. All agents are short-lived: they see actions in their neighborhoods one period ago, and then they take an action that reveals their posterior belief of the state.

If we had a *fixed* state but the same signals and observation structure, learning would trivially be perfect: media outlets would learn the state from the source and report it to the public. In the dynamic environment, given that P has no signal, she can at best hope to learn θ_{t-1} (and use that to estimate θ_t). Can this benchmark be achieved, and if so, when?

A typical estimate of a media outlet at time t is a linear combination of $s_{M_i,t}$ and θ_{t-1} (the latter being the social information that the media outlets learned from the source). In

particular, the estimate can be expressed as

$$a_{M_i,t} = w_i s_{M_i,t} + (1 - w_i) \theta_{t-1},$$

where the weight w_i on the media outlet's signal is increasing in the precision of that signal. We give the public no private signal, for simplicity only.

Suppose first that the media outlets have identically distributed private signals. Because the member of the public observes many symmetric media outlets, it turns out that her best estimate of the state, $a_{P,t}$, is simply the average of the estimates of the media outlets. Since each of these outlets uses the same weight $w_i = w$ on its private signal, we may write

$$a_{P,t} = w \sum_{i=1}^n \frac{s_{M_i,t-1}}{n} + (1 - w) \theta_{t-2} \approx w \theta_{t-1} + (1 - w) \theta_{t-2}.$$

That is, P 's estimate is an average of media private signals from last period, combined with what the media learned from the source, which tracks the state in the period before that. In the approximate equality, we have used the fact that an average of many private signals is approximately equal to the state, by our assumption of independent errors. No matter how many media outlets there are, and even though each has independent information about θ_{t-1} , the public's beliefs are confounded by older information.

What if, instead, half of the media outlets (say $M_1, \dots, M_{n/2}$) have more precise private signals than the other half, perhaps because these outlets have invested more in covering this topic? The media outlets with more precise signals will then place weight w_A on their private signals, while the media outlets with less precise signals use a smaller weight w_B . We will now argue that a member of the public can extract more information from the media in this setting. In particular, she can first compute the averages of the two groups' actions

$$\begin{aligned} \text{type } A \text{ average} \\ \text{action at time } t-1 &= w_A \sum_{i=1}^{n/2} \frac{s_{M_i,t-1}}{n/2} + (1 - w_A) \theta_{t-2} \approx w_A \theta_{t-1} + (1 - w_A) \theta_{t-2} \\ \text{type } B \text{ average} \\ \text{action at time } t-1 &= w_B \sum_{i=n/2+1}^n \frac{s_{M_i,t-1}}{n/2} + (1 - w_B) \theta_{t-2} \approx w_B \theta_{t-1} + (1 - w_B) \theta_{t-2}. \end{aligned}$$

Then, since $w_A > w_B$, the public knows two distinct linear combinations of θ_{t-1} and θ_{t-2} . The state θ_{t-1} is identified from these. So the member of the public can form a very precise estimate of θ_{t-1} —which, recall, is as well as she can hope to do. The key force is that the two groups of media outlets give different mixes of the old information and the more recent state, and by understanding this, the public can infer both. Indeed, to recover θ_{t-1} , the public puts a negative weight on the B group actions, which allows it to subtract off old information and focus on the recent state, θ_{t-1} . One can show that if, in contrast, agents are naive, e.g., if they think that all of the estimates of the media are uncorrelated (or only mildly correlated)

conditional on the state, they will put positive weights on their observations and will again be bounded from learning the state.

This illustration relied on a particular network with several special features: a very “central” source, one-directional links, and no communication among the media outlets or public. We will show that the same considerations determine learning quality in a large class of random networks in which agents have many neighbors, with complex connections among them. Quite generally, diversity of signal endowments in their neighborhoods allows agents to concentrate on new developments in the state while filtering out old, less relevant information and thus estimate the changing state as accurately as physical constraints allow.

Outline. Section 2 sets up the basic model and discusses its interpretation. Section 3 defines our equilibrium concept and shows that equilibria exist. Section 4 reports our main theoretical results on the quality of information aggregation. In Section 5, we discuss learning outcomes under a variety of non-Bayesian models. Section 6 defines and analyzes social influence. Section 7 relates our model and results to the social learning literature. In Section 8, we describe our numerical exercise with network data from Indian villages and discuss a simple extension to multi-dimensional states to interpret our results on signal diversity.

2. MODEL

State of the world. There is a discrete set of instants, $\mathcal{T} = \mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$. At each time $t \in \mathcal{T}$, there is a *state*, a random variable θ_t taking values in $\mathbb{R}$. This state evolves as an AR(1) stochastic process. That is,

$$\theta_{t+1} = \rho\theta_t + \nu_{t+1}, \quad (2.1)$$

where ρ is a constant with $0 < |\rho| \leq 1$ and $\nu_{t+1} \sim \mathcal{N}(0, \sigma_\nu^2)$ are independent innovations. When $|\rho| < 1$ we have the explicit formula

$$\theta_t = \sum_{\ell=0}^{\infty} \rho^\ell \nu_{t-\ell},$$

and thus the state at any time t has the stationary distribution $\theta_t \sim \mathcal{N}\left(0, \frac{\sigma_\nu^2}{1-\rho^2}\right)$. We make the normalization throughout that innovations have variance 1, i.e., $\sigma_\nu = 1$.

As an alternative, we will also sometimes consider a specification with a starting time, $\mathcal{T} = \mathbb{Z}_{\geq 0} = \{0, 1, 2, \dots\}$, where the state process can be defined as in (2.1) starting at time 0 with some specified distribution for θ_0 .

Information and observations. The set of *nodes* is $N = \{1, 2, \dots, n\}$. Each node i can be thought of as a location, and is associated with a set $N_i \subseteq N$ of nodes that i can observe, called its *neighborhood*.⁶

⁶For all results, a node i ’s neighborhood can, but need not, include i itself.

Each node is populated by a sequence of *agents* in overlapping generations. At each time t , there is a node- i agent, labeled (i, t) , who takes that node's action $a_{i,t}$. This agent is born at time $t - m$ at a certain node and has m periods to observe the actions taken around her before she acts. Thus, when taking her action, the agent (i, t) knows $a_{j,t-\ell}$ for all nodes $j \in N_i$ and lags $\ell \in \{1, 2, \dots, m\}$. We call m the *memory*; it reflects how many periods of actions in her neighborhood an agent passively observes before acting. One interpretation is that a node corresponds to a role in an organization. A worker in that role has some time to observe colleagues in related roles before choosing a once-and-for-all action herself. Much of our analysis is done for an arbitrary finite m ; we view the restriction to finite memory as useful for avoiding technical complications, but because m can be arbitrarily large, this restriction has little substantive content.⁷

In addition to social information from her neighborhood, each agent also sees a private signal,

$$s_{i,t} = \theta_t + \eta_{i,t},$$

where the error term $\eta_{i,t} \sim \mathcal{N}(0, \sigma_i^2)$ has a variance $\sigma_i^2 > 0$ that depends on the node but not on the time period. All the errors $\eta_{i,t}$ and state innovations ν_t are independent of one another. An agent's information is a vector consisting of her private signal and all of her social observations. An important special case will be $m = 1$, where agents observe only one period of others' behavior before acting themselves, so that the agent's information is $(s_{i,t}, (a_{j,t-1})_{j \in N_i})$.

The observation structure is common knowledge, as is the informational environment (i.e., the joint distribution of all exogenous random variables). The network $G = (N, E)$ is the set of nodes N together with the (fixed) set of *links* E , defined as the subset of pairs $(i, j) \in N \times N$ such that $j \in N_i$.

An *environment* is specified by $(G, \boldsymbol{\sigma})$, where $\boldsymbol{\sigma} = (\sigma_i)_{i \in N}$ is the profile of signal variances.

Preferences and best responses. As stated above, in each period t , each agent (i, t) makes her once-and-for-all choice $a_{i,t} \in \mathbb{R}$, seeking to make this action close to the current state. Utility is given by

$$u_{i,t}(a_{i,t}) = -\mathbb{E}[(a_{i,t} - \theta_t)^2]. \quad (2.2)$$

By a standard fact about squared-error loss functions, given the distribution of $(\mathbf{a}_{N_i, t-\ell})_{\ell=1}^m$, the optimal choice of agent (i, t) is to set

$$a_{i,t} = \mathbb{E}[\theta_t \mid \underbrace{s_{i,t}, (\mathbf{a}_{N_i, t-\ell})_{\ell=1}^m}_{i\text{'s information}}]. \quad (2.3)$$

⁷It is worth noting that even when the memory m is small, observed actions can indirectly incorporate signals from much further in the past.

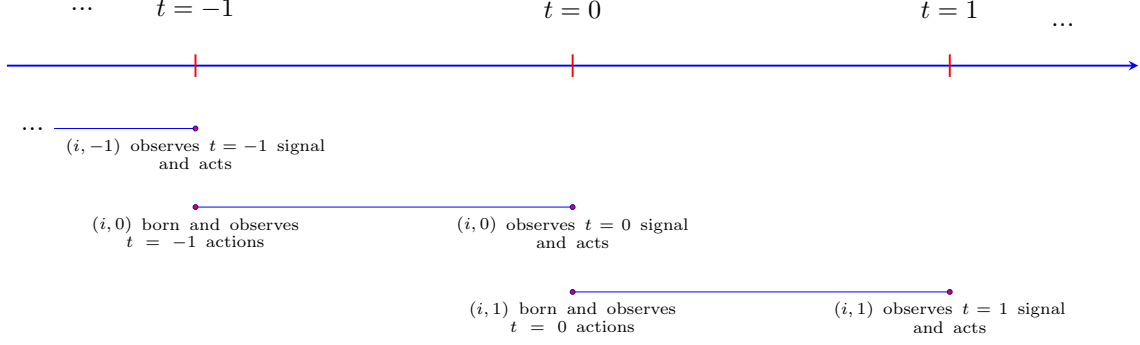


FIGURE 2.1. An illustration of the overlapping generations structure of the model for $m = 1$. At time $t - 1$, agent (i, t) is born and observes actions taken at time $t - 1$ in her neighborhood. Then, at time t , she observes her private signal $s_{i,t}$ and takes her action $a_{i,t}$.

Here the notation $\mathbf{a}_{N_i,t'}$ refers to the vector $(a_{j,t'})_{j \in N_i}$ of time- t' actions in the agent's neighborhood. An action can be interpreted as an agent's estimate of the state, and we will sometimes use this terminology.

The conditional expectation (2.3) depends, of course, on the prior of agent (i, t) about θ_t , which, under correctly specified beliefs, has distribution $\theta_t \sim \mathcal{N}\left(0, \frac{\sigma_\nu^2}{1-\rho^2}\right)$. We actually allow the prior to be any normal distribution or a uniform improper prior.⁸ It saves on notation to analyze the case where all agents have improper priors. Because actions under a normal prior are related to actions under the improper prior by a simple linear bijection—and thus have the same information content for other agents—all results immediately extend to the general case.

The doubly-infinite time axis introduces some subtleties into the definition of strategy profiles; complete details are formalized in Appendix E.

3. UPDATING AND EQUILIBRIUM

In this section we study agents' learning behavior and present a notion of stationary equilibrium. We begin with the canonical case of Bayesian agents with correct models of others' behavior; we study other behavioral assumptions in Section 5 below.

3.1. Best-response behavior. The first step is to analyze optimal updating behavior in response to others' strategies. A strategy of an agent is *linear* if the action taken is a linear function of the variables in her information set. We will analyze agents' best responses to linear strategies, showing that they are linear and computing them explicitly.

⁸We take priors, like the information structure and network, to be common knowledge.

Suppose predecessors have played linear strategies up to time t .⁹ Then we can express each action up until time t as a weighted summation of past signals $s_{i,t}$, for various i and t . Since these signals themselves are linear combinations of the Gaussian innovations ν_t and signal errors $\eta_{i,t}$, the joint distribution of $(a_{i,t-\ell} - \theta_t)_{i \in N, \ell \geq 1}$ is multivariate Gaussian. It follows that $\mathbb{E}[\theta_t \mid s_{i,t}, (\mathbf{a}_{N_i,t-\ell})_{\ell=1}^m]$ is a linear function of $s_{i,t}$ and $(\mathbf{a}_{N_i,t-\ell})_{\ell=1}^m$. The rest of this subsection analyzes this conditional expectation.

3.1.1. Covariance matrices. The optimal weights for an agent to place on past actions depends on their distribution. It will be useful to look at the *errors* $\rho^\ell a_{i,t-\ell} - \theta_t$ of past actions as predictors of the time- t state. Here $\rho^\ell a_{i,t-\ell}$ is the conditional expectation of θ_t given $a_{i,t-\ell}$, and the error is the difference between this prediction and θ_t . Given a linear strategy profile played up until time t , let $\mathbf{V}_t$ be the covariance matrix of the vector of errors $(\rho^\ell a_{i,t-\ell} - \theta_t)_{i \in N, 0 \leq \ell \leq m-1}$. (This covariance matrix has dimensions $nm \times nm$.) The entries of $\mathbf{V}_t$ are denoted by $V_{ij,t}$. In the case $m = 1$, this is simply the covariance matrix $\mathbf{V}_t = \text{Cov}((a_{i,t} - \theta_t)_{i \in N})$.

3.1.2. Best-response weights. A strategy profile is a best response if the weights each agent places on her observations maximize her utility, (2.2), i.e. minimize the squared error of her action. We now characterize such weights in terms of the covariance matrices we have defined. Consider an agent at time t , and suppose some linear strategy profile has been played up until time t . Let $\mathbf{V}_{N_i,t-1}$ be a sub-matrix of $\mathbf{V}_{t-1}$ that contains only the rows and columns corresponding to neighbors¹⁰ of i and consider the following covariance matrix constructed from all of (i, t) 's observations, including her private signal $s_{i,t}$:

$$\mathbf{C}_{i,t-1} = \begin{pmatrix} & & & 0 \\ & \mathbf{V}_{N_i,t-1} & & 0 \\ & & & \vdots \\ 0 & 0 & \dots & \sigma_i^2 \end{pmatrix}.$$

⁹We will discuss this below in the context of our equilibrium concept; but one immediate motivation is that, in the model with a starting time, where $\mathcal{T} = \mathbb{Z}_{\geq 0}$, Bayesian agents' updating at $t = 0$ is a single-agent problem where optimal behavior is a linear function of own signals only, and thus the hypothesis holds. At later times it holds by induction.

¹⁰Explicitly, $\mathbf{V}_{N_i,t-1}$ are the covariances of $(\rho^\ell a_{j,t-\ell} - \theta_t)$ for all $j \in N_i$ and $\ell \in \{1, \dots, m\}$.

Conditional on observations $(\mathbf{a}_{N_i,t-\ell})_{\ell=1}^m$ and $s_{i,t}$, agent (i, t) 's posterior belief about the state θ_t is a normal distribution with a mean that can be calculated as

$$\mathbb{E}[\theta_t \mid s_{i,t}, (\mathbf{a}_{N_i,t-\ell})_{\ell=1}^m] = \underbrace{\frac{\mathbf{1}^T \mathbf{C}_{i,t-1}^{-1}}{\mathbf{1}^T \mathbf{C}_{i,t-1}^{-1} \mathbf{1}}}_{\text{agent } i\text{'s weights}}} \cdot \underbrace{\begin{pmatrix} \rho \mathbf{a}_{N_i,t-1} \\ \vdots \\ \rho^m \mathbf{a}_{N_i,t-m} \\ s_{i,t} \end{pmatrix}}_{\text{agent } i\text{'s observations}}. \quad (3.1)$$

(see Example 4.4 of Kay (1993)). By equation (2.3), this mean is the action that (i, t) plays. Expression (3.1) is a linear combination of the agent's signal and the observed actions; the weights in this linear combination depend on the matrix $\mathbf{V}_{t-1}$, but *not* on realizations of any random variables. In (3.1) we use our assumption of an improper prior.¹¹

We denote by $(\mathbf{W}_t, \mathbf{w}_t^s)$ a *weight profile* in period t , with $\mathbf{w}_t^s \in \mathbb{R}^n$ being the weights agents place on their private signals and $\mathbf{W}_t$ being the weights they place on their other information.

3.1.3. The evolution of covariance matrices under best-response behavior. Assuming agents best-respond according to the optimal weights just described in (3.1), we can compute the resulting next-period covariance matrix $\mathbf{V}_t$ from the previous covariance matrix. This defines a map $\Phi : \mathcal{V} \rightarrow \mathcal{V}$, given by

$$\Phi : \mathbf{V}_{t-1} \mapsto \mathbf{V}_t. \quad (3.2)$$

This map gives the basic dynamics of the model: how an arbitrary variance-covariance matrix $\mathbf{V}_{t-1}$ maps to a new one when all agents best-respond to $\mathbf{V}_{t-1}$. The variance-covariance matrix $\mathbf{V}_{t-1}$ (along with parameters of the model) determines (i) the weights agents place on their observations, via (3.1), and (ii) the distributions of the random variables that are being combined in this operation. This yields the deterministic updating dynamic Φ . A consequence is that the weights agents place on observations are (commonly) known, and do not depend on any random realizations.

Example 1. We compute the map Φ explicitly in the case $m = 1$. We refer to the weight agent i places on $\rho a_{j,t-1}$ as $W_{ij,t}$ and the weight on $s_{i,t}$, her private signal, as $w_{i,t}^s$. Note we have, from (3.1) above, explicit expressions for these weights. Then

$$[\Phi(\mathbf{V})]_{ii} = (w_i^s)^2 \sigma_i^2 + \sum W_{ik} W_{ik'} (\rho^2 V_{kk'} + 1) \quad \text{and} \quad [\Phi(\mathbf{V})]_{ij} = \sum W_{ik} W_{jk'} (\rho^2 V_{kk'} + 1). \quad (3.3)$$

3.2. Stationary equilibrium in linear strategies. We will now turn our attention to *stationary equilibria in linear strategies*—ones in which all agents' strategies are linear with

¹¹As we have mentioned, this is for convenience and without loss of generality. Our analysis applies equally to any proper normal prior for θ_t : To get an agent's estimate of θ_t , the formula in (3.1) would simply be averaged with a constant term accounting for the prior, and everyone could invert this deterministic operation to recover the same information from others' actions.

time-invariant coefficients—though, of course, we will allow all agents to consider deviating to arbitrary strategies, including non-linear ones. Once we establish the existence of such equilibria, we will use the word *equilibrium* to refer to one of these unless otherwise noted.

A reason for focusing on equilibria in linear strategies comes from noting that, in the variant of the model with a starting time (i.e., the case $\mathcal{T} = \mathbb{Z}_{\geq 0}$) agents begin by using only private signals, and they do this linearly. After that, inductively applying the reasoning of Section 3.1, best-responses are linear at all future times. Taking time to extend infinitely backward is an idealization that allows us to focus on exactly stationary behavior.

We now show the existence of stationary equilibria in linear strategies.

Proposition 1. A stationary equilibrium in linear strategies exists, and is associated with a covariance matrix $\hat{\mathbf{V}}$ such that $\Phi(\hat{\mathbf{V}}) = \hat{\mathbf{V}}$.

The proof appears in Appendix A.

At such an equilibrium, the covariance matrix $\mathbf{V}_t$ and all agent strategies are time-invariant. Actions are linear combinations of observations with stationary weights (which we denote by $\hat{W}_{ij}$ and $\hat{w}_i^s$). The form of these rules has some resemblance to static equilibrium notions studied in the rational expectations literature (e.g., Vives, 1993; Babus and Kondor, 2018; Lambert, Ostrovsky, and Panov, 2018; Mossel, Mueller-Frank, Sly, and Tamuz, 2020). It also has a similar form to the DeGroot (1974) and Friedkin and Johnsen (1997) updating rules, typically imposed as behavioral heuristics. In our dynamic environment, such a solution emerges as a steady state.

3.2.1. Proof sketch for the existence result. The goal is to apply the Brouwer fixed-point theorem to show there is a covariance matrix $\hat{\mathbf{V}}$ that remains unchanged under updating. To find a convex, compact set to which we can apply the fixed-point theorem, we use the fact that when agents best respond to any beliefs about prior actions, all action variances are bounded above and bounded away from zero below. This is because all agents' actions must be at least as precise in estimating θ_t as their private signals, and cannot be more precise than estimates given perfect knowledge of θ_{t-1} combined with the private signal. This establishes bounds on action variances. The Cauchy-Schwartz inequality then bounds covariances in terms of the corresponding variances. All matrices respecting these bounds constitute a compact, convex set containing the image of Φ . This and the continuity of Φ allow us to apply the Brouwer fixed-point theorem.

3.2.2. Other remarks. In the case of $m = 1$, we can use the formula of Example 1, equation (3.3), to write the fixed-point condition $\Phi(\hat{\mathbf{V}}) = \hat{\mathbf{V}}$ explicitly. More generally, for any m , we can obtain a formula in terms of $\hat{\mathbf{V}}$ for the weights $\hat{W}_{ij}$ and $\hat{w}_i^s$ in the best response to $\hat{\mathbf{V}}$, and use this to describe the equilibrium $\hat{\mathbf{V}}_{ij}$ as solving a system of polynomial equations. These equations typically have large degree and cannot be solved analytically except in very simple

cases, but they can readily be used to solve for equilibria numerically. A related feature of the model is that standard methods can easily be applied to estimate it and test hypotheses within it (see Appendix D for details).

The main insight is that we can find equilibria by studying action covariances; this idea applies equally to many extensions and variations of our basic model. We give two examples: (1) We assume that agents observe neighbors perfectly, but one could define other observation structures. For instance, agents could observe actions with noise, or they could observe some set of linear combinations of neighbors' actions with noise. Similarly, agents could be observing predecessors' actions for heterogeneous durations before acting (i.e., node-specific m). (2) We assume agents are Bayesian and best-respond rationally to the distribution of actions, but the same proof would also show that equilibria exist under other behavioral rules (see Section 5.1).¹²

Proposition 1 shows that there exists a stationary linear equilibrium. We show later, as part of Proposition 2, that there is a *unique* stationary linear equilibrium in networks having a particular structure. In general, uniqueness of the equilibrium is an open question that we leave for future work.¹³ In Section 4.2.3 and Appendix G, we discuss a nonstationary variant of the model which has a unique equilibrium, and relate it to our main model.

How much information does each agent need to play her equilibrium strategy? In a stationary equilibrium, she only needs to know the steady-state variance-covariance matrix $\hat{\mathbf{V}}_{N_i}$ in her neighborhood. Then her problem of inferring θ_{t-1} becomes essentially a linear regression problem. If historical empirical data on neighbors' error variances and covariances are available, then $\hat{\mathbf{V}}_{N_i}$ can be estimated from such data.

4. HOW GOOD IS INFORMATION AGGREGATION IN EQUILIBRIUM?

In this section we analyze the quality of information aggregation in stationary equilibrium.

We begin with a definition. Recall that an agent at time t uses social information to form a belief about θ_{t-1} , which is a sufficient statistic for the past in the agent's decision problem. The conditional expectation of θ_{t-1} that an agent (i, t) forms based on social information is called her *social signal* and denoted by $r_{i,t}$:

$$r_{i,t} = \mathbb{E}[\theta_{t-1} \mid (\mathbf{a}_{N_i,t-\ell})_{\ell=1}^m].$$

¹²What is important in the proof is that actions depend continuously on the covariance structure of an agent's observations; the action variances are uniformly bounded under the rule agents play; and there is a decaying dependence of behavior on the very distant past.

¹³We have checked numerically that Φ is not, in general, a contraction in any of the usual norms (entrywise sup, Euclidean operator norm, etc.). In computing equilibria numerically for many examples, we have not been able to find a case of equilibrium multiplicity. Indeed, in all of our numerical examples, repeatedly applying Φ to an initial covariance matrix gives the same fixed point for any starting conditions.

Definition 1. For a given strategy profile, define the *aggregation error* $\kappa_{i,t}^2 = \text{Var}(r_{i,t} - \theta_{t-1})$ to be the variance of the social signal (equivalently, the expected squared error in the social signal as a prediction of θ_{t-1}).

The aggregation error measures how well an agent can extract information from social observations. Note that agent i 's aggregation error is a monotone transformation of her expected utility.¹⁴ We will be interested in this number, and how low error can be in equilibrium.

The environment features informational externalities: players do not internalize the impact that their learning rules have on others' learning. Consequently, there is no reason to expect outcomes to be efficient in any exact sense. And we have seen that the details of equilibrium in a particular network can be complicated. However, it turns out that much more can be said about the behavior of aggregation errors as neighborhood size (i.e., the number of social observations) grows. In this section, we study the asymptotic efficiency of information aggregation. We give conditions under which aggregation error decays as quickly as physically possible, and different conditions under which it remains far from efficient levels even when agents have arbitrarily many observations. We discuss the case $m = 1$ for simplicity but the reasoning extends easily to other values of m .

A benchmark lower bound on aggregation error. A first observation is a lower bound on the aggregation error (in terms of an asymptotic rate as a function of a node's degree) under *any* behavior of agents. This establishes a benchmark relative to which we can assess equilibrium outcomes.

Let d_i denote the out-degree of an agent i .

Fact 1. Fix $\rho \in (-1, 1)$ as well as upper and lower bounds for private signal variances, so that $\sigma_i^2 \in [\underline{\sigma}^2, \bar{\sigma}^2]$ for all i . On any network and for all strategy profiles, we have $\kappa_{i,t}^2 \geq c/d_i$ for all i and t , where c is a constant that depends only on $\rho, \underline{\sigma}^2$, and $\bar{\sigma}^2$.

The lower bound is reminiscent of the central limit theorem: if an agent had d_i conditionally independent noisy signals about θ_{t-1} (e.g., by observing neighbors' private signals directly), then the variance of her estimate would be of order $1/d_i$. Fact 1 notes that it is not possible for aggregation errors to decay (as a function of degree) any faster than that.

For an intuition, imagine that an agent sees neighbors' private signals (not just actions) one period after they are received, and all other private signals two periods after they are received; this clearly gives an upper bound on the quality of aggregation given physical communication constraints. The information that is two periods old cannot be very informative about θ_{t-1} because of the movement in the state from period $t-2$ to $t-1$; a large constant number z of signals about θ_{t-1} would be better. Thus, a lower bound on aggregation error is given by the

¹⁴In fact, for any decision dependent on θ_t , an agent is better off with a lower value of $\kappa_{i,t}^2$. This is a consequence of the fact that unidimensional Gaussian signals can be Blackwell ordered by their precision.

error that could be achieved with $d_i + z$ independent signals about θ_{t-1} of the best possible precision ($\underline{\sigma}^{-2}$). The bound follows from these observations.

Outline of results: When is aggregation comparable to the benchmark? Fact 1 places a lower bound on aggregation error given the physical constraints. Even efficient learning could not do better than this bound. We will examine when equilibrium learning can achieve aggregation of similar quality. More precisely, we ask when there is a stationary equilibrium where the aggregation error at node i satisfies $\widehat{\kappa}_i^2 \leq C/d_i$ for all i , for some constant C .

In Section 4.2 we establish a good-aggregation result: outcomes comparable to the benchmark are achieved in equilibrium in a class of networks. The key condition enabling the asymptotically efficient equilibrium outcome is called *signal diversity*: each individual has access to enough neighbors with multiple different kinds of private signals. The fact that neighbors use private information differently turns out to give the agents enough power to identify θ_{t-1} with equilibrium aggregation error that decays at a rate matching the lower bound up to a multiplicative constant.

In Section 4.3, we turn to negative results. Without signal diversity, equilibrium aggregation can be extremely bad. Our first negative result shows that when signals are exchangeable, it may be that the aggregation error $\widehat{\kappa}_i^2$ does not approach zero in any equilibrium no matter how large neighborhoods are, though a social planner could achieve good aggregation by prescribing different updating weights. We prove this in highly symmetric networks. Once we move away from such networks, one might ask whether diversity in individuals' network positions could play a role analogous to signal diversity and enable approximately efficient learning. Our second negative result shows that this is impossible. When signals are homogeneous and all agents' degrees in network G_n are bounded by $d(n)$ (where $d(n)$ is any unbounded sequence) then in any equilibrium, aggregation errors cannot vanish at rate $C/d(n)$ for any $C > 0$ as the network grows.

4.1. Distributions of networks and signals. For our good-aggregation result, we study large populations and specify two aspects of the environment: *network distributions* and *signal distributions*. In terms of network distributions, we work with a standard type of random network model—a stochastic block model (see, e.g., Holland, Laskey, and Leinhardt, 1983). It makes the structure of equilibrium tractable while also allowing us to capture rich heterogeneity in network positions. We also specify *signal distributions*: how signal precisions are allocated to agents, in a way that may depend on network position. We now formalize these two primitives of the model and state the assumptions we work with.

Fix a set of *network types* $k \in \mathcal{K} = \{1, 2, \dots, K\}$. For each pair of network types, there is a given probability $p_{kk'}$ that each agent of network type k has a link to each agent of network type k' . An assumption we maintain on these probabilities is that each network type k

observes at least one network type (possibly k itself) with positive probability. There is also a vector $(\alpha_1, \dots, \alpha_K)$ of *population shares* of each type, which we assume are all positive. Jointly, $(p_{kk'})_{k,k' \in \mathcal{K}}$ and α specify the network distribution. These parameters can encode differences in expected degree and also features such as homophily (where some groups of types are linked to each other more densely than to others).

We next define signal distributions, which describe the allocation of signal variances to network types. Fix a finite set $\mathbb{S}$ of private signal variances, which we call signal types.¹⁵ We let $q_{k\tau}$ be the share of agents of network type k with signal type τ ; $(q_{k\tau})_{k \in \mathcal{K}, \tau \in \mathbb{S}}$ defines the signal distribution.

Let the nodes in network n be partitioned into the network types $N_n^1, N_n^2, \dots, N_n^K$, with the cardinality $|N_n^k|$ equal to $\lfloor \alpha_k n \rfloor$ or $\lceil \alpha_k n \rceil$ (rounding so that there are n agents in the network). We (deterministically) set the signal variances σ_i^2 equal to elements of $\mathbb{S}$ in accordance with the signal shares (again rounding as needed). Let $(G_n)_{n=1}^\infty$ be a sequence of directed or undirected random networks with these nodes, so that $i \in N_n^k$ and $j \in N_n^{k'}$ are linked with probability $p_{kk'}$; these realizations are all independent.

A *stochastic block model* D is specified by the linking probabilities $(p_{kk'})_{k,k' \in \mathcal{K}}$, the type shares α , and the signal distribution $(q_{k\tau})_{k \in \mathcal{K}, \tau \in \mathbb{S}}$. We let $(G_n(D), \sigma_n(D))$ denote a realization of the network and signal variances under a given stochastic block model. We say that a signal type τ is represented in a network type k if $q_{k\tau} > 0$.

Definition 2. A stochastic block model satisfies *signal diversity* if each network type has a positive probability of linking with at least one network type containing two distinct signal types.

We will discuss stochastic block models that satisfy this condition as well as ones that do not, and show that the condition is pivotal for information-aggregation.

4.2. Good aggregation under diverse signals. Our first main result is that signal diversity is sufficient for good aggregation in the networks described in the previous section. Aggregation error decays at a rate C/d_i for each node i independently of the structural properties of the network.

We first define a notion of good aggregation for an agent in terms of a bound on that agent's aggregation error.

Definition 3. Given $\varepsilon > 0$, we say that agent i achieves the ε -aggregation benchmark in a given equilibrium if the aggregation error satisfies $\widehat{\kappa}_i^2 \leq \varepsilon$.

We say an event (indexed by n) occurs *asymptotically almost surely* if the probability of the event converges to 1 as $n \rightarrow \infty$.

¹⁵The assumptions of finitely many signal types and network types are for technical convenience only, and could be relaxed.

Theorem 1. *Fix any stochastic block model D satisfying signal diversity. There exists $C > 0$ such that asymptotically almost surely the environment $(G_n(D), \sigma_n(D))$ has an equilibrium where all agents achieve the C/n -aggregation benchmark.*

So for large enough n , society is very likely to aggregate information very well. The uncertainty in this statement is over the network, as there is always a small probability of a realized network which prevents learning (e.g., an agent has no neighbors). We give an outline of the argument next, and the proof appears in Appendix B.

The constant C in the theorem statement can depend on the stochastic block model D . However, given any compact set of stochastic block models D , we can choose a single $C > 0$ for which the result holds uniformly across D .¹⁶ Thus, the theorem can be applied without detailed information on how the random graphs are generated, as long as some bounds are known about which models are possible.

4.2.1. Discussion of the proof. To give intuition for Theorem 1, we first describe why the theorem holds on the complete network¹⁷ with two signal types A and B in the $m = 1$ case. This echoes the intuition of the example in the introduction. We then discuss the challenges involved in generalizing the result to our general stochastic block model networks, and the techniques we use to overcome those challenges.

Consider a time- t agent, (i, t) . Recall that the social signal $r_{i,t}$ is the optimal estimate of θ_{t-1} based on the actions (i, t) has observed in her neighborhood. In the complete network, all players have the same social signal, which we call r_t .¹⁸

At any equilibrium, each agent's action is a weighted average of her private signal and this social signal.

$$a_{i,t} = \hat{w}_i^s s_{i,t} + (1 - \hat{w}_i^s) r_t. \quad (4.1)$$

The weights on the observations $s_{i,t}$ and $\rho a_{j,t-1}$ (which constitute the social signal) sum to 1, because the optimal action is an unbiased estimate of θ_t . The weight $\hat{w}_i^s$ on the private signal depends on the precision of this signal relative to the social signal. We call the weights used by agents of the two distinct signal types $\hat{w}_A^s$ and $\hat{w}_B^s$. Suppose signal type A is more accurate than signal type B , so that $\hat{w}_A^s > \hat{w}_B^s$.

Now, turning our attention to the next period of updating, observe that each time- $(t+1)$ agent can compute two averages of the time- t actions—one for each type. Using (4.1) to

¹⁶The reason is that the distribution of aggregation errors is upper hemicontinuous in model parameters, so if the desired bounds hold for each point in a compact set, they can be made uniform.

¹⁷Note this is a special case of the stochastic block model.

¹⁸In particular, agent (i, t) sees everyone's past action, including the one taken last period at the same node, $a_{i,t-1}$.

rewrite $a_{i,t}$ and then plugging in $s_{i,t} = \theta_t + \eta_{i,t}$:

$$\begin{aligned} \text{type } A \text{ average} &= \frac{1}{n_A} \sum_{i: \sigma_i^2 = \sigma_A^2} a_{i,t} = \hat{w}_A^s \theta_t + (1 - \hat{w}_A^s) r_t + O(n^{-1/2}), \\ \text{action at time } t & \\ \text{type } B \text{ average} &= \frac{1}{n_B} \sum_{i: \sigma_i^2 = \sigma_B^2} a_{i,t} = \hat{w}_B^s \theta_t + (1 - \hat{w}_B^s) r_t + O(n^{-1/2}). \\ \text{action at time } t & \end{aligned}$$

Here n_A and n_B denote the numbers of agents of each type (recalling we assumed each type is a positive share of the population size, n). The $O(n^{-1/2})$ error terms come from the average signal noises $\eta_{i,t}$ of agents in each group; the bound holds with high probability by the central limit theorem. In other words, each time- $(t+1)$ agent can obtain precise estimates of two different convex combinations of θ_t and r_t . Because the two weights, $\hat{w}_A^s$ and $\hat{w}_B^s$, are distinct, she can approximately solve for θ_t as a linear combination of the average actions taken by each type she observes (up to signal error). It follows that in the equilibrium we are considering, the agent must have an estimate at least as precise as what she can obtain by the strategy we have described, and will thus be very close the benchmark. The estimator of θ_t in this strategy places negative weight on $\frac{1}{n_B} \sum_{i: \sigma_i^2 = \sigma_B^2} a_{i,t-1}$, thus *anti-imitating* the agents of signal type B—those with the less precise private signal. Proposition 3 below implies that the actual equilibrium in which agents learn will also have agents anti-imitating others.

To use the same approach in general, we need to show that each individual observes a large number of neighbors of at least two signal types who also have similar social signals. More precisely, the proof shows that agents with the same network type have highly correlated social signals. Showing this is much more subtle than it was in the above illustration. In general, the social signals in an arbitrary network realization are endogenous objects that depend to some extent on all the links.

A key insight allowing us to overcome this difficulty is a useful general fact about sufficiently dense stochastic block models: despite a lot of idiosyncratic randomness in direct connections, the law of large numbers implies the number of paths of length two between any i of type k and j of type k' going through agents of type k'' is nearly determined by the types k , k' , and k'' , with a small relative error.¹⁹ We can leverage this to deduce some important facts about the updating map Φ (recall Section 3.1.3) in the realized random network, and specifically about the evolution of social signals.

In particular, if we look at the set of covariance matrices where all social signals are close to perfect, we can show that the composition $\Phi^2 := \Phi \circ \Phi$ maps this set to itself. In other words, if social signals are very precise, then they will remain very precise two periods later. If the two-step path counts were determined by types exactly, the reasoning of the complete graph

¹⁹For simplicity we begin by illustrating the argument in a random graph family where the number of two-step paths is nearly deterministic. The argument extends to a larger class of models where the same property applies to longer paths, as we discuss in the next subsection.

example above would generalize, with all neighbors of i of the same type updating in the same way. We show that despite the path counts being known only approximately, the desired conclusion holds. This is nontrivial because the weights agents use in their updating—and thus the evolution of social signals—could depend sensitively on realized network structure; small relative errors could matter. A key step is to develop results on matrix perturbations to show that small relative changes in the network do not affect Φ^2 too much. A fixed-point theorem then implies there is a fixed point of Φ^2 in the set of outcomes with very precise social signals. With some further analysis we can deduce that this implies the existence of an equilibrium (corresponding to a fixed point of Φ) with nearly perfect aggregation.

4.2.2. Sparser random graphs. In the random graphs we have defined in Section 4.1, the group-level linking probabilities $(p_{kk'})$ are, for simplicity, held fixed as n grows. This yields expected degrees that grow linearly in the population size, which may not be the desired asymptotic model. We can, however, establish versions of our results in a class of models much more flexible with respect to degrees. While it is important to have neighborhoods “large enough” (i.e., growing in n) to permit the application of laws of large numbers, their rate of growth can be considerably slower than linear. For example, our proof can be extended directly to degrees that scale as n^α for any $\alpha > 0$ to show that asymptotically almost surely, there exists an equilibrium where the C/n^α -aggregation benchmark is achieved for all agents. Instead of studying Φ^2 and two-step paths, one can apply the same analysis to the L -fold composition Φ^L , which reflects L -step paths. In order to do this, one uses the fact that for L larger than $1/\alpha$, the number of paths of length L between any two nodes is determined by their types with a small relative error. Extending our proof above, we can then characterize the behavior of Φ^L and then deduce the claimed aggregation property for Φ .

4.2.3. The good-aggregation outcome as a unique prediction. The theorem above says good aggregation is supported in an equilibrium but does not state that this is the unique equilibrium outcome. To deal with this issue, we study the alternative model with $\mathcal{T} = \mathbb{Z}_{\geq 0}$ (where agents begin with only their own signals and then best-respond to the previous distribution of behavior at each time). We show that its long-run outcomes get arbitrarily close to the good-aggregation equilibrium of Theorem 1 as $n \rightarrow \infty$, under the same conditions. Thus, even if there were other equilibria of the stationary model, they could not be approached via the natural iterative procedure coming from the $\mathcal{T} = \mathbb{Z}_{\geq 0}$ model. Formal statements and details are in Appendix G.

4.3. Aggregation under homogeneous signals. Having established conditions for good aggregation under signal diversity, we now explore what happens without signal diversity. Our general message is that aggregation is worse.

To gain an intuition for this, note that it is essential to the argument described in the previous subsection that different agents have different signal precisions. Recall the complete graph case. From the perspective of an agent $(i, t + 1)$, the fact that type A and type B neighbors place different weights on the social signal r_t allows the agent to avoid a collinearity problem and separate θ_t from a confound. In that example, if type A and B agents had the same signal types, they would use the same weights, and our agent trying to learn from them would face a collinearity problem.

We begin by studying graphs having a symmetric structure and show that learning outcomes are necessarily bounded very far from good aggregation. We then turn to arbitrary large graphs and prove a lower bound on aggregation error that implies the homogeneous-signals regime has, quite generally, worse outcomes for some agents than those achieved by everyone in our good-aggregation result.

4.3.1. Aggregation in graphs with symmetric neighbors.

Definition 4. A network G has *symmetric neighbors* if whenever $j, j' \in N_i$ for some i , then $N_j = N_{j'}$.

In the undirected case, the graphs with symmetric neighbors are the complete network and complete bipartite networks.²⁰ For directed graphs, the condition allows a larger variety of networks.

Proposition 2. Consider a sequence $(G_n)_{n=1}^\infty$ of strongly connected graphs with symmetric neighbors. Assume that all signal variances are equal, and that $m = 1$. Then there is a unique equilibrium on each G_n . Moreover, there exists $\varepsilon > 0$ such that the ε -aggregation benchmark is not achieved by any agent i at this equilibrium for any n .

All agents have non-vanishing aggregation errors at the unique equilibrium. So all agents learn poorly compared to the diverse signals case. The proof of this proposition, and the proofs of all subsequent results, appear in Appendix F.

This failure of good aggregation is not due simply to a lack of sufficient information in the environment: On the complete graph with exchangeable (i.e., non-diverse) signals, a social planner who set weights for all agents could achieve ε -aggregation for any $\varepsilon > 0$ when n is large. See Appendix I for a formal statement, proof and numerical results.²¹ In this sense, the social learning externalities are quite severe: a fairly small change in weights for each individual could yield a very large benefit in a world of homogeneous signal types.

We now give intuition for Proposition 2. In a graph with symmetric neighbors and homogeneous signals, in the unique equilibrium,²² actions of any agent's neighbors are exchangeable.

²⁰These are both special cases of our stochastic block model from Section 4.2, so Theorem 1 applies to these network structures when signal diversity is satisfied.

²¹We thank Alireza Tahbaz-Salehi for suggesting this analysis.

²²The proof of the proposition establishes uniqueness by showing that Φ is a contraction in a suitable sense.

So Bayesian estimates (and thus actions) must weight all neighbors equally. This prevents the sort of inference of θ_t that occurred with diverse signals. This is easiest to see on the complete graph, where *all* observations are exchangeable. So, in any equilibrium, each agent's action at time $t + 1$ is equal to a weighted average of her own signal and the average action $\frac{1}{|N_i|} \sum_{j \in N_i} a_{j,t}$:

$$a_{i,t+1} = \hat{w}_i^s s_{i,t+1} + (1 - \hat{w}_i^s) \frac{1}{|N_i|} \sum_{j \in N_i} a_{j,t}. \quad (4.2)$$

By iteratively using this equation, we can see that actions must place substantial weight on the average of signals from, e.g., two periods ago, and indeed farther back. Note that all signals $s_{j,t'}$ at past times t' take the form $\theta_{t'} + \eta_{i,t'}$. Thus, although the effect of *signal errors* $\eta_{i,t'}$ vanishes (by averaging) as n grows large, the correlated error from *past changes in the state* $\nu_{t'}$ never “washes out” of estimates, and this is what prevents vanishing aggregation errors.

The bad-aggregation result as stated applies to exactly homogeneous signal types only. In fact, in finite networks we need sufficiently heterogeneous signals to avoid the bad-learning obstruction; this is illustrated in Appendix C. In Section 4.4 we discuss the welfare implications of this failure of aggregation.

As a consequence of Theorem 1 and Proposition 2, we can give an example where making one node's private information less precise helps all agents.

Corollary 1. There exists a network G and an agent $i \in G$ such that increasing σ_i^2 gives a Pareto improvement at the unique equilibrium.

To prove the corollary, we consider the complete graph with homogeneous signals and large n . By Proposition 2, all agents have non-vanishing aggregation errors. If we instead give agent 1 a very uninformative signal, all players can anti-imitate agent 1 and achieve vanishing aggregation errors. When the signals at the initial configuration are sufficiently imprecise, this gives a Pareto improvement. There are also examples where severing links in the observational network can yield a Pareto improvement, as reported in an earlier version of the present paper (Dasaratha, Golub, and Hak, 2018).

4.3.2. Aggregation in arbitrary networks. Section 4.3.1 showed aggregation errors are non-vanishing when signal endowments and neighborhoods are symmetric. A natural question is whether asymmetry in network positions can substitute for asymmetry in signal endowments. In Section 4.2 the key point was that different neighbors' actions were informative about different linear combinations of θ_t and θ_{t-1} , and this permitted filtering. Perhaps different network positions can achieve the same effect?

We thus move to *arbitrary* networks and show a weaker but much more general result. Consider any sequence of equilibria on any networks with symmetric signal endowments.

Our result here is that no equilibrium achieves C/n -aggregation for almost all agents, no matter what C is. In particular, this implies that the rate of learning is slower than at the good-learning equilibrium with diversity of signal endowments from Theorem 1. Moreover, if degrees are bounded above by some $\bar{d}$ growing at rate slower than n , we prove the stronger statement that no equilibrium achieves $C/\bar{d}$ -aggregation for almost all agents.

Theorem 2. *Let $C > 0$. Let $(G_n)_{n=1}^\infty$ be an arbitrary sequence of networks and suppose all private signals have variance σ^2 . If all agents' in-degrees and out-degrees are bounded above by some $\bar{d}(n) \rightarrow \infty$, then in any sequence of equilibria, $\hat{\kappa}_i^2 > C/\bar{d}(n)$ for a non-vanishing fraction of agents i .*

In addition to considering arbitrary networks, we allow the memory m to be arbitrary (yet finite). Because the assumptions are much weaker, we obtain a weaker conclusion than in Proposition 2. While Proposition 2 shows that aggregation errors are non-vanishing, the theorem shows that aggregation errors cannot vanish quickly, but does not rule out aggregation errors vanishing more slowly.

The basic intuition is that to avoid putting substantial weight on θ_{t-2} , an agent at time t must anti-imitate some neighbors. If all or almost all neighbors achieve C/n -aggregation for some C and have identical types of private signals, there is not much diversity among neighbors. So more and more anti-imitation is needed as n grows large in the sense that the total positive weight and total negative weight on neighbors both grow large. But then the contribution to the agent's variance from neighbors' private signal errors cannot vanish quickly.

We can combine Theorems 1 and 2 to compare the value of signal diversity and network diversity. With diversity of signal endowments, there exists $C > 0$ such that asymptotically almost surely there is a good-learning equilibrium achieving the C/n -aggregation benchmark for all agents under the stochastic block model. With exchangeable signals, it is not possible to find equilibria so close to the benchmark for *any* sequence of networks.

Theorem 2 shows that network heterogeneity cannot improve learning outcomes as much as signal heterogeneity. The formal results establish a gap between asymptotic rates of convergence (in contrast to the stronger bad-learning result for networks with symmetric neighbors). Section 8.1 shows that in real-world (highly asymmetric) social networks, signal heterogeneity improves learning outcomes much more than choosing a very favorable network structure but homogeneous signals.

4.4. The welfare loss associated with homogeneity. The results derived so far in this section show that there is a qualitative difference in how well agents are able to infer recent states across the homogeneous and heterogeneous signal settings. How important is this difference for welfare? We illustrate next that the welfare loss associated with signal homogeneity can be arbitrarily severe.

To gain an intuition for this, note that with homogeneous signals, period- t actions are confounded by previous states. These confounds include θ_{t-2} , which all $t-1$ agents use in the same way (as illustrated in the example of the introduction). But the confounds also include θ_{t-3} , which could not be filtered out by $t-1$ agents, and so forth. The more weight agents place on social information (i.e., the more informative the past is), the more severe this confounding is. If the state is highly persistent and private signals are not very precise, then the confounds from periods even very long ago are substantial. The following corollary quantifies this effect.

Corollary 2. Consider a complete graph with all signal variances equal to σ^2 , and let $m = 1$. Then, in any symmetric strategy profile,

$$\text{Var}(a_{i,t} - \theta_t) \geq \frac{(1 - \hat{w}^s)^2}{1 - (1 - \hat{w}^s)^2 \rho^2}.$$

As $\rho \rightarrow 1$ from below and $\sigma^{-2} \rightarrow 0$ and, agent i 's action error in the unique equilibrium converges to infinity. Moreover, this convergence is uniform in n .²³

In contrast, recall that our main positive results show the C/n -aggregation benchmark would be achieved with signal heterogeneity. When this benchmark is achieved, each individual obtains a variance $\text{Var}(a_{i,t} - \theta_t)$ that is at worst 1 if n is large enough.²⁴ This bound on variance does not depend on σ^2 or ρ . This implies *welfare can be arbitrarily worse in environments with signal homogeneity compared to ones with heterogeneity*. In particular, the corollary guarantees that we can choose (σ^{-2}, ρ) so that the error is arbitrarily large, uniformly in n . If we modify the signal distribution so that it is heterogeneous²⁵ then error variance will be at most 1 for large enough n .

In large complete graphs with homogeneous signals, we can explicitly characterize the limit action variance (and therefore welfare). Let V^∞ denote the limit, as n grows large, of $\text{Var}(a_{i,t} - \theta_t)$. Let Cov^∞ denote the limit covariance of any two agents' errors. By direct computation using equation (3.3), these can be seen to be related by the following equations, which have a unique solution:

$$V^\infty = \frac{1}{\sigma^{-2} + (\rho^2 \text{Cov}^\infty + 1)^{-1}}, \quad \text{Cov}^\infty = \frac{(\rho^2 \text{Cov}^\infty + 1)^{-1}}{[\sigma^{-2} + (\rho^2 \text{Cov}^\infty + 1)^{-1}]^2}. \quad (4.3)$$

These equations also let us extend Corollary 2 beyond the complete graph. The variance and covariance in (4.3) describe the limits of all variances and covariances in any graph with symmetric neighbors where degrees tend uniformly to infinity. Indeed, it can be deduced (as

²³For any $\underline{v}$, there are $\underline{\rho} < 1$ and $\bar{\sigma}^{-2} > 0$ such that if $\rho > \underline{\rho}$ and $\sigma^{-2} < \bar{\sigma}^{-2}$, then $\text{Var}(a_{i,t} - \theta_t) \geq \underline{v}$ for all n .

²⁴Note the agent can use the estimate of last period's state, which has an error of order C/n . If the agent simply guessed this estimate, then she would achieve $\text{Var}(\theta_t - \theta_{t-1}) = 1$, since the state innovation has variance 1. Combining this with her private signal does strictly better than this.

²⁵For example, by making half the agents' signals strictly worse.

in the proof of Corollary 2) that agents' actions are equal to an appropriately discounted sum of past $\theta_{t-\ell}$, up to error terms (arising from $\eta_{i,t-\ell}$) that vanish asymptotically. As $\sigma^{-2} \rightarrow 0$ and $\rho \rightarrow 1$ from below, equations (4.3) show that Cov^∞ and therefore also V^∞ diverge to infinity. This shows the welfare loss from homogeneity can also be arbitrarily severe in graphs with symmetric neighbors and large degrees.

5. THE IMPORTANCE OF UNDERSTANDING CORRELATIONS

In the positive result on achieving the C/n -aggregation benchmark (Theorem 1), a key aspect of the argument involved agents filtering out confounding information from their neighbors' estimates—i.e., responding in a sophisticated way to the correlation structure of those estimates. In this section, we demonstrate that this sort of behavior is essential for nearly perfect aggregation, and that more naively imitative heuristics yield outcomes far from the benchmark. Empirical studies have found evidence (depending on the setting and the subjects) consistent with both equilibrium behavior and naive inference in the presence of correlated observations (e.g., [Eyster, Rabin, and Weizsacker, 2015](#); [Dasaratha and He, 2019](#); [Enke and Zimmermann, 2019](#)).

We begin with a canonical model of agents who do not account for correlations among their neighbors' estimates conditional on the state, and show by example that naive agents achieve much worse learning than Bayesian agents, and thus have non-vanishing aggregation errors. We then formalize the idea that accounting for correlations in neighbors' actions is crucial to reaching the benchmark. This is done by demonstrating a general lack of asymptotic learning by agents who use imitative strategies, rather than filtering in a sophisticated way. Finally, we show that even in fixed, finite networks, any positive weights chosen by optimizing agents will be Pareto-dominated.

5.1. Naive agents. In this part we introduce agents who misunderstand the distribution of the signals they are facing and who therefore do not update as Bayesians with a correct understanding of their environment. We consider a particular form of misspecification that simplifies solving for equilibria analytically:²⁶

Definition 5. We call an agent *naive* if she believes that all neighbors choose actions equal to their private signals and maximizes her expected utility given these incorrect beliefs.

Equivalently, a naive agent believes her neighbors all have empty neighborhoods. This is the analogue, in our model, of “best-response trailing naive inference” ([Eyster and Rabin, 2010](#)). So naive agents understand that their neighbors' actions from the previous period are estimates of θ_{t-1} . But they think each such estimate is independent given the state, and

²⁶There are a number of possible variants of our behavioral assumption, and it is straightforward to numerically study alternative specifications of behavior in our model ([Alatas et al., 2016](#) consider one such variant).

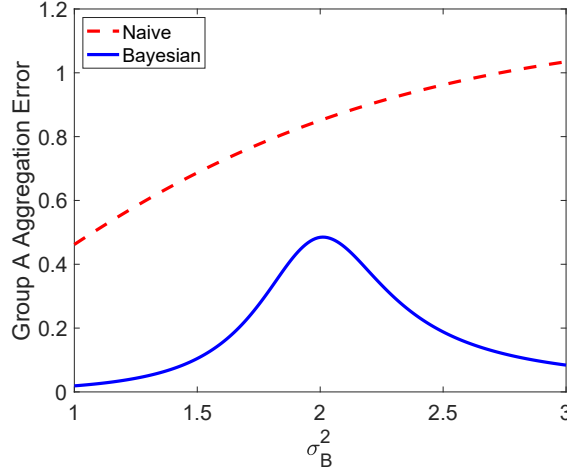


FIGURE 5.1. Bayesian and naive learning on a complete graph and $n = 600$ agents divided into two equal sized groups. The plot shows the aggregation error in group A as group B's private signal variance varies, fixing group A's private signal variance at $\sigma_A^2 = 2$.

that the precision of the estimate is equal to the signal precision of the corresponding agent. They then play their expectation of the state given this misspecified theory of others' play.

In Figure 5.1, we compare Bayesian and naive learning outcomes. We consider a complete network with 600 agents and $\alpha = 0.9$. Half of agents have signal variance $\sigma_A^2 = 2$, while we vary the signal variance σ_B^2 of the remaining agents. The figure shows the average social signal variance for the group of agents with private signal variance $\sigma_A^2 = 2$. The figure suggests that naive agents learn substantially worse than rational agents, whether signals are diverse or not. We prove this holds for general stochastic block models and provide formulas for variances under naive learning in Appendix H.

5.2. More general learning rules: Understanding correlation is essential for good aggregation. We now show more generally that a sophisticated response to correlation is needed to achieve vanishing aggregation errors on any sequence of growing networks. To this end, we make the following definition:

Definition 6. The *steady state* associated with weights $\mathbf{W}$ and $\mathbf{w}^s$ is the (unique) covariance matrix $\mathbf{V}^*$ such that if actions have a variance-covariance matrix given by $\mathbf{V}_t = \mathbf{V}^*$ and next-period actions are set using weights $(\mathbf{W}, \mathbf{w}^s)$, then $\mathbf{V}_{t+1} = \mathbf{V}^*$ as well.

In this definition of steady state, instead of best-responding to others' actual distributions of play, agents use exogenous weights $\mathbf{W}$ in all periods.

By a straightforward application of the contraction mapping theorem, if agents use any non-negative weights under which covariances remain bounded at all times, there is a unique steady state.

Consider a sequence of networks $(G_n)_{n=1}^\infty$ with n agents in G_n .

Proposition 3. Fix any sequence of steady states under non-negative weights on G_n . Suppose that all private signal variances are bounded below by $\underline{\sigma}^2 > 0$ and that all agents place weight at most $\bar{w} < 1$ on their private signals. Then there is an $\varepsilon > 0$ such that, for all n , the ε -aggregation benchmark is not achieved by any agent i at steady state.

The essential idea is that at time $t+1$ observed time- t actions all put weight on actions from period $t-1$, which causes θ_{t-1} to have a (positive weight) contribution to all observed actions. Agents do not know θ_{t-1} and, with positive weights, cannot take any linear combination that would recover it. Even with a very large number of observations, this confound prevents agents from learning yesterday’s state precisely.

To see why the weights on private signals must be bounded away from one, note that an individual agent could learn well without adjusting for correlations by observing many autarkic agents who simply report their private signals. But in this case, all of these autarkic agents would have non-vanishing aggregation errors. If we did not impose a bound on private signal weights, learning would fail in a weaker sense: in that case, *some* agent must still fail to achieve the ε -aggregation benchmark for small enough ε .

On undirected networks, the proposition implies that aggregation errors do not vanish under naive inference or under various other specifications of non-Bayesian inference. Moreover, the same argument shows that in any sequence of *Bayesian* equilibria on undirected networks where all agents use positive weights, no agent can learn well.

5.3. Without anti-imitation, outcomes are Pareto-inefficient. The previous section argued that anti-imitation is critical to achieving vanishing aggregation errors. We now show that even in small networks, where that benchmark is not relevant, any equilibrium without anti-imitation is Pareto-inefficient relative to another steady state. This result complements our asymptotic analysis by showing a different sense (relevant for small networks) in which anti-imitation is necessary to make the best use of information.

Proposition 4. Suppose the network G is strongly connected and some agent has more than one neighbor. Given any naive equilibrium or any Bayesian equilibrium where all weights are positive, the action variances at that equilibrium are Pareto-dominated by action variances at another steady state.

The basic argument behind Proposition 4 is that if agents place marginally more weight on their private signals, this introduces more independent information that eventually benefits everyone. In a review of sequential learning experiments, [Weizsäcker \(2010\)](#) finds that subjects weight their private signals more heavily than is optimal (given the empirical behavior of others they observe). Proposition 4 implies that in our environment with optimizing agents, it is actually welfare-improving for individuals to “overweight” their own information relative to best-response behavior.

The condition on equilibrium weights says that no agent anti-imitates any of her neighbors. This assumption makes the analysis tractable, but we believe the basic force also works in finite networks with some anti-imitation. In the proof in Appendix F, we state and prove a more general result where weights are non-negative but need not all arise from Bayesian or naive updating.

Proof sketch. The idea of the proof of the rational case is to begin at the steady state and then marginally shift the rational agent’s weights toward her private signal. By the envelope theorem, this means agents’ actions are less correlated but not significantly worse in the next period. We show that if all agents continue using these new weights, the decreased correlation eventually benefits everyone. In the last step, we use the absence of anti-imitation, which implies that the updating function associated with agents using fixed (as opposed to best-response) weights is monotonic in terms of the variances of guesses. To first order, some covariances decrease while others do not change after one period under the new weights. Monotonicity of the updating function and strong connectedness imply that eventually all agents’ variances decrease.

The proof in the naive case is simpler. Here a naive agent is overconfident about the quality of her social information, so she would benefit from shifting some weight from her social information to her signal. This deviation also reduces her correlation with other agents, so it is Pareto-improving.

6. SOCIAL INFLUENCE

It is often of interest how influential agents are in terms of affecting aggregate behavior and how this depends on the environment. For example, these issues are a focus of studies including DeMarzo, Vayanos, and Zweibel (2003) and Golub and Jackson (2010) in the DeGroot model with an unchanging state. In this section, we define a suitable analogue of social influence for our dynamic environment. We then study how an agent’s influence depends on her signal precision and degree. We find that, relative to benchmark results from the DeGroot model, signal precisions are more important, while social connectedness plays a similar role in both models.

6.1. Defining social influence. We define the *total influence* of node i in a stationary equilibrium to be the total weight that all actions place on the private signal of agent (i, t) . The total influence measures the total increase in actions if $s_{i,t}$ increases by 1 (due to an

idiosyncratic shock, say).²⁷ At equilibrium, the total influence of i is:

$$\text{TI}(i) = \sum_{j \in N} \sum_{k=0}^{\infty} (\rho \widehat{W})_{ji}^k \widehat{w}_i^s.$$

This expression for total influence is a version of Katz-Bonacich centrality with respect to the matrix $\widehat{W}$ of weights. The decay parameter is the persistence ρ of the of the AR(1) state process.

We define the *social influence* of i to be the total weight that all actions *in future periods* place on the private signal of agent (i, t) . At equilibrium, the social influence of i is:

$$\text{SI}(i) = \sum_{j \in N} \sum_{k=1}^{\infty} (\rho \widehat{W})_{ji}^k \widehat{w}_i^s = \text{TI}(i) - \widehat{w}_i^s.$$

The social influence measures the influence of an agent at node i on other agents. Social influence and total influence differ only by the weight (i, t) places on her own current signal, because an agent's signal realization does not affect others' actions in the same period. Note that agent i 's social influence depends on the weight $\widehat{w}_i^s$ she places on her own signal as well as the weights agents place on each others' actions.

The next result, which follows from Proposition 1 on equilibrium existence, shows that the summation that defines social influence is guaranteed to converge at equilibrium, which makes social influence (and similarly total influence) well-defined.²⁸

Proposition 5. The social influence $\text{SI}(i)$ is well-defined at any equilibrium and is equal to $[\mathbf{1}'(I - \rho \widehat{W})^{-1} - \mathbf{1}]_i \widehat{w}_i^s$.

We show this as follows: if social influence did not converge, some agents would have actions with very large variances (because their actions would depend sensitively on small idiosyncratic shocks). But then these agents would have simple deviations that would improve their accuracy, such as following their private signals. So this could not happen in equilibrium.

In general, social influence can be negative: an agent's net effect on others can be in the opposite direction of her signal.

6.2. Which agents are influential? We now ask how the social influence $\text{SI}(i)$ of an agent depends on her signal precision and degree. To facilitate the most direct comparison with standard results in models with a fixed state, such as DeMarzo, Vayanos, and Zweibel (2003), we focus on cases where social influences are positive.

To examine the effect of signal precision on social influence, we first study complete networks with $n \geq 2$ agents and two private signal variances: half the agents have more precise

²⁷Note that in a stationary equilibrium, this depends on the node and not the time, so we speak interchangeably of the influence of a node and that of an agent at this node.

²⁸Since $\widehat{W}$ can contain both positive and negative numbers, some of them potentially quite large, it is not immediately obvious that the summation converges.

signals, and the other half have less precise signals. We call the two groups' signal variances σ_A^2 and σ_B^2 and the corresponding agents' social influences $\text{SI}(A)$ and $\text{SI}(B)$. We show that the ratio between the two groups' social influences in equilibrium is larger than the ratio between their signal precisions (whenever the imprecise group has positive social influence).

Proposition 6. On a complete network with $m = 1$ and signal variances $\sigma_A^2 < \sigma_B^2$, in the unique equilibrium it holds that

$$\frac{\text{SI}(A)}{\text{SI}(B)} > \frac{\sigma_A^{-2}}{\sigma_B^{-2}}$$

whenever $\text{SI}(B) > 0$.

The proposition says that increasing a group's precision increases their influence more than proportionately. As we have seen in our main results, if the precision difference is large enough, then it is optimal to place zero or negative weight on the less precise group. The result says that even before this happens, imprecision hurts a group's influence considerably—and, as we will discuss below, more than in benchmark models of social influence.

The proposition assumes the network is complete, but numerical evidence suggests that on other networks, too, agents with more precise signals tend to be much more influential. We simulate a configuration model with $n = 40$ nodes, each with degree $d = 5$.²⁹ Nodes are randomly assigned to a precise signal with variance σ_A^2 or an imprecise signal with variance σ_B^2 (with equal probability).

We are interested in the ratio $\text{SI}(A)/\text{SI}(B)$ in this more complicated environment. If social influence were approximately proportional to precision, then $\text{SI}(A)/\text{SI}(B)$ would be approximately $\sigma_A^{-2}/\sigma_B^{-2}$. To assess by how much the better information of the precise group exceeds this benchmark, we will look at the ratio

$$R_\sigma = \frac{\text{SI}(A)/\text{SI}(B)}{\sigma_A^{-2}/\sigma_B^{-2}}.$$

Table 1 reports this ratio over 100 runs of the simulation model for various pairs of σ_A^2 and σ_B^2 each in the interval $[0.5, 5]$. The entries of the table would be equal to one if influence is proportional to precision. Instead, all off-diagonal entries are greater than one (or negative), meaning social influence depends more (and often much more) on signal precision than in the proportional benchmark.

Having examined how influence depends on precisions, we turn to how it depends on degrees. We again use a configuration model, which allows us to fix any degree distribution, but generate the graphs uniformly conditional on the degrees. We will find that social influence depends less on degree than on precision. We simulate a configuration model with $n = 40$ nodes, each randomly assigned degrees d_A or d_B (with equal probability of each) and with

²⁹This model works by creating n nodes, each with d “stubs” sticking out of it, and then performing a random matching of the stubs to create a graph. See Jackson (2010), Section 4.5.10, for details.

	Precision	σ_B^2									
		0.5	1	1.5	2	2.5	3	3.5	4	4.5	5
σ_A^2	0.5	1	2.14	4.53	9.06	168.09	-29.69	-15.54	-8.22	-7.41	-6.19
	1		1	1.80	3.15	5.62	9.95	26.35	-29.89	-15.23	-13.67
	1.5			1	1.64	2.56	3.80	6.27	10.64	44.97	-31.32
	2				1	1.51	2.22	3.13	4.35	7.05	11.34
	2.5					1	1.43	2.04	2.69	3.88	5.50
	3						1	1.38	1.86	2.46	3.35
	3.5							1	1.34	1.74	2.27
	4								1	1.30	1.67
	4.5									1	1.28
	5										1

TABLE 1. The table shows how far influence ratios are from a benchmark of being proportional to precision. We use a configuration model with a regular network and heterogeneous signal variances; there are $n = 40$ agents and the degree is $d = 5$. Agents are randomly assigned to signal variances σ_A^2 or σ_B^2 . Each entry is computed from 100 runs with persistence $\rho = 0.9$. Each table entry reports the ratio $R_\sigma = \frac{\text{SI}(A)/\text{SI}(B)}{\sigma_A^2/\sigma_B^2}$ for the precision parameters corresponding to that entry.

	Degree	d_B									
		1	2	3	4	5	6	7	8	9	10
d_A	1	1	1.11	1.02	0.96	0.94	0.92	0.96	1.00	1.03	1.09
	2		1	1.02	1.01	0.98	0.96	0.92	0.93	0.91	0.95
	3			1	1.01	1.01	1.01	0.98	0.96	0.94	0.93
	4				1	1.01	1.01	1.00	0.99	0.98	0.96
	5					1	1.01	1.01	1.02	1.01	1.00
	6						1	1.01	1.01	1.01	1.01
	7							1	1.01	1.01	1.02
	8								1	1.01	1.01
	9									1	1.01
	10										1

TABLE 2. The table shows how far influence ratios are from a benchmark of being proportional to degree. We use a configuration model with two possible degrees on $n = 40$ agents with homogeneous signal variance $\sigma^2 = 2$. Agents are randomly assigned to degrees d_A or d_B . Each entry is computed from 100 runs with persistence $\rho = 0.9$. Each table entry reports the ratio $R_d = \frac{\text{SI}(A)/\text{SI}(B)}{d_A/d_B}$ for the degree parameters corresponding to that entry.

$\sigma^2 = 2$ for all agents. Table 2 reports the ratio

$$R_d = \frac{\text{SI}(A)/\text{SI}(B)}{d_A/d_B}.$$

over 100 runs of the simulation model for degrees between 1 and 10. Again, the entries would be equal to one if social influence were proportional to degree. Social influence is indeed approximately proportional to degree: the entries in the table range between 0.91 and 1.11.

Remark 1. A simple intuition explains why social influence depends more on private information than network position. Increasing an agent's private signal precision and her degree both (tend to) make her action more accurate. Increasing private signal precision also implies the agent places more weight on her private information, which is recent and independent of other agents' actions. This provides more reason for others to place weight on her actions,

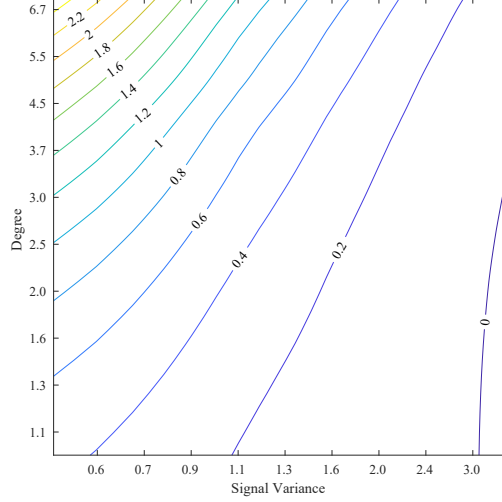


FIGURE 6.1. Level curves for average social influence of agents in a configuration model with 5,000 networks with $n = 40$ agents in each and persistence $\rho = 0.9$. Degrees are chosen uniformly from $\{1, 2, \dots, 7\}$ and private signal variances are chosen uniformly from $\{0.5, 1, \dots, 3.5\}$. The figure shows level curves for average social influence (drawn via cubic interpolation) on a log-log plot.

amplifying the effect of the increased accuracy. In contrast, increasing degree tends to make an agent place more weight on her social information, which is older and more correlated with others. This countervails the effect of increased accuracy, making the agent a less appealing source for others.

The exercises so far varied only one of signal precision or degree, and we now explore how social influence depends on precision and degree jointly. To do so, we compute equilibrium social influences on 5,000 networks with $n = 40$ agents in each. Each agent is randomly assigned a degree chosen uniformly from $\{1, 2, \dots, 7\}$ and a private signal variance chosen uniformly and independently from $\{0.5, 1, \dots, 3.5\}$. Networks are then drawn via the configuration model. Figure 6.1 plots the level curves for average social influence. The steepness of the level curves shows that social influence again depends more on signal variance than degree, especially when signals are less precise.

6.3. Comparison with a DeGroot benchmark. The results above are interesting to compare with what we might expect in canonical network models with a fixed state. A standard benchmark in the networks literature to express influence as a function of network position is the DeGroot model.

To allow us to also consider the effect of signal precision, we consider a fixed-state setting studied by DeMarzo, Vayanos, and Zweibel (2003) as a foundation for DeGroot updating, where weights depend on precisions. Agents start with an improper prior, receive independent normal private signals s_i about the state once, and then each takes an action $a_{i,0}$ equal to her expectation of θ . After this, agents observe their neighbors' actions and take actions $a_{i,1}$,

which are Bayesian expectations of the state θ given their observations. In all subsequent periods $t > 1$, agents observe their neighbors' actions $a_{j,t-1}$ and take actions $a_{i,t}$ as if $a_{j,t-1}$ has the same distribution as j 's private signal (i.e., they naively repeat their optimal strategy from the first period).

One natural measure of social influence is the influence of s_i on the long-run consensus beliefs $\lim_{t \rightarrow \infty} a_{j,t}$. In an undirected, connected, and aperiodic network, this limit exists and the influence of agent i is proportional to her private signal precision σ_i^{-2} and to her degree d_i . Compared to this benchmark, social influence in our changing-state model is more sensitive to signal precision (in the complete graph and in our simulations for configuration models). On the other hand, the impact of degree on social influence is very similar to the DeGroot benchmark—approximately proportional to degree. That is, influence depends more on an agent's private information while the dependence on network position is remarkably similar. The difference between the benchmark and our model is explained in Remark 1.

7. RELATED LITERATURE

Whether decentralized communication can facilitate efficient adaptation to a changing world is a fundamental question in economic theory, related to questions raised by Hayek (1945)³⁰ and of primary interest in some applied problems, e.g., in real business cycle models with consumers and firms learning about evolving states.³¹ Nevertheless, there is relatively little modeling of Bayesian learning of dynamic states in the large literature on social learning and information aggregation in networks, whose most relevant papers we now review.³²

Play in the stationary linear equilibria of our model closely resembles behavior in the DeGroot (1974) model, where agents update by linearly aggregating network neighbors' past estimates, with constant weights on neighbors over time. DeMarzo, Vayanos, and Zweibel (2003), in a Gaussian environment with an unchanging state, derive DeGroot learning as the Bayesian behavior in the first round of communication, and use that as a foundation for a DeGroot rule as a boundedly-rational heuristic. Molavi, Tahbaz-Salehi, and Jadbabaie (2018) have offered new bounded-rationality foundations for the DeGroot rule. Our different environment offers a different foundation for averaging rules with constant weights, as a

³⁰“If ... the economic problem of society is mainly one of rapid adaptation to changes in the particular circumstances of time and place ... there still remains the problem of communicating to [each individual] such further information as he needs.” Hayek's main concern was aggregation of information through markets, but the same questions apply more generally.

³¹See Angeletos and La'O (2010) for a survey of related models that are used to study real business cycles. More recent developments include Angeletos and Lian (2018) and Molavi (2019), with the latter allowing a form of misspecification.

³²For more complete surveys of different parts of this literature, see, among others, Acemoglu and Ozdaglar (2011), Golub and Sadler (2016), and Mossel and Tamuz (2017). See Moscarini, Ottaviani, and Smith (1998) for an early model in a binary-action environment, where it is shown that a changing state can break information cascades.

stationary equilibrium of a stationary environment.³³ Though the updating rule resembles those derived in fixed-state environments, we have stressed that the learning implications are quite different.

Several recent papers in computer science and engineering study dynamic environments similar to ours. [Shahrampour, Rakhlin, and Jadbabaie \(2013\)](#) study an exogenous-weights version, interpreted as a set of Kalman filters under the control of a planner. They bound measures of welfare in terms of the persistence of the state process (ρ) and network invariants, such as the spectral gap. [Frongillo, Schoenebeck, and Tamuz \(2011\)](#) study a $\rho = 1$ model of the state. They characterize the steady-state distribution of behavior for any weights, and calculate equilibrium weights on a complete network, which they show are inefficient. Our Proposition 4 documents a related inefficiency, while the quality of equilibrium learning in large, incomplete networks and social influence in equilibrium are topics not considered in these papers. In economics, [Alatas, Banerjee, Chandrasekhar, Hanna, and Olken \(2016\)](#) perform an empirical exercise in a similar model with a quasi-Bayesian learning rule. Their estimation assumes agents ignore the correlation between social observations, similarly to our naive models.³⁴ Our results show that the degree of rationality can be pivotal for the outcomes of such processes, and provides foundations for structural inference to test various behavioral assumptions.

Our results about when agents learn well are related to two phenomena that have played an important role in the social learning literature. A theme in the social learning literature is that heterogeneity—in agents’ preferences or in whom they observe—can be helpful for learning. One manifestation of this is the phenomenon (typically in sequential social learning models with a fixed state) of *sacrificial lambs*: a fairly sparse set of agents who observe nobody can help everyone else learn well, because their actions are then informative only about their private signals, and unconfounded by an information cascade ([SgROI, 2002](#), [Arieli and Mueller-Frank, 2019](#)). Heterogeneity in preferences can also serve a similar purpose: if preferences have full support, there is a positive probability that preference bias counteracts available social information, causing an agent to follow her private signal ([Goeree, Palfrey, and Rogers, 2006](#), [Lobel and Sadler, 2016](#)). A crucial difference is that our mechanism does not rely on any agents simply revealing their private signals: heterogeneity helps by changing

³³Indeed, agents behaving according to the DeGroot heuristic in other environments might have to do with their experiences in stationary environments where it is closer to optimal.

³⁴The paper’s focus is estimating parameters of social learning rules using data from Indonesian villages, where agents are trying to assess each other’s wealth.

how neighbors use their social information, which in turn aids an agent in inferring a common confound.³⁵

Second, a robust aspect of rational learning in sequential models is the phenomenon of anti-imitation, as discussed, e.g., by [Eyster and Rabin \(2014\)](#). They give general conditions for fully Bayesian agents to anti-imitate in the sequential model. We find that anti-imitation is also an important feature in our dynamic model, and in our context is crucial for good learning. Despite this similarity, there is an important contrast between our environment and standard sequential models. In those models, while rational agents *do* prefer to anti-imitate, in many cases individuals, and society as a whole, could obtain good outcomes using heuristics without any anti-imitation: for instance, by combining the information that can be inferred from a single neighbor with one’s own private signal. [Acemoglu, Dahleh, Lobel, and Ozdaglar \(2011\)](#) and [Lobel and Sadler \(2015\)](#) show that such a heuristic leads to asymptotic learning in a sequential model. Our dynamic learning environment is different, as shown in [Proposition 3](#): to have any hope of approaching good aggregation benchmarks, agents must respond in a sophisticated way, with anti-imitation, to their neighbors’ (correlated) estimates.

8. DISCUSSION AND EXTENSIONS

8.1. Aggregation and its absence without asymptotics: Numerical results. The message of [Section 4](#) is that signal diversity enables good aggregation, and signal homogeneity obstructs it. The theoretical results were asymptotic, and relied on various assumptions about network structure. It is natural to ask whether our main conclusions hold up in realistic finite networks. To analyze this, we numerically study equilibria of our model in actual social networks from Indian villages ([Banerjee, Chandrasekhar, Duflo, and Jackson, 2013](#)). This subsection briefly summarizes our findings; we describe the exercise fully in [Appendix C](#).

We examine the benefits of signal heterogeneity for equilibrium aggregation. The network data are essentially the only empirical input to our exercise.³⁶ Given a network, we compute equilibria using our model and parameters chosen for illustration. We compare two environments that differ in signal allocations: (i) a homogeneous case, with all signal variances set to 2, and (ii) a heterogeneous case, where half of the nodes have a signal variance less than 2 (which we vary) and half of the nodes have a signal variance greater than 2.³⁷

We first compare the value of a good network with the value of heterogeneous signals. Some networks have better learning than others even with homogeneous signals. We define

³⁵A bit farther afield, in [Sethi and Yildiz \(2012\)](#), learning outcomes when two individuals repeatedly learn from each other depend on whether their (heterogeneous) priors are independent or correlated; the common thread is that a natural assumption about agents’ attributes (independent priors in their case) leads to an identification problem. The mechanics are otherwise quite different.

³⁶In particular, we have no data on signal qualities, and simply posit that households without electricity have worse access to external information.

³⁷We choose the larger signal variance so that the average precision in each village is $\frac{1}{2}$, which holds the total inflow of information constant in a sense made precise in the appendix.

the *network-driven variation* in learning to be the standard deviation of learning quality (aggregation error) across villages in the homogeneous case. Our main finding is that increasing the private signal variance for half of the agents by 50% changes social signal error variance by 6.5 times the network-driven variation. That is, introducing this amount of private signal heterogeneity improves learning much more than the most favorable network among the villages.

We also quantify *how much* signal diversity is needed to approach benchmark aggregation in finite networks. Though the asymptotic prediction changes starkly depending on whether signal precisions are identical or not, considerable diversity is actually required to achieve these benefits in a finite network. Starting from homogeneous signals and increasing signal diversity, aggregation error changes very slightly at first. Once the variance of the less precise signal has increased by 50% relative to the starting point, learning quality has moved about halfway to what is achievable with the most extreme signal heterogeneity.

8.2. Multidimensional states and informational specialization. Our formal analysis assumed a one-dimensional state and one-dimensional signals, which varied only in their precisions. Our message about the value of diversity is, however, better interpreted in a mathematically equivalent multidimensional model.

Consider Bayesian agents who learn and communicate about two independent dimensions simultaneously (each one working as in our model). If all agents have equally precise signals about both dimensions, then society may not learn well about either of them. In contrast, if half the agents have superior signals about one dimension and inferior signals about the other (and the other half has the reverse), then society can learn well about both dimensions. Thus, the designer has a strong preference for an organization with informational specialization where some, but not all, agents are expert in a particular dimension.³⁸

Of course, there are many familiar reasons for specialization, in information or any other activity. For instance, it may be that more total information can be collected in this case, or that incentives are easier to provide. Crucially, specialization is valuable in our setting for a reason distinct from all these: it helps agents with their inference problems.

More generally, one can readily extend our model and equilibrium concept to a multi-dimensional state $\theta_t \in \mathbb{R}^d$ and arbitrary Gaussian signals about it, with flexible correlations. We would expect to find suitable generalizations of the basic message that sufficient diversity of neighborhoods (in terms of signal types) facilitates learning. The assumption that agents know neighbors' signal distributions is clearly very helpful for tractability; it would be interesting to consider models in which agents are also uncertain about these distributions.

³⁸This raises important questions about what information agents would acquire, and whom they would choose to observe, which are the focus of a growing literature. For recent papers on this in the context of networks, see Sethi and Yildiz (2016) and Myatt and Wallace (2017), among others.

REFERENCES

- ACEMOGLU, D., M. DAHLEH, I. LOBEL, AND A. OZDAGLAR (2011): “Bayesian Learning in Social Networks,” *The Review of Economic Studies*, 78, 1201–1236.
- ACEMOGLU, D. AND A. OZDAGLAR (2011): “Opinion Dynamics and Learning in Social Networks,” *Dynamic Games and Applications*, 1, 3–49.
- ALATAS, V., A. BANERJEE, A. G. CHANDRASEKHAR, R. HANNA, AND B. A. OLKEN (2016): “Network Structure and the Aggregation of Information: Theory and Evidence from Indonesia,” *The American Economic Review*, 106, 1663–1704.
- ANGELETOS, G.-M. AND J. LA’O (2010): “Noisy Business Cycles,” *NBER Macroeconomics Annual*, 24, 319–378.
- ANGELETOS, G.-M. AND C. LIAN (2018): “Forward Guidance without Common Knowledge,” *American Economic Review*, 108, 2477–2512.
- ARIELI, I. AND M. MUELLER-FRANK (2019): “Multidimensional Social Learning,” *The Review of Economic Studies*, 86, 913–940.
- BABUS, A. AND P. KONDOR (2018): “Trading and Information Diffusion in Over-the-Counter Markets,” *Econometrica*, 86, 1727–1769.
- BALA, V. AND S. GOYAL (1998): “Learning from Neighbours,” *The Review of Economic Studies*, 65, 595–621.
- BANERJEE, A., A. G. CHANDRASEKHAR, E. DUFLO, AND M. O. JACKSON (2013): “The Diffusion of Microfinance,” *Science*, 341, 1236–1239.
- BANERJEE, A. AND D. FUDENBERG (2004): “Word-of-mouth Learning,” *Games and economic behavior*, 46, 1–22.
- DASARATHA, K., B. GOLUB, AND N. HAK (2018): “Bayesian Social Learning in a Dynamic Environment,” *CoRR*, abs/1801.02042.
- DASARATHA, K. AND K. HE (2019): “An Experiment on Network Density and Sequential Learning,” *arXiv preprint arXiv:1909.02220*.
- DEGROOT, M. H. (1974): “Reaching a Consensus,” *Journal of the American Statistical Association*, 69, 118–121.
- DEMARZO, P., D. VAYANOS, AND J. ZWEIBEL (2003): “Persuasion Bias, Social Influence, and Unidimensional Opinions,” *The Quarterly Journal of Economics*, 118, 909–968.
- ENKE, B. AND F. ZIMMERMANN (2019): “Correlation Neglect in Belief Formation,” *The Review of Economic Studies*, 86, 313–332.
- EYSTER, E. AND M. RABIN (2010): “Naive Herding in Rich-information Settings,” *American Economic Journal: Microeconomics*, 2, 221–243.
- (2014): “Extensive Imitation is Irrational and Harmful,” *Quarterly Journal of Economics*, 129, 1861–1898.
- EYSTER, E., M. RABIN, AND G. WEIZSACKER (2015): “An Experiment on Social Mislearning,” *Working Paper*.
- FRIEDKIN, N. E. AND E. C. JOHNSEN (1997): “Social positions in influence networks,” *Social Networks*, 19, 209–222.
- FRONGILLO, R., G. SCHOENEBECK, AND O. TAMUZ (2011): “Social Learning in a Changing World,” *Internet and Network Economics*, 146–157.
- GOEREE, J. K., T. R. PALFREY, AND B. W. ROGERS (2006): “Social Learning with Private and Common Values,” *Economic theory*, 28, 245–264.

- GOLUB, B. AND M. JACKSON (2010): “Naive Learning in Social Networks and the Wisdom of Crowds,” *American Economic Journal: Microeconomics*, 2, 112–149.
- GOLUB, B. AND E. SADLER (2016): “Learning in Social Networks,” in *The Oxford Handbook of the Economics of Networks*, ed. by Y. Bramoullé, A. Galeotti, B. Rogers, and B. Rogers, Oxford University Press, chap. 19, 504–542.
- HAREL, M., E. MOSSEL, P. STRACK, AND O. TAMUZ (2021): “Rational groupthink,” *The Quarterly Journal of Economics*, 136, 621–668.
- HAYEK, F. A. (1945): “The Use of Knowledge in Society,” *The American Economic Review*, 35, 519–530.
- HOLLAND, P. W., K. B. LASKEY, AND S. LEINHARDT (1983): “Stochastic Blockmodels: First Steps,” *Social Networks*, 5, 109–137.
- JACKSON, M. O. (2010): *Social and economic networks*, Princeton university press.
- JADBABAIE, A., P. MOLAVI, A. SANDRONI, AND A. TAHBAZ-SALEHI (2012): “Non-Bayesian Social Learning,” *Games and Economic Behavior*, 76, 210–225.
- KAY, S. M. (1993): *Fundamentals of Statistical Signal Processing*, Prentice Hall PTR.
- LAMBERT, N. S., M. OSTROVSKY, AND M. PANOV (2018): “Strategic Trading in Informationally Complex Environments,” *Econometrica*, 86, 1119–1157.
- LOBEL, I. AND E. SADLER (2015): “Information Diffusion in Networks through Social Learning,” *Theoretical Economics*, 10, 807–851.
- (2016): “Preferences, Homophily, and Social Learning,” *Operations Research*, 64, 564–584.
- MANRESA, E. (2013): “Estimating the Structure of Social Interactions Using Panel Data,” *Working Paper, CEMFI, Madrid*.
- MOLAVI, P. (2019): “Macroeconomics with Learning and Misspecification: A General Theory and Applications,” *Working Paper, Kellogg School of Management, Evanston, IL*.
- MOLAVI, P., A. TAHBAZ-SALEHI, AND A. JADBABAIE (2018): “A Theory of Non-Bayesian Social Learning,” *Econometrica*, 86, 445–490.
- MOSCARINI, G., M. OTTAVIANI, AND L. SMITH (1998): “Social Learning in a Changing World,” *Economic Theory*, 11, 657–665.
- MOSSEL, E., M. MUELLER-FRANK, A. SLY, AND O. TAMUZ (2020): “Social Learning Equilibria,” *Econometrica*, 88, 1235–1267.
- MOSSEL, E. AND O. TAMUZ (2017): “Opinion Exchange Dynamics,” *Probability Surveys*, 14, 155–204.
- MYATT, D. AND C. WALLACE (2017): “Information Acquisition and Use by Networked Players,” University of Warwick. Department of Economics. Creta Discussion Paper Series (32), available at <http://wrap.warwick.ac.uk/90449/>.
- PINELIS, I. (2018): “Inverse of Matrix with Blocks of Ones,” MathOverflow, URL:<https://mathoverflow.net/q/296933> (version: 2018-04-04).
- SETHI, R. AND M. YILDIZ (2012): “Public Disagreement,” *American Economic Journal: Microeconomics*, 4, 57–95.
- (2016): “Communication with Unknown Perspectives,” *Econometrica*, 84, 2029–2069.
- SGROI, D. (2002): “Optimizing Information in the Herd: Guinea Pigs, Profits, and Welfare,” *Games and Economic Behavior*, 39, 137–166.
- SHAHRAMPOUR, S., S. RAKHLIN, AND A. JADBABAIE (2013): “Online Learning of Dynamic Parameters in Social Networks,” *Advances in Neural Information Processing Systems*.
- VIVES, X. (1993): “How Fast do Rational Agents Learn?” *The Review of Economic Studies*, 60, 329–347.

- WEIZSÄCKER, G. (2010): “Do We Follow Others When We Should? A Simple Test of Rational Expectations,” *The American Economic Review*, 100, 2340–2360.
- WOLITZKY, A. (2018): “Learning from Others’ Outcomes,” *American Economic Review*, 108, 2763–2801.

APPENDIX A. EXISTENCE OF EQUILIBRIUM: PROOF OF PROPOSITION 1

Recall from Section 3.1 the map Φ , which gives the next-period covariance matrix $\Phi(\mathbf{V}_t)$ for any $\mathbf{V}_t$. The expression given there for this map ensures that its entries are continuous functions of the entries of $\mathbf{V}_t$. Our strategy is to show that this function maps a convex, compact set, $\mathcal{K}$, to itself, which, by Brouwer’s fixed-point theorem, ensures that Φ has a fixed point $\hat{\mathbf{V}}$. We will then argue that this fixed point corresponds to a stationary linear equilibrium.

We begin by defining the compact set $\mathcal{K}$. Because memory is arbitrary, entries of $\mathbf{V}_t$ are covariances between pairs of neighbor actions from any periods available in memory. Let k, l be two indices of such actions, corresponding to actions taken at nodes i and j respectively (at potentially different times), and let $\bar{\sigma}_i^2 = \max \left\{ \sigma_i^2, \rho^{m-1} \sigma_i^2 + \frac{1-\rho^{m-1}}{1-\rho} \right\}$. Now let $\mathcal{K} \subset \mathcal{V}$ be the subset of symmetric positive semi-definite matrices $\mathbf{V}_t$ such that, for any such k, l ,

$$V_{kk,t} \in \left[\min \left\{ \frac{1}{1 + \sigma_i^{-2}}, \frac{\rho^{m-1}}{1 + \sigma_i^{-2}} + \frac{1 - \rho^{m-1}}{1 - \rho} \right\}, \max \left\{ \sigma_i^2, \rho^{m-1} \sigma_i^2 + \frac{1 - \rho^{m-1}}{1 - \rho} \right\} \right]$$

$$V_{kl,t} \in [-\bar{\sigma}_i \bar{\sigma}_j, \bar{\sigma}_i \bar{\sigma}_j].$$

This set is closed and convex, and we claim that $\Phi(\mathcal{K}) \subset \mathcal{K}$.

To show this claim, we will first find upper and lower bounds on the variance of any neighbor’s action (at any period in memory). For the upper bound, note that a Bayesian agent will not choose an action with a larger variance than her signal, which has variance σ_i^2 . For a lower bound, note that if she knew the previous period’s state and her own signal, then the variance of her action would be $\frac{1}{1 + \sigma_i^{-2}}$. Thus an agent observing only noisy estimates of θ_t and her own signal can do no better.

By the same reasoning applied to the node- i agent from m periods ago, the error variance of $\rho^m a_{i,t-m} - \theta_t$ is at most $\rho^m \sigma_i^2 + \frac{1-\rho^m}{1-\rho}$ and at least $\frac{\rho^m}{1 + \sigma_i^{-2}} + \frac{1-\rho^m}{1-\rho}$. This establishes bounds on $V_{kk,t}$ for observations k from either the most recent or the oldest available period. The corresponding bounds from the periods between $t - m + 1$ and t are always weaker than at least one of the two bounds we have described, so we need only take minima and maxima over two terms.

This established the claimed bound on the variances. The bounds on covariances follow from Cauchy-Schwartz.

We have now established that there is a variance-covariance matrix $\hat{\mathbf{V}}$ such that $\Phi(\hat{\mathbf{V}}) = \hat{\mathbf{V}}$. By definition of Φ , this means there exists some weight profile $(\hat{\mathbf{W}}, \hat{\mathbf{w}}^s)$ such that, when

applied to prior actions that have variance-covariance matrix $\widehat{\mathbf{V}}$, produce variance-covariance matrix $\widehat{\mathbf{V}}$. However, it still remains to show that this is the variance-covariance matrix reached when agents have been using the weights $(\widehat{\mathbf{W}}, \widehat{\mathbf{w}}^s)$ forever.

To show this, first observe that if agents have been using the weights $(\widehat{\mathbf{W}}, \widehat{\mathbf{w}}^s)$ forever, the variance-covariance matrix $\mathbf{V}_t$ in any period is uniquely determined and does not depend on t ; call this $\check{\mathbf{V}}$.³⁹ This is because actions can be expressed as linear combinations of private signals with coefficients depending only on the weights. Second, it follows from our construction above of the matrix $\widehat{\mathbf{V}}$ and the weights $(\widehat{\mathbf{W}}, \widehat{\mathbf{w}}^s)$ that there is a distribution of actions where the variance-covariance matrix is $\widehat{\mathbf{V}}$ in every period and agents are using weights $(\widehat{\mathbf{W}}, \widehat{\mathbf{w}}^s)$ in every period. Combining the two statements shows that in fact $\check{\mathbf{V}} = \widehat{\mathbf{V}}$, and this completes the proof. Note that this argument also establishes that the response profile we have constructed is a strategy profile: under the responses used, we can write formally the dependence of actions on all prior signals, and verify using the observations on decay of dependence across time that the formula is summable and hence defines unique actions.

APPENDIX B. PROOF OF THEOREM 1

B.1. Notation and key notions. Let $\mathbb{S}$ be the (by assumption finite) set of all possible signal variances, and let $\bar{\sigma}^2$ be the largest of them. The proof will focus on the covariances of errors in social signals. Suppose that all agents have at least one neighbor. Take two arbitrary agents i and j . Recall that both $r_{i,t}$ and $r_{j,t}$ have mean θ_{t-1} , because each is an unbiased estimate⁴⁰ of θ_{t-1} ; we will thus focus on the errors $r_{i,t} - \theta_{t-1}$. Let A_t denote the variance-covariance matrix $(\text{Cov}(r_{i,t} - \theta_{t-1}, r_{j,t} - \theta_{t-1}))_{i,j}$ and let $\mathcal{W}$ be the set of such covariance matrices. For all i, j note that $\text{Cov}(r_{i,t} - \theta_{t-1}, r_{j,t} - \theta_{t-1}) \in [-\bar{\sigma}^2, \bar{\sigma}^2]$ using the Cauchy-Schwarz inequality and the fact that $\text{Var}(r_{i,t} - \theta_{t-1}) \in [0, \bar{\sigma}^2]$ for all i . This fact about variances says that no social signal is worse than putting all weight on an agent who follows only her private signal. Thus the best-response map Φ is well-defined and induces a map $\tilde{\Phi}$ on $\mathcal{W}$.

Next, for any $\psi, \zeta > 0$ we will define the subset $\mathcal{W}_{\psi, \zeta} \subset \mathcal{W}$ to be the set of covariance matrices in $\mathcal{W}$ such that both of the following hold:

1. for any pair of distinct agents⁴¹ $i \in G_n^k$ and $j \in G_n^{k'}$,

$$\text{Cov}(r_{i,t} - \theta_{t-1}, r_{j,t} - \theta_{t-1}) = \psi_{kk'} + \zeta_{ij}$$

³⁹The variance-covariance matrices are well-defined because the (W, w^s) weights yield unambiguous strategy profiles in the sense of Appendix E.

⁴⁰This is because it is a linear combination, with coefficients summing to 1, of unbiased estimates of θ_{t-1} .

⁴¹Throughout this proof, we abuse terminology by referring to agents and nodes interchangeably when the relevant t is clear or specified nearby.

- where (i) $\psi_{kk'}$ depends only on the network types of the two agents (k and k' , which may be the same); (ii) $|\psi_{kk'}| < \psi$; and (iii) $|\zeta_{ij}| < \zeta$;
2. for any single agent $i \in G_n^k$,

$$\text{Var}(r_{i,t} - \theta_{t-1}) = \psi_k + \zeta_{ii}$$

where (i) ψ_k only depends on the network type of the agent; (ii) $|\psi_k| < \psi$, and (iii) $|\zeta_{ii}| < \zeta$.

This is the space of covariance matrices such that each covariance is split into two parts. Considering (1) first, $\psi_{kk'}$ is an effect that depends only on i 's and j 's network types, while ζ_{ij} adjusts for the individual-level heterogeneity arising from different link realizations. The description of the decomposition in (2) is analogous.

B.2. Proof strategy.

B.2.1. *A set $\mathcal{W}_{\bar{\psi}, \bar{\zeta}}$ of outcomes with good learning.* Our goal is to show that as n grows large, there is an equilibrium in which $\text{Var}(r_{i,t} - \theta_{t-1})$ becomes very small, which then implies that the agents asymptotically learn. To this end we define a set of covariances with this property as well as some other useful properties. We will take $\bar{\psi}$ and $\bar{\zeta}$ to be arbitrarily small numbers and show that for large enough n , with high probability (which we abbreviate “asymptotically almost surely” or “a.a.s.”) there is an equilibrium with a social error covariance matrix A_t in the set $\mathcal{W}_{\bar{\psi}, \bar{\zeta}}$. That will imply that, in this equilibrium, $\text{Var}(r_{i,t} - \theta_{t-1})$ becomes arbitrarily small as we take the constants $\bar{\psi}$ and $\bar{\zeta}$ to be small. In our constructions, the ζ_{ij} (resp., ζ_i) terms will be set to much smaller values than the $\psi_{kk'}$ (resp., ψ_k) terms, because group-level covariances are more predictable and less sensitive to idiosyncratic realizations than individual-level covariances.

B.2.2. *Approach to showing that $\mathcal{W}_{\bar{\psi}, \bar{\zeta}}$ contains an equilibrium.* To show that there is (a.a.s.) an equilibrium outcome with a social error covariance matrix A_t in the set $\mathcal{W}_{\bar{\psi}, \bar{\zeta}}$, the plan is to construct a set so that (a.a.s.) $\bar{\mathcal{W}} \subset \mathcal{W}_{\bar{\psi}, \bar{\zeta}}$ and $\tilde{\Phi}(\bar{\mathcal{W}}) \subset \bar{\mathcal{W}}$. This set will contain an equilibrium by the Brouwer fixed point theorem, and therefore so will $\mathcal{W}_{\bar{\psi}, \bar{\zeta}}$.

To construct the set $\bar{\mathcal{W}}$, we will fix a positive constant β (to be determined later), and define

$$\bar{\mathcal{W}} = \mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}} \cup \tilde{\Phi} \left(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}} \right).$$

We will then prove that, for large enough n , (i) $\tilde{\Phi}(\bar{\mathcal{W}}) \subseteq \bar{\mathcal{W}}$ and (ii) for another suitable positive constant λ ,

$$\bar{\mathcal{W}} \subset \mathcal{W}_{\frac{\beta}{n}, \frac{\lambda}{n}}.$$

This will allow us to establish that (a.a.s.) $\bar{\mathcal{W}} \subset \mathcal{W}_{\bar{\psi}, \bar{\zeta}}$ and $\tilde{\Phi}(\bar{\mathcal{W}}) \subset \bar{\mathcal{W}}$, with $\bar{\psi}$ and $\bar{\zeta}$ being arbitrarily small numbers.

The following two lemmas will allow us to deduce (immediately after stating them) properties (i) and (ii) of $\overline{\mathcal{W}}$.

Lemma 1. There is a function $\underline{\lambda}(\beta) \geq 1$ such that the following holds. For all large enough β and all $\lambda \geq \underline{\lambda}(\beta)$, for n sufficiently large we have $\tilde{\Phi}\left(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}\right) \subset \mathcal{W}_{\frac{\beta}{n}, \frac{\lambda}{n}}$ with probability at least $1 - \frac{1}{n}$.

Lemma 2. For all large enough β , for n sufficiently large, $\tilde{\Phi}^2\left(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}\right) \subset \mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}$, with probability at least $1 - \frac{1}{n}$.⁴²

Putting these lemmas together, a.a.s. we have,

$$\tilde{\Phi}^2\left(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}\right) \subset \mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}} \quad \text{and} \quad \tilde{\Phi}\left(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}\right) \subset \mathcal{W}_{\frac{\beta}{n}, \frac{\lambda}{n}}.$$

From this it follows that $\overline{\mathcal{W}} = \mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}} \cup \tilde{\Phi}\left(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}\right)$ is mapped to a subset of itself by $\tilde{\Phi}$, and contained in $\mathcal{W}_{\frac{\beta}{n}, \frac{\lambda}{n}}$, as claimed.

B.2.3. Proving the lemmas by analyzing how $\tilde{\Phi}$ and $\tilde{\Phi}^2$ act on sets $\mathcal{W}_{\psi, \zeta}$. The lemmas are about how $\tilde{\Phi}$ and $\tilde{\Phi}^2$ act on the covariance matrix A_t , assuming it is in a certain set $\mathcal{W}_{\psi, \zeta}$, to yield new covariance matrices. Thus, we will prove these lemmas by studying two periods of updating. The analysis will come in five steps.

Step 1: No-large-deviations (NLD) networks and the high-probability event. Step 1 concerns the “with high probability” part of the lemmas. In the entire argument, we condition on the event of a *no-large-deviations* (NLD) network realization, which says that certain realized statistics in the network (e.g., number of paths between two nodes) are close to their expectations. The expectations in question depend only on agents’ types. Therefore, on the NLD realization, the realized statistics do not vary much based on which exact agents we focus on, but rather depend only on their types. Step 1 defines the NLD event E formally and shows that it has high probability. We use the structure of the NLD event throughout our subsequent steps, as we mention below.

Step 2: Weights in one step of updating are well-behaved. We are interested in $\tilde{\Phi}$ and $\tilde{\Phi}^2$, which describe how the covariance matrix A_t of social signal errors changes under updating. How this works is determined by the “basic” updating map Φ , and so we begin by studying the weights involved in it and then make deductions about the implications for the evolution of the variance-covariance matrix A_t .

The present step establishes that in one step of updating, the weight $W_{ij, t+1}$ that agent $(i, t+1)$ places on the action of another agent j in period t , does not depend too much on the identities of i and j . It only depends on their (network and signal) types. This is established by using our explicit formula for weights in terms of covariances. We rely on (i) the fact that

⁴²The notation $\tilde{\Phi}^2$ means the operator $\tilde{\Phi}$ applied twice.

covariances are assumed to start out in a suitable $\mathcal{W}_{\psi, \zeta}$, and (ii) our conditioning on the NLD event E . The NLD event is designed so that the network quantities that go into determining the weights depend only on the types of i and j (because the NLD event forbids too much variation within type). The restriction to $A_t \in \mathcal{W}_{\psi, \zeta}$ ensures that covariances in the initial period t do not vary too much with type, either.

Step 3: Lemma 1: $\tilde{\Phi}(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}) \subset \mathcal{W}_{\frac{\beta}{n}, \frac{\lambda}{n}}$. Once we have analyzed one step of updating, it is natural to consider the implications for the covariance matrix. Because we now have a bound on how much weights can vary after one step of updating, we can compute bounds on covariances. We show that if covariances A_t are in $\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}$, then after one step, covariances are in $\mathcal{W}_{\frac{\beta}{n}, \frac{\lambda}{n}}$. Note that the introduction of another parameter λ on the right-hand side implies that this step might worsen our control on covariances somewhat, but in a bounded way. This establishes Lemma 1.

Step 4: Weights in two steps of updating are well-behaved. The fourth step establishes that the statement made in Step 2 remains true when we replace $t + 1$ by $t + 2$. By the same sort of reasoning as in Step 2, an additional period of updating cannot create too much further idiosyncratic variation in weights. Proving this requires analyzing the covariance matrices of various social signals (i.e., the A_{t+1} that the updating induces), which is why we needed to do Step 3 first.

Step 5: Lemma 2: $\tilde{\Phi}^2(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}) \subset \mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}$. Now we use our understanding of weights from the previous steps, along with additional structure, to show the key remaining fact. What we have established so far about weights allows us to control the weight that a given agent's estimate at time $t + 2$ places on the social signal of another agent at time t . This is Step 5(a). In the second part, Step 5(b), we use that to control the covariances in A_{t+2} . It is important in this part of the proof that different agents have very similar “second-order neighborhoods”: the paths of length 2 beginning from an agent are very similar, in terms of their counts and what types of agents they go through. We use our control of second-order neighborhoods, as well as the assumptions on variation across entries of A_t to bound this variation well enough to conclude that $A_{t+2} \in \mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}$.

B.3. Carrying out the steps.

B.3.1. Step 1. Here we formally define the NLD event, which we call E . It is given by $E = \cap_{i=1}^5 E_i$, where the events E_i will be defined next.

(E_1) Let $X_{i, \tau k}^{(1)}$ be the number of agents having signal type τ and network type k who are observed by i . The event E_1 is that this quantity is close to its expected value in the following sense, simultaneously for all possible values of the subscript:

$$(1 - \zeta^2)\mathbb{E}[X_{i, \tau k}^{(1)}] \leq X_{i, \tau k}^{(1)} \leq (1 + \zeta^2)\mathbb{E}[X_{i, \tau k}^{(1)}].$$

(E_2) Let $X_{ii',\tau k}^{(2)}$ be the number of agents having signal type τ and network type k who are observed by *both* i and i' . The event E_2 is that this quantity is close to its expected value in the following sense, simultaneously for all possible values of the subscript:

$$(1 - \zeta^2)\mathbb{E}[X_{ii',\tau k}^{(2)}] \leq X_{ii',\tau k}^{(2)} \leq (1 + \zeta^2)\mathbb{E}[X_{ii',\tau k}^{(2)}].$$

(E_3) Let $X_{i,\tau k,j}^{(3)}$ be the number of agents having signal type τ and network type k who are observed by agent i and who observe agent j . The event E_3 is that this quantity is close to its expected value in the following sense, simultaneously for all possible values of the subscript:

$$(1 - \zeta^2)\mathbb{E}[X_{i,\tau k,j}^{(3)}] \leq X_{i,\tau k,j}^{(3)} \leq (1 + \zeta^2)\mathbb{E}[X_{i,\tau k,j}^{(3)}].$$

(E_4) Let $X_{ii',\tau k,j}^{(4)}$ be the number of agents having signal type τ and network type k who are observed by both agent i and i' and who observe j . The event E_4 is that this quantity is close to its expected value in the following sense, simultaneously for all possible values of the subscript:

$$(1 - \zeta^2)\mathbb{E}[X_{ii',\tau k,j}^{(4)}] \leq X_{ii',\tau k,j}^{(4)} \leq (1 + \zeta^2)\mathbb{E}[X_{ii',\tau k,j}^{(4)}].$$

(E_5) Let $X_{i,\tau k,jj'}^{(5)}$ be the number of agents of signal type τ and network type k who are observed by agent i and who observe both j and j' . The event E_5 is that this quantity is close to its expected value in the following sense, simultaneously for all possible values of the subscript:

$$(1 - \zeta^2)\mathbb{E}[X_{i,\tau k,jj'}^{(5)}] \leq X_{i,\tau k,jj'}^{(5)} \leq (1 + \zeta^2)\mathbb{E}[X_{i,\tau k,jj'}^{(5)}].$$

We claim that the probability of the complement of the event E vanishes exponentially. We can check this by showing that the probability of each of the E_i vanishes exponentially. For E_1 , for example, the bounds will hold unless at least one agent has degree outside the specified range. The probability of this is bounded above by the sum of the probabilities of each individual agent having degree outside the specified range. By Chebyshev's inequality, the probability a given agent has degree outside this range vanishes exponentially. Because there are n agents in G_n , this sum vanishes exponentially as well. The other cases are similar.

For the rest of the proof, we condition on the event E .

B.3.2. Step 2. As a shorthand, let $\psi = \beta/n$ for a sufficiently large constant β , and let $\zeta = 1/n$.

Lemma 3. Suppose that in period t the matrix $A = A_t$ of covariances of social signals satisfies $A \in \mathcal{W}_{\psi,\zeta}$ and all agents are optimizing in period $t+1$. Then there is a γ so that for all n sufficiently large,

$$\frac{W_{ij,t+1}}{W_{i'j',t+1}} \in \left[1 - \frac{\gamma}{n}, 1 + \frac{\gamma}{n}\right].$$

whenever i and i' have the same network and signal types and j and j' have the same network and signal types.

To prove this lemma, we will use our weights formula:

$$W_{i,t+1} = \frac{\mathbf{1}^T \mathbf{C}_{i,t}^{-1}}{\mathbf{1}^T \mathbf{C}_{i,t}^{-1} \mathbf{1}}.$$

This says that in period $t + 1$, agent i 's weight on agent j is proportional to the sum of the entries of column j of $\mathbf{C}_{i,t}^{-1}$. We want to show that the change in weights is small as the covariances of observed social signals vary slightly. To do so we will use the Taylor expansion of $f(A) = \mathbf{C}_{i,t}^{-1}$ around the covariance matrix $A(0)$ at which all $\psi_{kk'} = 0$, $\psi_k = 0$ and $\zeta_{ij} = 0$.

We begin with the first partial derivative of f at $A(0)$ in an arbitrary direction. Let $A(x)$ be any perturbation of A_0 in one parameter, i.e., $A(x) = A(0) + xM$ for some constant matrix M with entries in $[-1, 1]$. Let $\mathbf{C}_i(x)$ be the matrix of covariances of the actions observed by i given that the covariances of agents' social signals were $A(x)$. There exists a constant γ_1 depending only on the possible signal types such that each entry of $\mathbf{C}_i(x) - \mathbf{C}_i(x')$ has absolute value at most $\gamma_1(x - x')$ whenever both x and x' are small.

We will now show that the column sums of $\mathbf{C}_i(x)^{-1}$ are close to the column sums of $\mathbf{C}(0)_i^{-1}$. To do so, we will evaluate the formula

$$\frac{\partial f(A(x))}{\partial x} = \frac{\partial \mathbf{C}_i(x)^{-1}}{\partial x} = \mathbf{C}_i(x)^{-1} \frac{\partial \mathbf{C}_i(x)}{\partial x} \mathbf{C}_i(x)^{-1} \quad (\text{B.1})$$

at zero. If we can bound each column sum of this expression (evaluated at zero) by a constant (depending only on the signal types and the number of network types K), then the first derivative of f will also be bounded by a constant.

Recall that $\mathbb{S}$ is the set of signal types and let $S = |\mathbb{S}|$; index the signal types by numbers ranging from 1 to S . To bound the column sums of $\mathbf{C}_i(0)^{-1}$, suppose that the agent observes r_i agents from each signal type $1 \leq i \leq S$. Reordering so that all agents of each signal type are grouped together, we can write

$$\mathbf{C}_i(0) = \begin{pmatrix} a_{11}\mathbf{1}_{r_1 \times r_1} + b_1 I_{r_1} & a_{12}\mathbf{1}_{r_1 \times r_2} & & a_{1S}\mathbf{1}_{r_1 \times r_S} \\ a_{12}\mathbf{1}_{r_2 \times r_1} & a_{22}\mathbf{1}_{r_2 \times r_2} + b_2 I_{r_2} & & \vdots \\ & & \ddots & \\ a_{1S}\mathbf{1}_{r_S \times r_1} & \dots & & a_{SS}\mathbf{1}_{r_S \times r_S} + b_S I_{r_S} \end{pmatrix}$$

Therefore, $\mathbf{C}_i(0)$ can be written as a block matrix with blocks $a_{ij}\mathbf{1}_{r_i \times r_j} + b_i \delta_{ij} I_{r_i}$ where $1 \leq i, j \leq S$ and $\delta_{ij} = 1$ for $i = j$ and 0 otherwise.

We now have the following important approximation of the inverse of this matrix.⁴³

Lemma 4 (Pinelis (2018)). Let C be a matrix consisting of $S \times S$ blocks, with its (i, j) block given by

$$a_{ij}\mathbf{1}_{r_i \times r_j} + b_i \delta_{ij} I_{r_i}$$

⁴³We are very grateful to Iosif Pinelis for suggesting this argument.

and let $A = a_{ij}1_{r_i \times r_j}$ be an invertible matrix. As $n \rightarrow \infty$, then the (i, i) block of C^{-1} equals

$$\frac{1}{b_i}I_{r_i} - \frac{1}{b_i r_i}1_{r_i \times r_i} + O(1/n^2)$$

while the off-diagonal blocks are $O(1/n^2)$.

Proof. First note that the ij -block of C^{-1} has the form

$$c_{ij}1_{r_i \times r_j} + d_i \delta_{ij} I_{r_i}$$

for some real c_{ij} and d_i .

Therefore, CC^{-1} can be written in matrix form as

$$\begin{aligned} \sum_k (a_{ik}1_{r_i \times r_k} + b_i \delta_{ik} I_{r_i})(c_{kj}1_{r_k \times r_j} + d_k \delta_{kj} I_{r_k}) = \\ (a_{ij}d_j + \sum_k (a_{ik}r_k + \delta_{ik}b_k)c_{kj})1_{r_i \times r_j} + b_i d_i \delta_{ij} I_{r_i}. \end{aligned} \quad (\text{B.2})$$

Note that the last summand is the identity matrix.

Let D_d denote the diagonal matrix with d_i in the (i, i) diagonal entry, let $D_{1/b}$ denote the diagonal matrix with $1/b_i$ in the (i, i) diagonal entry, etc. Breaking up the previous display (B.2) into its diagonal and off-diagonal parts, we can write

$$AD_d + (AD_r + D_b)C = 0 \text{ and } D_d = D_{1/b}.$$

Hence,

$$\begin{aligned} C &= -(AD_r + D_b)^{-1}AD_d \\ &= -(I_q + D_r^{-1}A^{-1}D_b)^{-1}(AD_r)^{-1}AD_{1/b} \\ &= -(I_q + D_r^{-1}A^{-1}D_b)^{-1}D_{1/(br)} \\ &= -D_{1/(br)} + O(1/n^2) \end{aligned}$$

where $br := (b_1r_1, \dots, b_qr_q)$. Therefore as $n \rightarrow \infty$ the off-diagonal blocks will be $O(1/n^2)$ while the diagonal blocks are

$$\frac{1}{b_i}I_{r_i} - \frac{1}{b_i r_i}1_{r_i \times r_i} + O(1/n^2)$$

as desired. □

Using Lemma 4 we can analyze the column sums of⁴⁴

$$C_i(0)^{-1}MC_i(0)^{-1}.$$

In more detail, we use the formula of the lemma to estimate both copies of $C_i(0)^{-1}$, and then expand this to write an expression for any column sum of $C_i(0)^{-1}MC_i(0)^{-1}$. It follows

⁴⁴Recall we wrote $A(x) = A(0) + xM$, and in (B.1) we expressed the derivative of f in x in terms of the matrix we exhibit here.

straightforwardly from this calculation that all these column sums are $O(1/n)$ whenever all entries of M are in $[-1, 1]$.

We can bound the higher-order terms in the Taylor expansion by the same technique: by differentiating equation B.1 repeatedly in x , we obtain an expression for the k^{th} derivative in terms of $\mathbf{C}_i(0)^{-1}$ and M :

$$f^{(k)}(0) = k! \mathbf{C}_i(0)^{-1} M \mathbf{C}_i(0)^{-1} M \mathbf{C}_i(0)^{-1} \cdot \dots \cdot M \mathbf{C}_i(0)^{-1},$$

where M appears k times in the product. By the same argument as above, we can show that the column sums of $\frac{f^{(k)}(0)}{k!}$ are bounded by a constant independent of n . The Taylor expansion is

$$f(A) = \sum_k \frac{f^{(k)}(0)}{k!} x^k.$$

Since we take $A \in \mathcal{W}_{\psi, \zeta}$, we can assume that x is $O(1/n)$. Because the column sums of each summand are bounded by a constant times x^k , the column sums of $f(A)$ are bounded by a constant.

Finally, because the variation in the column sums is $O(1/n)$ and the weights are proportional to the column sums, each weight varies by at most a multiplicative factor of γ_1/n for some γ_1 . We find that the first part of the lemma, which bounded the ratios between weights $W_{ij,t+1}/W_{i'j',t+1}$, holds.

B.3.3. Step 3. We complete the proof of Lemma 1, which states that the covariance matrix of $r_{i,t+1}$ is in $\mathcal{W}_{\psi, \zeta'}$. Recall that $\zeta' = \lambda/n$ for some constant n , so we are showing that if the covariance matrix of the $r_{i,t}$ is in a neighborhood $\mathcal{W}_{\psi, \zeta}$, then the covariance matrix in the next period is in a somewhat larger neighborhood $\mathcal{W}_{\psi, \zeta'}$. The remainder of the argument then follows by the same arguments as in the proof of the first part of the lemma: we now bound the change in time- $(t+2)$ weights as we vary the covariances of time- $(t+1)$ social signals within this neighborhood.

Recall that we decomposed each covariance $\text{Cov}(r_{i,t} - \theta_{t-1}, r_{j,t} - \theta_{t-1}) = \psi_{kk'} + \zeta_{ij}$ into a term $\psi_{kk'}$ depending only on the types of the two agents and a term ζ_{ij} , and similarly for variances. To show the covariance matrix is contained in $\mathcal{W}_{\psi, \zeta'}$, we bound each of these terms suitably.

We begin with ζ_{ij} (and ζ_i). We can write

$$r_{i,t+1} = \sum_j \frac{W_{ij,t+1}}{1 - w_{i,t+1}^s} \rho a_{i,t} = \sum_j \frac{W_{ij,t+1}}{1 - w_{i,t+1}^s} \rho (w_{j,t}^s s_{j,t} + (1 - w_{j,t}^s) r_{j,t}).$$

By the first part of the lemma, the ratio between any two weights (both of the form $W_{ij,t+1}$, $w_{i,t+1}^s$, or $w_{j,t}^s$) corresponding to pairs of agents of the same types is in $[1 - \gamma_1/n, 1 + \gamma_1/n]$ for a constant γ_1 . We can use this to bound the variation in covariances of $r_{i,t+1}$ within types

by ζ' : we take the covariance of $r_{i,t+1}$ and $r_{j,t+1}$ using the expansion above and then bound the resulting summation by bounding all coefficients.

Next we bound $\psi_{kk'}$ (and ψ_k). It is sufficient to show that $\text{Var}(r_{i,t+1} - \theta_t)$ is at most ψ . To do so, we will give an estimator of θ_t with variance less than β/n , and this will imply $\text{Var}(r_{i,t+1} - \theta_t) < \beta/n = \psi$ (recall $r_{i,t+1}$ is the estimate of θ_t given agent i 's social observations in period $t+1$). Since this bounds all the variance terms by ψ , the covariance terms will also be bounded by ψ in absolute value.

Fix an agent i of network type k and consider some network type k' such that $p_{kk'} > 0$. Then there exists two signal types, which we call A and B , such that i observes $\Omega(n)$ agents of each of these signal types in G_n^k .⁴⁵ The basic idea will be that we can approximate θ_t well by taking a linear combination of the average of observed agents of network type k and signal type A and the average of observed agents of network type k and signal type B .

In more detail: Let $N_{i,A}$ be the set of agents of type A in network type k observed by i and $N_{i,B}$ be the set of agents of type B in network type k observed by i . Then fixing some agent j_0 of network type k ,

$$\frac{1}{|N_{i,A}|} \sum_{j \in N_{i,A}} a_{j,t-1} = \frac{\sigma_A^{-2}}{1 + \sigma_A^{-2}} \theta_t + \frac{1}{1 + \sigma_A^{-2}} r_{j_0,t-1} + \text{noise}$$

where the noise term has variance of order $1/n$ and depends on signal noise, variation in $r_{j,t}$, and variation in weights. These bounds on the noise term follow from the assumption that the covariance matrix of the $r_{i,t}$ is in a neighborhood $\mathcal{W}_{\psi,\zeta}$ and our analysis of variation in weights. Similarly

$$\frac{1}{|N_{i,B}|} \sum_{j \in N_{i,B}} a_{j,t-1} = \frac{\sigma_B^{-2}}{1 + \sigma_B^{-2}} \theta_t + \frac{1}{1 + \sigma_B^{-2}} r_{j_0,t-1} + \text{noise}$$

where the noise term has the same properties. Because $\sigma_A^2 \neq \sigma_B^2$, we can write θ_t as a linear combination of these two averages with coefficients independent of n up to a noise term of order $1/n$. We can choose β large enough such that this noise term has variance most β/n for all n sufficiently large. This completes the Proof of Lemma 1.

B.3.4. Step 4: We now give the two-step version of Lemma 3.

Lemma 5. Suppose that in period t the matrix $A = A_t$ of covariances of social signals satisfies $A \in \mathcal{W}_{\psi,\zeta}$ and all agents are optimizing in periods $t+1$ and $t+2$. Then there is a γ so that for all n sufficiently large,

$$\frac{W_{ij,t+2}}{W_{i'j',t+2}} \in \left[1 - \frac{\gamma}{n}, 1 + \frac{\gamma}{n}\right].$$

⁴⁵We use the notation $\Omega(n)$ to mean greater than Cn for some constant $C > 0$ when n is large.

whenever i and i' have the same network and signal types and j and j' have the same network and signal types.

Given what we established about covariances in Step 3, the lemma follows by the same argument as the proof of Lemma 3.

Step 5: Now that Lemma 5 is proved, we can apply it to show that $\tilde{\Phi}^2(\mathcal{W}_{\psi,\zeta}) \subset \mathcal{W}_{\psi,\zeta}$.

We will do this by first writing the time- $(t+2)$ behavior in terms of agents' time- t observations (Step 5(a)), which comes from applying $\tilde{\Phi}$ twice. This gives a formula that can be used for bounding the covariances⁴⁶ of time- $(t+2)$ actions in terms of covariances of time- t actions. Step 5(b) then applies this formula to show we can take ζ_{ij} and ζ_i to be sufficiently small. (Recall the notation introduced in Section B.1 above.) We split our expression for $r_{i,t+2}$ into several groups of terms and show that the contribution of each group of terms depends only on agents' types up to a small noise term. Step 5(c) notes that we can also take $\psi_{kk'}$ and ψ_k to be sufficiently small.

Step 5(a): We calculate:

$$\begin{aligned} r_{i,t+2} &= \sum_j \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} \rho a_{j,t+1} \\ &= \rho \left(\sum_j \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} w_{j,t+1}^s s_{j,t+1} + \sum_{j,j'} \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} W_{jj',t+1} \rho a_{j',t} \right) \\ &= \rho \left(\sum_j \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} w_{j,t+1}^s s_{j,t+1} + \rho \left(\sum_{j,j'} \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} W_{jj',t+1} w_{j',t}^s s_{j',t} \right. \right. \\ &\quad \left. \left. + \sum_{j,j'} \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} W_{jj',t+1} (1 - w_{j',t}^s) r_{j',t} \right) \right). \end{aligned}$$

Let $c_{ij',t}$ be the coefficient on $r_{j',t}$ in this expansion of $r_{i,t+2}$. Explicitly,

$$c_{ij',t} = \sum_j \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} W_{jj',t+1} (1 - w_{j',t}^s). \quad (\text{B.3})$$

The coefficient $c_{ij',t}$ adds up the influence of $r_{j',t}$ on $r_{i,t+2}$ over all paths of length two.

First, we establish a lemma about how much these weights vary.

Lemma 6. There exists γ such that for n sufficiently large, when i and i' have the same network types and j' and j'' have the same network and signal types, the ratio $c_{ij',t}/c_{i'j'',t}$ is in $[1 - \gamma/n, 1 + \gamma/n]$.

Proof. Fix i and j' . For each network type k'' and signal type s , consider the number of agents j of network type k'' and signal type s who are observed by i and who observe j' . This

⁴⁶We take this term to refer to variances, as well.

number varies by at most a factor ζ^2 as we change i and j' , preserving signal and network types. For each such j , the contribution of that agent's action to $c_{ij',t}$ is (recalling (B.3))

$$\frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} W_{jj',t+1} (1 - w_{j',t}^s).$$

By applying Lemma 3 repeatedly, we can choose γ_1 such that each of these contributions varies by at most a factor of γ_1/n as we change i in G_k and j' in $G_{k'}$. Thus, $c_{ij',t}$ is a sum of terms which vary by at most a multiplicative factor of γ_1/n as we change i and j' preserving signal and network types. If we can show that the sum of the absolute values of these terms is bounded, then it will follow that $c_{ij',t}$ varies by at most a multiplicative factor of γ/n for some n . This bound on the sum of absolute values follows from the calculation of weights in the proof of Lemma 3. $\square$

Step 5(b): We first show that fixing the values of $\psi_{kk'}$ and ψ_k in period t , the variation in the covariances $\text{Cov}(r_{i,t+2} - \theta_{t+1}, r_{i',t+2} - \theta_{t+1})$ of these terms as we vary i and i' over network types is not larger than ζ . From the formula above, we observe that we can decompose $r_{i,t+2} - \theta_{t+1}$ as a linear combination of three mutually independent groups of terms:

- (i) signal error terms $\eta_{j,t+1}$ and $\eta_{j',t}$;
- (ii) the errors $r_{j',t} - \theta_t$ in the social signals from period t ; and
- (iii) changes in state ν_t and ν_{t+1} between periods t and $t+2$.

Note that the terms $r_{j',t} - \theta_t$ are linear combinations of older signal errors and changes in the state. We bound each of the three groups in turn:

(i) Signal Errors: We first consider the contribution of signal errors. When i and i' are distinct, the number of such terms is close to its expected value because we are conditioning on the events E_2 and E_4 defined in Section B.1. Moreover the weights are close to their expected values by Step 2, so the variation is bounded suitably. When i and i' are equal, we use the facts that the weights are close to their expected values and the variance of an average of $\Omega(n)$ signals is small.

(ii) Social Signals: We now consider terms $r_{j',t} - \theta_t$, which correspond to the third summand in our expression for $r_{i,t+2}$. Since we will analyze the weight on ν_t below, it is sufficient to study the terms $r_{j',t} - \theta_{t-1}$.

By Lemma 6, the coefficients placed on $r_{j',t}$ by i and on $r_{j'',t}$ by i' vary by a factor of at most $2\gamma/n$. Moreover, the absolute value of each of these covariances is bounded above by ψ and the variation in these terms is bounded above by ζ . We conclude that the variation from these terms has order $1/n^2$.

(iii) Innovations: Finally, we consider the contribution of the innovations ν_t and ν_{t+1} . We treat ν_{t+1} first. We must show that any two agents of the same types place the same weight on the innovation ν_{t+1} (up to an error of order $\frac{1}{n^2}$). This will imply that the contributions of

timing to the covariances $\text{Cov}(r_{i,t+2} - \theta_{t+1}, r_{i',t+2} - \theta_{t+1})$ can be expressed as a term that can be included in the relevant $\psi_{kk'}$ and a lower-order term which can be included in $\zeta_{ii'}$.

The weight an agent places on ν_{t+1} is equal to the weight she places on signals from period $t + 1$. So this is equivalent to showing that the total weight

$$\rho \sum_j \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} w_{j,t+1}^s$$

agent i places on period $t + 1$ depends only on the network type k of agent i and $O(1/n^2)$ terms. We will first show the average weight placed on time- $(t + 1)$ signals by agents of each signal type depends only on k . We will then show that the total weights on agents of each signal type do not depend on n .

Suppose for simplicity here that there are two signal types A and B ; the general case is the same. We can split the sum from the previous paragraph into the subgroups of agents with signal types A and B :

$$\rho \sum_{j:\sigma_j^2=\sigma_A^2} \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} w_{j,t+1}^s + \rho \sum_{j:\sigma_j^2=\sigma_B^2} \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s} w_{j,t+1}^s.$$

Letting $W_i^A = \sum_{j:\sigma_j^2=\sigma_A^2} \frac{W_{ij,t+2}}{1 - w_{i,t+2}^s}$ be the total weight placed on agents with signal type A and similarly for signal type B , we can rewrite this as:

$$W_i^A \rho \sum_{j:\sigma_j^2=\sigma_A^2} \frac{W_{ij,t+2}}{W_i^A(1 - w_{i,t+2}^s)} w_{j,t+1}^s + W_i^B \rho \sum_{j:\sigma_j^2=\sigma_B^2} \frac{W_{ij,t+2}}{W_i^B(1 - w_{i,t+2}^s)} w_{j,t+1}^s.$$

The coefficients $\frac{W_{ij,t+2}}{W_i^A(1 - w_{i,t+2}^s)}$ in the first sum now sum to one, and similarly for the second. We want to check that the first sum $\sum_{j:\sigma_j^2=\sigma_A^2} \frac{W_{ij,t+2}}{W_i^A(1 - w_{i,t+2}^s)} w_{j,t+1}^s$ does not depend on k , and the second sum is similar.

For each j in group A ,

$$w_{j,t+1}^s = \frac{\sigma_A^{-2}}{\sigma_A^{-2} + (\rho^2 \kappa_{j,t+1} + 1)^{-1}},$$

where we define $\kappa_{j,t+1}^2 = \text{Var}(r_{j,t+1} - \theta_t)$ to be the error variance of the social signal. Because $\kappa_{j,t+1}$ is close to zero, we can approximate $w_{j,t+1}^s$ locally as a linear function $\mu_1 \kappa_{j,t+1} + \mu_2$ where $\mu_1 < 1$ (up to order $\frac{1}{n^2}$ terms).

So we can write the sum of interest as

$$\sum_{j:\sigma_j^2=\sigma_A^2} \frac{W_{ij,t+2}}{W_i^A(1 - w_{i,t+2}^s)} \left(\mu_1 \sum_{j',j''} W_{jj',t+1} W_{jj'',t+1} (\rho^2 \mathbf{V}_{j'j'',t} + 1) + \mu_2 \right).$$

By Lemma 3, the weights vary by at most a multiplicative factor contained in $[1 - \gamma/n, 1 + \gamma/n]$. The number of paths from i to j' passing through agents of any network type k'' and any signal type is close to its expected value (which depends only on i 's network type), and the

weight on each path depends only on the types involved up to a factor in $[1 - \gamma/n, 1 + \gamma/n]$. The variation in $\mathbf{V}_{j',t}$ consists of terms of the form $\psi_{k'k''}$, $\psi_{k'}$, and $\zeta_{j'j''}$, all of which are $O(1/n)$, and terms from signal errors $\eta_{j',t}$. The signal errors only contribute when $j = j'$, and so only contribute to a fraction of the summands of order $1/n$. So we can conclude the total variation in this sum as we change i within the network type k has order $1/n^2$.

Now that we know each the average weight on private signals of the observed agents of each signal type depends only on k , it remains to check that W_i^A and W_i^B only depend on k . The coefficients W_i^A and W_i^B are the optimal weights on the group averages

$$\sum_{j:\sigma_j^2=\sigma_A^2} \frac{W_{ij,t+2}}{W_i^A(1-w_{i,t+2}^s)} \rho a_{j,t+1} \text{ and } \sum_{j:\sigma_j^2=\sigma_B^2} \frac{W_{ij,t+2}}{W_i^B(1-w_{i,t+2}^s)} \rho a_{j,t+1},$$

so we need to show that the variances and covariance of these two terms depend only on k . We check the variance of the first sum: we can expand

$$\sum_{\sigma_j^2=\sigma_A^2} \frac{W_{ij,t+2}}{W_i^A(1-w_{i,t+2}^s)} \rho a_{j,t+1} = \sum_{\sigma_j^2=\sigma_A^2} \frac{W_{ij,t+2}}{W_i^A(1-w_{i,t+2}^s)} \rho (w_{j,t+1}^s s_{j,t+1} + (1-w_{j,t+1}^s) r_{j,t+1}).$$

We can again bound the signal errors and social signals as in the previous parts of this proof, and show that the variance of this term depends only on k and $O(1/n^2)$ terms. The second variance and covariance are similar, so W_i^A and W_i^B depend only on k and $O(1/n^2)$ terms.

This takes care of the innovation ν_{t+1} . Because we have included any innovations prior to ν_t in the social signals $r_{j',t}$, to complete Step 5(b) we need only show the weight on ν_t depends only on the network type k of an agent.

The analysis is a simpler version of the analysis of the weight on ν_{t+1} . It is sufficient to show the total weight placed on period t social signals depends only on the network type of k of an agent i . This weight is equal to

$$\rho^2 \sum_{j,j'} \frac{W_{ij,t+2}}{1-w_{i,t+2}^s} \cdot W_{jj',t+1} \cdot (1-w_{j',t}^s).$$

As in the ν_{t+1} case, we can approximate $(1-w_{j',t}^s)$ as a linear function of $\kappa_{j',t}$ up to $O(1/n^2)$ terms. Because the number of paths to each agent j' through a given type and the weights on each such path cannot vary too much within types, the same argument shows that this sum depends only on k and $O(1/n^2)$ terms. Thus Step 5(b) is complete.

Step 5(c): The final step is to verify that we can take $\psi_{kk'}$ and ψ_k to be smaller than ψ . It is sufficient to show that the variance $\text{Var}(r_{i,t+2} - \theta_{t+1})$ of each social signal about θ_{t+1} is at most ψ . The proof is the same as in Step 2(b).

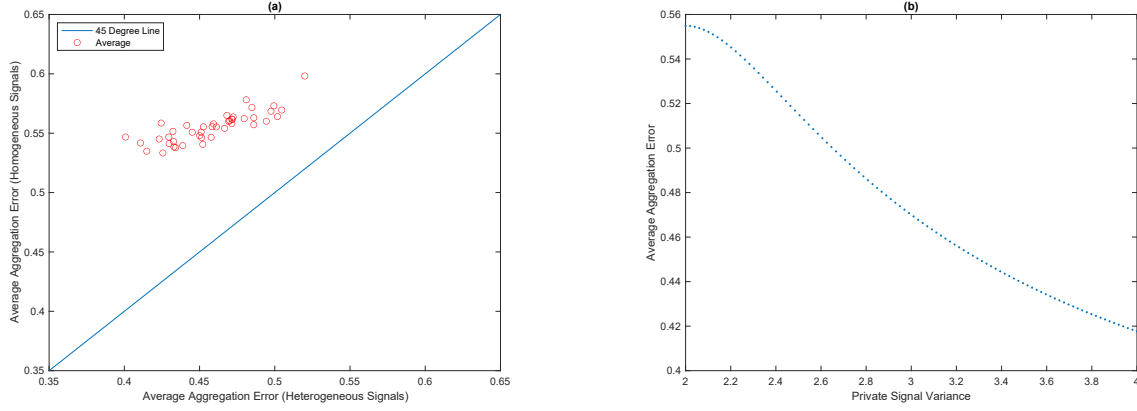


FIGURE C.1. Social signal variance in Indian villages. (a) The average social signal variance of agents in each village, in the homogeneous and heterogeneous cases. In the homogeneous case all agents have private signal variance 2. In the heterogeneous case, half of agents have private signal variance $\frac{3}{2}$ and half of agents have private signal variance 3. (b) The average social signal variance for all agents as we vary the worse private signal variance from 2 to 4 and hold fixed the average precision of private signals.

APPENDIX C. NUMERICAL RESULTS IN REAL NETWORKS (ONLINE APPENDIX)

The message of Section 4 is that signal diversity enables good aggregation, and signal homogeneity obstructs it. The theoretical results in that section, however, were asymptotic, and the good-aggregation result used some assumptions on the distribution of graphs. In this section we show that the substantive message applies to realistic networks with moderate degrees. We do this by computing equilibria for actual social networks from the data in Banerjee, Chandrasekhar, Duflo, and Jackson (2013). This data set contains the social networks of villages in rural India.⁴⁷ There are 43 networks in the data, with an average network size of 212 nodes (standard deviation = 53.5), and an average degree of 19 (standard deviation = 7.5).

Our simulation exercises measure the benefits of heterogeneity for equilibrium aggregation. For each network, we calculate the equilibrium with $\rho = 0.9$ for two types of environments. The first is the homogeneous case, with all signal variances set to 2. The second is a heterogeneous case, where half of the agents have a signal variance greater than 2 and half of villagers have a signal variance less than 2, chosen to hold constant the total amount of information that reaches the community via private signals. That is, we set the signal variances so that the average precision in each village is $\frac{1}{2}$, as in the homogeneous case. This signal assignment holds fixed the average utility when all villagers are autarkic, or equivalently holds fixed the average utility when all villagers know the state θ_{t-1} in the previous period exactly. At the

⁴⁷We take the networks that were used in the estimation in Banerjee, Chandrasekhar, Duflo, and Jackson (2013). As in their work, we take every reported relationship to be reciprocal for the purposes of sharing information. This makes the graphs undirected.

same time, it varies the level of heterogeneity in signal endowments. Villagers are randomly assigned to better or worse private signals, and the simulation results do not depend substantially on the realized random assignment. Our outcomes will be the average social signal error variance in each village and the average social signal error variance across all villages.

It is useful to begin by looking at the equilibrium average aggregation errors, i.e., social signal variances, in the case of homogeneous signals. This is the horizontal coordinate in Figure C.1(a); each village is a data point, and the points have a standard deviation of 0.013. In this case, differences in learning outcomes are due only to differences in the network structure, and we will call this number the *network-driven variation*. Now we introduce some private signal diversity. In our first exercise, we change the variance of the worse private signal from 2 (homogeneous signals) to 3 (heterogeneous signals), and adjust the other variance as discussed above to hold fixed the total amount of information coming into the network. The vertical coordinate in Figure C.1(b) depicts the equilibrium aggregation error in each village. The average of this number across all villages falls to 0.470, compared to 0.555 (in the homogeneous case). Therefore, adding heterogeneity by increasing the private signal variance for half of the agents by 50% changes social signal error variance by 6.5 times the network-driven variation. Learning is much better with some private signal heterogeneity than in villages with very favorable networks (i.e., those that achieve the best aggregation under homogeneous signals).

In Figure C.1(b), rather than working with the particular choice of 3 for the variance of the private signal, we look across all choices of this variance between 2 and 4 and plot the average equilibrium social signal variance across all villages.

Figure C.1(b) also sheds light on the value of a small amount of heterogeneity. The results in Section 4 can be summarized as saying that, to achieve the aggregation benchmark of essentially knowing the previous period's state, there need to be at least two different private signal variances in the network. Formally, this is a knife-edge result: As long as private signal variances differ at all, then as $n \rightarrow \infty$, aggregation errors vanish; with exactly homogeneous signal endowments, aggregation errors are much higher. The figure shows that the transition from the first regime to the second is actually gradual. In particular, a very small amount of heterogeneity provides little benefit in finite networks, as there is not enough diversity of signal endowments for villagers to anti-imitate. However, a 50% change in the variance of one of the signals (equivalently, a 22% change in its standard deviation) makes the community much better able to use the same total amount of information.

APPENDIX D. IDENTIFICATION AND TESTABLE IMPLICATIONS (ONLINE APPENDIX)

One of the main advantages of the parametrization we have studied is that standard methods can easily be applied to estimate the model and test hypotheses within it. The key feature making the model econometrically well-behaved is that, in the solutions we focus on, agents'

actions are linear functions of the random variables they observe. Moreover, the evolution of the state and arrival of information creates exogenous variation. We briefly sketch how these features can be used for estimation and testing.

Assume the following. The analyst obtains noisy measurements $\bar{a}_{i,t} = a_{i,t} + \xi_{i,t}$ of agent's actions (where $\xi_{i,t}$ are i.i.d., mean-zero error terms). He knows the parameter ρ governing the stochastic process, but may not know the network structure or the qualities of private signals $(\sigma_i)_{i=1}^n$. Suppose also that the analyst observes the state θ_t ex post (perhaps with a long delay).⁴⁸

Now, consider *any* steady state in which agents put constant weights W_{ij} on their neighbors and w_i^s on their private signals over time. We will discuss the case of $m = 1$ to save on notation, though all the statements here generalize readily to arbitrary m .

We first consider how to estimate the weights agents are using, and to back out the structural parameters of our model when it applies. The strategy does not rely on uniqueness of equilibrium. We can identify the weights agents are using through standard vector autoregression methods. In steady state,

$$\bar{a}_{i,t} = \sum_j W_{ij} \rho \bar{a}_{j,t-1} + w_i^s \theta_t + \zeta_{i,t}, \quad (\text{D.1})$$

where $\zeta_{i,t} = w_i^s \eta_{i,t} - \sum_j W_{ij} \rho \xi_{j,t-1} + \xi_{i,t}$ are error terms i.i.d. across time. The first term of this expression for $\zeta_{i,t}$ is the error of the signal that agent i receives at time t . The summation combines the measurement errors from the observations $\bar{a}_{j,t-1}$ from the previous period.⁴⁹ Thus, we can obtain consistent estimators $\widetilde{W}_{ij}$ and $\widetilde{w}_i^s$ for W_{ij} and w_i^s , respectively.

We now turn to the case in which agents are using *equilibrium* weights. First, and most simply, our estimates of agents' equilibrium weights allow us to recover the network structure. If the weight $\widehat{W}_{ij}$ is non-zero for any i and j , then agent i observes agent j . Generically the converse is true: if i observes j then the weight $\widehat{W}_{ij}$ is non-zero. Thus, network links can generically be identified by testing whether the recovered social weights are nonzero. For such tests (and more generally) the standard errors in the estimators can be obtained by standard techniques.⁵⁰

Now we examine the more interesting question of how structural parameters can be identified assuming an equilibrium is played, and also how to test the assumption of equilibrium.

The first step is to compute the empirical covariances of action errors from observed data; we call these $\widetilde{V}_{ij}$. Under the assumption of equilibrium, we now show how to determine the

⁴⁸We can instead assume that the analyst observes (a proxy for) the private signal $s_{i,t}$ of agent i ; we mention how below.

⁴⁹This system defines a VAR(1) process (or generally VAR(m) for memory length m).

⁵⁰Methods involving regularization may be practically useful in identifying links in the network. [Manresa \(2013\)](#) proposes a regularization (LASSO) technique for identifying such links (peer effects). In a dynamic setting such as ours, with serial correlation, the techniques required will generally be more complicated.

signal variances using the fact that equilibrium is characterized by $\Phi(\widehat{\mathbf{V}}) = \widehat{\mathbf{V}}$ and recalling the explicit formula (3.3) for Φ . In view of this formula, the signal variances σ_i^2 are uniquely determined by the other variables:

$$\widehat{V}_{ii} = \sum_j \sum_k \widehat{W}_{ij} \widehat{W}_{ik} \left(\rho^2 \widehat{V}_{jk} + 1 \right) + (\widehat{w}_i^s)^2 \sigma_i^2. \quad (\text{D.2})$$

Replacing the model parameters other than σ_i^2 by their empirical analogues, we obtain a consistent estimate $\widetilde{\sigma}_i^2$ of σ_i . This estimate could be directly useful—for example, to an analyst who wants to choose an “expert” from the network and ask about her private signals.

Note that our basic VAR for recovering the weights relies only on constant linear strategies and does not assume that agents are playing any particular strategy within this class. Thus, if agents are using some other behavioral rule (e.g., optimizing in a misspecified model) we can replace (D.2) by a suitable analogue that reflects the bounded rationality in agents’ inference. If such a steady state exists, and using the results in this section, one can create an econometric test that is suitable for testing how agents are behaving. For instance, we can test the hypothesis that they are Bayesian against the naive alternative of our Section 5.1.

APPENDIX E. DETAILS OF DEFINITIONS (ONLINE APPENDIX)

E.1. Exogenous random variables. Fix a probability space $(\Omega, \mathcal{F}, \mathbb{P})$. Let $(\nu_t, \eta_{i,t})_{t \in \mathbb{Z}, i \in N}$ be normal, mutually independent random variables, with ν_t having variance 1 and $\eta_{i,t}$ having variance σ_i^2 . Also take a stochastic process $(\theta_t)_{t \in \mathbb{Z}}$, such that for each $t \in \mathbb{Z}$, we have (for $0 < |\rho| \leq 1$)

$$\theta_t = \rho \theta_{t-1} + \nu_t.$$

Such a stochastic process exists by standard constructions of the AR(1) process or, in the case of $\rho = 1$, of the Gaussian random walk on a doubly infinite time domain. Define $s_{i,t} = \theta_t + \eta_{i,t}$.

E.2. Formal definition of game and stationary linear equilibria.

Players and strategies. The set of players (or agents) is $\mathcal{A} = \{(i, t) : i \in N, t \in \mathbb{Z}\}$. The set of (pure) *responses* of an agent (i, t) is defined to be the set of all Borel-measurable functions $\xi_{(i,t)} : \mathbb{R} \times (\mathbb{R}^{|N(i)|})^m \rightarrow \mathbb{R}$, mapping her own signal and her neighborhood’s actions, $(s_{i,t}, (\mathbf{a}_{N_i, t-\ell})_{\ell=1}^m)$, to a real-valued action $a_{i,t}$. We call the set of these functions $\widetilde{\Xi}_{(i,t)}$. Let $\widetilde{\Xi} = \prod_{(i,t) \in \mathcal{A}} \widetilde{\Xi}_{(i,t)}$ be the set of response profiles. We now define the set of (*unambiguous*) *strategy profiles*, $\Xi \subset \widetilde{\Xi}$. We say that a response profile $\xi \in \widetilde{\Xi}$ is a strategy profile if the following two conditions hold

1. There is a tuple of real-valued random variables $(a_{i,t})_{i \in N, t \in \mathbb{Z}}$ on $(\Omega, \mathcal{F}, \mathbb{P})$ such that for each $(i, t) \in \mathcal{A}$, we have

$$a_{i,t} = \xi_{(i,t)}(s_{i,t}, (\mathbf{a}_{N_i, t-\ell})_{\ell=1}^m).$$

2. Any two tuples of real-valued random variables $(a_{i,t})_{i \in N, t \in \mathbb{Z}}$ satisfying Condition 1 are equal almost surely.

That is, a response profile is a strategy profile if there is an essentially unique specification of behavior that is consistent with the responses: i.e., if the responses uniquely determine the behavior of the population, and hence payoffs.⁵¹ Note that if $\xi \in \Xi$, then it can be checked that $\tilde{\xi} = (\xi'_{(i,t)}, \xi_{-(i,t)}) \in \Xi$ whenever $\xi'_{(i,t)} \in \tilde{\Xi}_{(i,t)}$. Thus, if we start with a strategy profile and consider agent (i, t) 's deviations, they are unrestricted: she may consider any response.

Payoffs. The payoff of an agent (i, t) under any strategy profile $\xi \in \Xi$ is

$$u_{i,t}(\xi) = -\mathbb{E} [(a_{i,t} - \theta_t)^2] \in [-\infty, 0],$$

where the actions $a_{i,t}$ are taken according to $\xi_{(i,t)}$ and the expectation is taken in the probability space we have described. This expectation is well-defined because inside the expectation there is a non-negative, measurable random variable, for which an expectation is always defined, though it may be infinite.

Equilibria. A (Nash) *equilibrium* is defined to be a strategy profile $\xi \in \Xi$ such that, for each $(i, t) \in \mathcal{A}$ and each $\tilde{\xi} \in \Xi$ such that $\tilde{\xi} = (\xi'_{(i,t)}, \xi_{-(i,t)})$ for some $\xi'_{(i,t)} \in \tilde{\Xi}_{(i,t)}$, we have

$$u_{i,t}(\tilde{\xi}) \leq u_{i,t}(\xi).$$

For $p \in \mathbb{Z}$, we define the shift operator $\mathfrak{T}_p$ to translate variables to time indices shifted p steps forward. This definition may be applied, for example, to Ξ .⁵² A strategy profile $\xi \in \Xi$ is *stationary* if, for all $p \in \mathbb{Z}$, we have $\mathfrak{T}_p \xi = \xi$.

We say $\xi \in \Xi$ is a *linear* strategy profile if each ξ_i is a linear function. Our analysis focuses on *stationary, linear equilibria*.

APPENDIX F. REMAINING PROOFS (ONLINE APPENDIX)

F.1. Proof of Proposition 2. We first check there is a unique equilibrium and then prove the remainder of Proposition 2.

Lemma 7. Suppose G has symmetric neighbors. Then there is a unique equilibrium.

Proof of Lemma 7. We will show that when the network satisfies the condition in the proposition statement, Φ induces a contraction on a suitable space. For each agent, we can consider the variance of the best estimator for yesterday's state based on observed actions. We can

⁵¹Condition 1 is necessary to rule out response profiles such as the one given by $\xi_{i,t}(s_{i,t}, a_{i,t-1}) = |a_{i,t-1}| + 1$. This profile, despite consisting of well-behaved functions, does not correspond to any specification of behavior for the whole population (because time extends infinitely backward). Condition 2 is necessary to rule out response profiles such as the one given by $\xi_{i,t}(s_{i,t}, a_{i,t-1}) = a_{i,t-1}$, which have many satisfying action paths, leaving payoffs undetermined.

⁵²I.e., $\sigma' = \mathfrak{T}_p \sigma$ is defined by $\sigma'_{(i,t)} = \sigma_{(i,t-p)}$.

analyze these variances using the envelope theorem. Moreover, the space of these variances is a sufficient statistic for determining all agent strategies and action variances.

Let $r_{i,t}$ be i 's *social signal*—the best estimator of θ_{t-1} based on the period $t-1$ actions of agents in N_i —and let $\kappa_{i,t}^2$ be the variance of $r_{i,t} - \theta_{t-1}$.

We claim that Φ induces a map $\tilde{\Phi}$ on the space of variances $\kappa_{i,t}^2$, which we denote $\tilde{\mathcal{V}}$. We must check the period t variances $(\kappa_{i,t}^2)_i$ uniquely determine all period $t+1$ variances $(\kappa_{i,t+1}^2)_i$: The variance $\mathbf{V}_{ii,t}$ of agent i 's action, as well as the covariances $\mathbf{V}_{ii',t}$ of all pairs of agents i, i' with $N_i = N_{i'}$, are determined by $\kappa_{i,t}^2$. Moreover, by the condition on our network, these variances and covariances determine all agents' strategies in period $t+1$, and this is enough to pin down all period $t+1$ variances $\kappa_{i,t+1}^2$.

The proof proceeds by showing $\tilde{\Phi}$ is a contraction on $\tilde{\mathcal{V}}$ in the sup norm.

For each agent j , we have $N_i = N_{i'}$ for all $i, i' \in N_j$. So the period t actions of an agent i' in N_j are

$$a_{i',t} = \frac{(\rho^2 \kappa_{i,t}^2 + 1)^{-1}}{\sigma_{i'}^{-2} + (\rho^2 \kappa_{i,t}^2 + 1)^{-1}} \cdot r_{i,t} + \frac{\sigma_{i'}^{-2}}{\sigma_{i'}^{-2} + (\rho^2 \kappa_{i,t}^2 + 1)^{-1}} \cdot s_{i',t} \quad (\text{F.1})$$

where $s_{i',t}$ is agent (i') 's signal in period t and $r_{i,t}$ the social signal of i (the same one that i' has). It follows from this formula that each action observed by j is a linear combination of a private signal and a *common* estimator $r_{i,t}$, with positive coefficients which sum to one. For simplicity we write

$$a_{i',t} = b_0 \cdot r_{i,t} + b_{i'} \cdot s_{i',t} \quad (\text{F.2})$$

(where b_0 and $b_{i'}$ depend on i' and t , but we omit these subscripts). We will use the facts $0 < b_0 < 1$ and $0 < b_{i'} < 1$.

We are interested in how $\kappa_{j,t+1}^2 = \text{Var}(r_{j,t+1} - \theta_t)$ depends on $\kappa_{i,t}^2 = \text{Var}(r_{i,t} - \theta_{t-1})$. The estimator $r_{j,t+1}$ is a linear combination of observed actions $a_{i',t}$, and therefore can be expanded as a linear combination of signals $s_{i',t}$ and the estimator $r_{i,t}$. We can write

$$r_{j,t+1} = c_0 \cdot (\rho r_{i,t}) + \sum_{i'} c_{i'} s_{i',t} \quad (\text{F.3})$$

and therefore (taking variances of both sides)

$$\begin{aligned} \kappa_{j,t+1}^2 &= \text{Var}(r_{j,t+1} - \theta_t) = c_0^2 \text{Var}(\rho r_{i,t} - \theta_t) + \sum_{i'} c_{i'}^2 \sigma_{i'}^2 \\ &= c_0^2 (\rho^2 \kappa_{i,t}^2 + 1) + \sum_{i'} c_{i'}^2 \sigma_{i'}^2 \end{aligned}$$

The desired result, that $\tilde{\Phi}$ is a contraction, will follow if we can show that the derivative $\frac{d\kappa_{j,t+1}^2}{d\kappa_{i,t}^2} = c_0^2 \rho^2 \in [0, \delta]$ for some $\delta < 1$. By the envelope theorem, when calculating this derivative, we can assume that the weights placed on actions $a_{i',t}$ by the estimator $r_{j,t+1}$ do

not change as we vary $\kappa_{i,t}^2$, and therefore c_0 and the $c_{i'}$ above do not change. So it is enough to show the coefficient c_0 is in $[0, 1]$.

The intuition for the lower bound is that *anti-imitation* (agents placing negative weights on observed actions) only occurs if observed actions put too much weight on public information. But if $c_0 < 0$, then the weight on public information is actually negative so there is no reason to anti-imitate. This is formalized in the following lemma.

Lemma 8. Suppose j has symmetric neighbors. Then agent j 's social signal places non-negative weight on a neighbor i 's social signal from the previous period, i.e., $c_0 \geq 0$.

Proof. To check this formally, suppose that c_0 is negative. Then the social signal $r_{j,t+1}$ puts negative weight on some observed action—say the action $a_{k,t}$ of agent k . We want to check that the covariance of $r_{j,t+1} - \theta_t$ and $a_{k,t} - \theta_t$ is negative. Using (F.2) and (F.3), we compute that

$$\begin{aligned} \text{Cov}(r_{j,t+1} - \theta_t, a_{k,t} - \theta_t) &= \text{Cov}\left(c_0(\rho r_{i,t} - \theta_t) + \sum_{i' \in N_j} c_{i'}(s_{i',t} - \theta_t), b_0(\rho r_{i,t} - \theta_t) + b_k(s_{k,t} - \theta_t)\right) \\ &= c_0 b_0 \text{Var}(\rho r_{i,t} - \theta_t) + c_k b_k \text{Var}(s_{k,t} - \theta_t) \end{aligned}$$

because all distinct summands above are mutually independent. We have $b_0, b_k > 0$, while $c_0 < 0$ by assumption and $c_k < 0$ because the estimator $r_{j,t+1}$ puts negative weight on $a_{k,t}$. So the expression above is negative. Therefore, it follows from the usual Gaussian Bayesian updating formula that the best estimator of θ_t given $r_{j,t+1}$ and $a_{k,t}$ puts positive weight on $a_{k,t}$. However, this is a contradiction: the best estimator of θ_t given $r_{j,t+1}$ and $a_{k,t}$ is simply $r_{j,t+1}$, because $r_{j,t+1}$ was defined as the best estimator of θ_t given observations that included $a_{k,t}$. This completes the proof of Lemma 8. $\square$

We now complete the proof of Lemma 7. For the upper bound $c_0 \leq 1$, the idea is that $r_{j,t+1}$ puts more weight on agents with better signals while these agents put little weight on public information, which keeps the overall weight on public information from growing too large.

Note that $r_{j,t+1}$ is a linear combination of actions $\rho a_{i',t}$ for $i' \in N_j$, with coefficients summing to 1. The only way the coefficient on $\rho r_{i,t}$ in $r_{j,t+1}$ could be at least 1 would be if some of these coefficients on $\rho a_{i',t}$ were negative and the estimator $r_{j,t+1}$ placed greater weight on actions $a_{i',t}$ which placed more weight on $r_{i,t}$.

Applying the formula (F.1) for $a_{i',t}$, we see that the coefficient b_0 on $\rho r_{i,t}$ is less than 1 and increasing in $\sigma_{i'}$. On the other hand, it is clear that the weight on $a_{i',t}$ in the social signal $r_{j,t+1}$ is decreasing in $\sigma_{i'}$: more weight should be put on more precise individuals. So in fact the estimator $r_{j,t+1}$ places less weight on actions $a_{i',t}$ which placed more weight on $r_{i,t}$.

Moreover, the coefficients placed on private signals are bounded below by a positive constant when we restrict to covariances in the image of $\tilde{\Phi}$ (because all covariances are bounded

as in the proof of Proposition 1). Therefore, each agent $i' \in N_j$ places weight at most one on the estimator $\rho r_{i,t-1}$. Agent j 's social signal $r_{j,t+1}$ is a sum of these agents' actions with coefficients summing to 1 and satisfying the monotonicity property above. We conclude that the coefficient on $\rho r_{i,t}$ in the expression for $r_{j,t+1}$ is at most one. This completes the proof of Lemma 7. $\square$

We now prove Proposition 2.

Proof of Proposition 2. By Lemma 7 there is a unique equilibrium on any network G with symmetric neighbors. Let $\varepsilon > 0$.

Consider any agent i . Her neighbors have the same private signal qualities and the same neighborhoods (by the symmetric neighbors assumption). So there exists an equilibrium where for all i , the actions of agent i 's neighbors are exchangeable. By uniqueness, this in fact holds at the sole equilibrium.

So agent i 's social signal is an average of her neighbors' actions:

$$r_{i,t} = \frac{1}{|N_i|} \sum_{j \in N_i} a_{j,t-1}.$$

Suppose the ε -aggregation benchmark is achieved. Then all agents must place weight at least $\frac{(1+\varepsilon)^{-1}}{(1+\varepsilon)^{-1} + \sigma^{-2}}$ on their social signals. So at time t , the social signal $r_{i,t}$ places weight at least $\frac{(1+\varepsilon)^{-1}}{(1+\varepsilon)^{-1} + \sigma^{-2}}$ on signals from at least two periods ago. Since the variance of any linear combination of signals from at least two periods ago, with weights summing to one, is at least $1 + \rho$, it follows that for ε sufficiently small the social signal $r_{i,t}$ is bounded away from a perfect estimate of θ_{t-1} . This gives a contradiction. $\square$

F.2. Proof of Corollary 1. Consider a complete graph in which all agents have signal variance σ^2 and memory $m = 1$. By Proposition 2, as n grows large the variances of all agents converge to $A > (1 + \sigma^{-2})^{-1}$.

Choose σ^2 large enough such that $A > 1$. To see that we can do this, note that as σ^2 grows large, the weight each agent places on their private signal vanishes. So the weight on signals from at least k periods ago approaches one for any k . Taking σ^2 such that this holds for k sufficiently large, we have $A > 1$.

Now suppose that we increase σ_1^2 to ∞ . Then $a_{1,t} = r_{1,t}$ in each period, so all agents can infer all private signals from the previous period. As n grows large, the variance of agent 1 converges to 1 and the variances of all other agents converge to $(1 + \sigma^{-2})^{-1}$. By our choice of σ^2 , this gives a Pareto improvement. We can see by continuity that the same argument holds for σ_1^2 finite but sufficiently large.

F.3. Proof of Corollary 2. Our goal is to estimate $\text{Var}(a_{i,t} - \theta_t)$.

First, observe that

$$\begin{aligned} a_{i,t} &= w_s s_{i,t} + (1 - w_s) \frac{1}{n} \sum_j a_{j,t-1} \\ &= w_s (\theta_t + \varepsilon_{i,t}) + (1 - w_s) \frac{1}{n} \sum_j a_{j,t-1}. \end{aligned}$$

This implies, inductively, that

$$a_{i,t} - \theta_t = w_s \varepsilon_{i,t} + \sum_{\ell=1}^{\infty} (w_s (1 - w_s)^\ell (\theta_t - \theta_{t-\ell} + \zeta_t)),$$

where the ζ_t are mean-zero random variables independent of all other random variables in the expression. (They are linear combinations of agents' signal noise realizations.) Thus,

$$\text{Var}(a_{i,t} - \theta_t) \geq \text{Var} \left(w_s \sum_{\ell=1}^{\infty} (1 - w_s)^\ell (\theta_t - \theta_{t-\ell}) \right)$$

Noting that $\theta_t = \sum_{k=0}^{\infty} \rho^k \nu_{t-k}$, we may write

$$\theta_t - \theta_{t-\ell} = \sum_{k=0}^{\ell-1} \rho^k \nu_{t-k}$$

and therefore

$$\begin{aligned} \text{Var}(a_{i,t} - \theta_t) &\geq \text{Var} \left(w_s \sum_{\ell=1}^{\infty} (1 - w_s)^\ell \sum_{k=0}^{\ell-1} \rho^k \nu_{t-k} \right) \\ &= w_s^2 \text{Var} \left(\sum_{k=0}^{\infty} \nu_{t-k} \rho^k \sum_{\ell=k+1}^{\infty} (1 - w_s)^\ell \right) \\ &= w_s^2 \text{Var} \left(\sum_{k=0}^{\infty} \nu_{t-k} \left(\rho^k \sum_{\ell=k+1}^{\infty} (1 - w_s)^\ell \right) \right) \\ &= w_s^2 \sum_{k=0}^{\infty} \left(\rho^k \sum_{\ell=k+1}^{\infty} (1 - w_s)^\ell \right)^2 \\ &= \frac{(1 - w_s)^2}{1 - (1 - w_s)^2 \rho^2}. \end{aligned}$$

This proves the bound

$$\text{Var}(a_{i,t} - \theta_t) \geq \frac{(1 - w_s)^2}{1 - (1 - w_s)^2 \rho^2}.$$

It remains to show the variances diverge to infinity as $\sigma^2 \rightarrow \infty$ and $\rho \rightarrow 1$ from below. Choose a sequence of pairs $(\sigma^2, \rho) \rightarrow (\infty, 1)$. If $w_s \rightarrow 0$ along any subsequence of this sequence, then along the subsequence we have $\frac{(1-w_s)^2}{1-(1-w_s)^2 \rho^2} \rightarrow \infty$ and so $\text{Var}(a_{i,t} - \theta_t) \rightarrow \infty$ as well. If w_s is non-vanishing, then $\text{Var}(a_{i,t} - \theta_t) \rightarrow \infty$ since the action variance is at

least $w_s^2 \sigma^2$ and $\sigma^2 \rightarrow \infty$. Finally, note that these bounds are both independent of n , so $\text{Var}(a_{i,t} - \theta_t) \rightarrow \infty$ uniformly in n .

F.4. Proof of Theorem 2. Suppose that all private signals have variance $\sigma^2 > 0$. Fix a sequence of networks G_n and an equilibrium on each G_n . We will show that given any constant $C > 0$ and any sequence of equilibria, the fraction of agents i such that

$$\widehat{\kappa}_i^2 \leq \frac{C}{\overline{d}}$$

is bounded away from one.

We first prove the result in the case $m = 1$. For each n , let $\mathcal{G}_n$ be the set of agents i satisfying

$$\widehat{\kappa}_i^2 \leq \frac{C}{\overline{d}},$$

i.e., the set of agents who do learn well. Assume for the sake of contradiction that $\frac{|\mathcal{G}_n|}{n} \rightarrow 1$ as $n \rightarrow \infty$ along some subsequence and pass to that subsequence.

For each j , we can express the action $a_{j,t}$ as a weighted sum of innovations and signal errors,⁵³ with all terms on the right-hand side conditionally independent:

$$a_{j,t} = \theta_t - \sum_{l=0}^{\infty} w_{j,t}(\nu_{t-l})(\rho^l \nu_{t-l}) + \sum_{l,j'} w_{j,t}(\eta_{j',t-l})(\rho^l \eta_{j',t-l}).$$

This expression is unique.

Lemma 9. For all $j \in \mathcal{G}_n$ we must have

$$w_{j,t}(\nu_t) \in \left(\frac{1}{\sigma^{-2} + 1} - \frac{C'}{\overline{d}}, \frac{1}{\sigma^{-2} + 1} \right)$$

for some $C' > 0$ (independent of j and n).

Proof. By the standard updating formula, the optimal weight $w_{j,t}(\nu_t)$ is $\frac{(\rho^2 \kappa_{j,t}^2 + 1)^{-1}}{(\rho^2 \kappa_{j,t}^2 + 1)^{-1} + \sigma^{-2}}$, where $\kappa_{j,t}^2$ is the variance of the best estimator of θ_{t-1} based on (j, t) 's social observations. The upper bound follows because this is minimized when $\kappa_{j,t}^2 = 0$. For the lower bound,

$$\begin{aligned} w_{j,t}(\nu_t) &= \frac{(\rho^2 \kappa_{j,t}^2 + 1)^{-1}}{(\rho^2 \kappa_{j,t}^2 + 1)^{-1} + \sigma^{-2}} \\ &= \frac{1}{(1 + \sigma^{-2}) + \sigma^{-2} \rho^2 \kappa_{j,t}^2} \\ &= \frac{1}{1 + \sigma^{-2}} - \frac{\sigma^{-2} \rho^2}{(1 + \sigma^{-2})^2} \kappa_{j,t}^2 + O(\kappa_{j,t}^4). \end{aligned}$$

⁵³To simplify calculations, we write this expression with a negative coefficient on the first sum so that the terms $w_{j,t}(\nu_{t-l})$ are positive. The weight that j places on ν_{t-l} is in fact $-w_{j,t}(\nu_{t-l})$.

For $\kappa_{j,t}^2$ in any neighborhood of zero, we can choose C'' such that the non-constant terms in the final expression are bounded below by $-C''\kappa_{j,t}^2$. Since by assumption we have $\kappa_{j,t}^2 \leq \frac{C}{d}$, the lemma follows with $C' = C \cdot C''$. $\square$

There are at most $(n - |\mathcal{G}_n|)\bar{d}$ in-coming links to agents outside $\mathcal{G}_n$. Each agent who observes at least $\frac{2(n-|\mathcal{G}_n|)}{n} \cdot \bar{d}$ agents outside $\mathcal{G}_n$ accounts for at least $\frac{2(n-|\mathcal{G}_n|)}{n} \cdot \bar{d}$ of those links, so there can be at most $\frac{n}{2}$ such agents. Since $\frac{|\mathcal{G}_n|}{n} \rightarrow 1$, there is an agent $i \in \mathcal{G}_n$ who observes fewer than $\frac{2(n-|\mathcal{G}_n|)}{n} \cdot \bar{d}$ agents outside $\mathcal{G}_n$.

Consider the action of such an agent $i \in \mathcal{G}_n$ in period $t + 1$. Since

$$\kappa_{i,t+1}^2 \leq \frac{C}{\bar{d}},$$

the weight on the innovation from the *previous* period satisfies

$$(w_{i,t+1}(\nu_t)\rho)^2 \leq \frac{C}{\bar{d}}. \quad (\text{F.4})$$

On the other hand, we can express this weight in terms of neighbors' weights as

$$w_{i,t+1}(\nu_t) = \sum_j \rho w_{ij,t+1} w_{j,t}(\nu_t).$$

We will show that if this weight $w_{i,t+1}(\nu_t)$ vanishes, then the contribution of private signal errors to $\kappa_{i,t+1}$ must be larger than $O(\frac{1}{\bar{d}})$.

We can split this summation as

$$w_{i,t+1}(\nu_t) = \rho \sum_{j \in \mathcal{G}_n} w_{ij,t+1} w_{j,t}(\nu_t) + \rho \sum_{j \notin \mathcal{G}_n} w_{ij,t+1} w_{j,t}(\nu_t).$$

We now consider two cases, depending on whether $\sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}| \rightarrow 0$, i.e., whether the sum of the absolute values of the weights on agents outside $\mathcal{G}_n$ is vanishing.

Case 1: $\liminf_n \sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}| = 0$. We can pass to a subsequence along which $\sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}| \rightarrow 0$.

We claim that it follows from the bounds on $w_{j,t}(\nu_t)$ in Lemma 9 that this can only occur if $\sum_j |w_{ij,t+1}| \rightarrow \infty$. If $\sum_j |w_{ij,t+1}|$ is bounded,

$$w_{i,t+1}(\nu_t) = \rho \sum_{j \in \mathcal{G}_n} w_{ij,t+1} w_{j,t}(\nu_t) + \rho \sum_{j \notin \mathcal{G}_n} w_{ij,t+1} w_{j,t}(\nu_t) = \rho \sum_{j \in \mathcal{G}_n} w_{ij,t+1} w_{j,t}(\nu_t) + o(1).$$

The second equality holds because $\sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}| \rightarrow 0$ and $w_{j,t}(\nu_t) \in [0, 1]$ for all j . Therefore,

$$w_{i,t+1}(\nu_t) = \rho \sum_{j \in \mathcal{G}_n} w_{ij,t+1} w_{j,t}(\nu_t) + o(1) \geq \left(\rho \frac{1}{1 + \sigma^{-2}} \frac{\sigma^{-2}}{\sigma^{-2} + 1} - o(1) \right) + o(1),$$

and the right-hand side is non-vanishing. Here the first term on the right-hand side is the limit of the sum if all of the terms $w_{j,t}(\nu_t)$ were equal to the upper bound $\frac{\sigma^{-2}}{\sigma^{-2}+1}$. The first $o(1)$ error term corresponds to the variation in $w_{j,t}(\nu_t)$ across j , which is $O(\frac{1}{d})$ by Lemma 9 and has bounded coefficients. Thus $w_{i,t+1}(\nu_t)$ is non-vanishing, but this contradicts the inequality (F.4). We have proven the claim.

The contribution to $\kappa_{i,t+1}^2$ from signal errors $\eta_{j,t}$ is $\sum_j |w_{ij,t+1}|^2 (w_{j,t}^s)^2 \sigma^2$. Since $w_{j,t}^s = 1 - w_{j,t}(\nu_t)$ converge uniformly to a constant $\frac{\sigma^{-2}}{\sigma^{-2}+1}$, we can bound this contribution below by an expression that is proportional to

$$\sum_j |w_{ij,t+1}|^2.$$

The summation has at most $\bar{d}$ non-zero terms. Applying the standard bound $\|v\|_1 \leq \sqrt{n}\|v\|_2$ on L^p norms on $\mathbb{R}^n$,

$$\sum_j |w_{ij,t+1}|^2 \geq \frac{1}{\bar{d}} \left(\sum_j |w_{ij,t+1}| \right)^2.$$

The right-hand side of this inequality grows at a rate faster than $\frac{1}{d}$ by the claim $\sum_j |w_{ij,t+1}| \rightarrow \infty$, and so the social signal error grows at a rate faster than $\frac{1}{d}$. This gives a contradiction.

Case 2: $\liminf_n \sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}| > 0$.

As in Case 1, the contribution to signal errors from neighbors $j \notin \mathcal{G}_n$ is proportional to

$$\sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}|^2.$$

By our choice of the agent i , she observes at most $\frac{2(n-|\mathcal{G}_n|)}{n} \cdot \bar{d}$ agents outside $\mathcal{G}_n$. The same standard bound on L^p norms gives

$$\sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}|^2 \geq \frac{1}{\bar{d}} \cdot \frac{n}{2(n-|\mathcal{G}_n|)} \left(\sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}| \right)^2.$$

By assumption, the cardinality $n - |\mathcal{G}_n|$ of the complement of $\mathcal{G}_n$ is $o(n)$ and $(\sum_{j \notin \mathcal{G}_n} |w_{ij,t+1}|)^2$ is non-vanishing. So the right-hand side grows at a rate faster than $\frac{1}{d}$. Thus the social signal error grows at a rate faster than $\frac{1}{d}$, which again gives a contradiction. This completes the proof in the case $m = 1$, and we next turn to the general argument.

Now, suppose $m \geq 1$ is arbitrary. As before, for each agent (j, t) , we can write:

$$a_{j,t} = \theta_t - \sum_{l=0}^{\infty} w_{j,t}(\nu_{t-l})(\rho^l \nu_{t-l}) + \sum_{l,j'} w_{j,t}(\eta_{j',t-l})(\rho^l \eta_{j',t-l}).$$

For each n , let $\mathcal{G}_n$ be the set of i satisfying

$$\hat{\kappa}_i^2 \leq \frac{C}{\bar{d}}.$$

Suppose $\limsup_n |\mathcal{G}_n|/n = 1$. Passing to a subsequence, we can assume that $\lim_n |\mathcal{G}_n|/n = 1$, i.e., the fraction of agents in $\mathcal{G}_n$ converges to one.

As in the $m = 1$ proof above, we can choose $i \in \mathcal{G}_n$ who observes fewer than $\frac{2(n-|\mathcal{G}_n|)}{n} \cdot \bar{d}$ agents outside $\mathcal{G}_n$. Choose any such i and consider the agent (i, t) with $i \in \mathcal{G}_n$, who observes neighbors' actions in periods $t - 1, \dots, t - m$. For each $1 \leq l \leq m$, we will write $w_{(i,t),(j,t-l)}$ for the weight that agent (i, t) places on the action of agent $(j, t - l)$. By the same argument as in Case 2 of the $m = 1$ proof above, $\liminf_n \sum_{j \notin \mathcal{G}_n} |w_{(i,t),(j,t-l)}| = 0$ for each l (since the fraction of agents outside $\mathcal{G}_n$ is vanishing). Passing to a subsequence, we can assume that $\lim_n \sum_{j \notin \mathcal{G}_n} |w_{(i,t),(j,t-l)}| = 0$.

We can express agent (i, t) 's action:

$$a_{i,t} = \sum_{1 \leq l \leq m} \left(\sum_{j \in \mathcal{G}_n} w_{(i,t),(j,t-l)} \rho^l a_{j,t-l} + \sum_{j \notin \mathcal{G}_n} w_{(i,t),(j,t-l)} \rho^l a_{j,t-l} \right).$$

We will show that this expression places non-vanishing weight on the innovation ν_{t-l} for some $l \geq 1$. This will contradict our assumption that $i \in \mathcal{G}_n$.

Since $\lim_n \sum_{j \notin \mathcal{G}_n} |w_{(i,t),(j,t-l)}| = 0$ and the weight each agent places on ν_{t-l} is bounded, it is sufficient to show that

$$\sum_{1 \leq l \leq m} \sum_{j \in \mathcal{G}_n} w_{(i,t),(j,t-l)} \rho^l a_{j,t-l}$$

places non-vanishing weight on the innovation ν_{t-l} for some $l \geq 1$.

For each (j, t') such that $j \in \mathcal{G}_n$, we have

$$a_{j,t'} = \frac{\theta_{t'-1} + \sigma^{-2} s_{j,t'}}{1 + \sigma^{-2}} + \epsilon_{j,t'},$$

where $\text{Var}(\epsilon_{j,t'}) \rightarrow 0$. This is because

$$a_{j,t'} = \frac{(\rho^2 \kappa_{i,t}^2 + 1)^{-1} r_{i,t} + \sigma^{-2} s_{j,t'}}{(\rho^2 \kappa_{i,t}^2 + 1)^{-1} + \sigma^{-2}},$$

and we have $\kappa_{i,t}^2 = \text{Var}(r_{i,t} - \theta_{t-1}) \rightarrow 0$.

Using this expression for $a_{j,t'}$, we obtain

$$\sum_{1 \leq l \leq m} \sum_{j \in \mathcal{G}_n} w_{(i,t),(j,t-l)} \rho^l a_{j,t-l} = \sum_{1 \leq l \leq m} \sum_{j \in \mathcal{G}_n} w_{(i,t),(j,t-l)} \rho^l \left(\frac{\theta_{t-l-1} + \sigma^{-2} s_{j,t-l}}{1 + \sigma^{-2}} + \epsilon_{j,t-l} \right)$$

By the same argument as in Case 1 of the $m = 1$ proof above,

$$\sum_{1 \leq l \leq m} \sum_{j \in \mathcal{G}_n} |w_{(i,t),(j,t-l)}|$$

must be bounded (or else the contributions of signal errors to $\widehat{\kappa}_i^2$ would be too large to have $i \in \mathcal{G}_n$). Therefore, it is sufficient to show that

$$\sum_{1 \leq l \leq m} \sum_{j \in \mathcal{G}_n} w_{(i,t),(j,t-l)} \rho^l \cdot \frac{\theta_{t-l-1} + \sigma^{-2} s_{j,t-l}}{1 + \sigma^{-2}}$$

places non-vanishing weight on the innovation ν_{t-l} for some $l \geq 1$.

This holds for the largest l such that $\sum_{j \in \mathcal{G}_n} w_{(i,t),(j,t-l)}$ is non-vanishing. Such an l must exist, because

$$\sum_{1 \leq l \leq m} \sum_{j \in \mathcal{G}_n} w_{(i,t),(j,t-l)} \rightarrow \frac{1}{1 + \sigma^{-2}}$$

since $i \in \mathcal{G}_n$.

F.5. Proof of Proposition 3. For each agent i , we can write

$$a_{i,t} = w_i^s s_{i,t} + \sum_j W_{ij} \rho a_{j,t-1} = w_i^s s_{i,t} + \sum_j W_{ij} \left(\rho w_j^s s_{j,t} + \sum_{j'} W_{jj'} \rho a_{j',t-2} \right).$$

Because we assume $w_i^s < \bar{w} < 1$ and $w_j^s < \bar{w} < 1$ for all j , the total weight $\sum_{j,j'} W_{ij} W_{jj'} \rho$ on terms $a_{j',t-2}$ is bounded away from zero. Because the error variance of each of these terms is greater than 1, this implies agent i fails to achieve the ε -aggregation benchmark for $\varepsilon > 0$ sufficiently small.

F.6. Proof of Proposition 4. We prove the following statement, which includes the proposition as special cases.

Proposition 7. Suppose the network G is strongly connected.⁵⁴ Consider weights $\mathbf{W}$ and $\mathbf{w}^s$ and suppose they are all positive, with an associated steady state $\mathbf{V}_t$. Suppose either

(1) there is an agent i whose weights are a Bayesian best response to $\mathbf{V}_t$, and some agent observes that agent and at least one other neighbor; or

(2) there is an agent whose weights are a naive best response to $\mathbf{V}_t$, and who observes multiple neighbors.

Then the steady state $\mathbf{V}_t$ is Pareto-dominated by another steady state.

We provide the proof in the case $m = 1$ to simplify notation. The argument carries through with arbitrary finite memory.

Case (1): Consider an agent l who places positive weight on a rational agent k and positive weight on at least one other agent. Define weights $\bar{\mathbf{W}}$ by $\bar{W}_{ij} = W_{ij}$ and $\bar{w}_i^s = w_i^s$ for all $i \neq k$, $\bar{W}_{kj} = (1 - \epsilon)W_{kj}$ for all $j \leq n$, and $\bar{w}_k^s = (1 - \epsilon)w_k^s + \epsilon$, where W_{ij} and w_i^s are the weights at the initial steady state. In words, agent k places weight $(1 - \epsilon)$ on her equilibrium strategy and extra weight ϵ on her private signal. All other players use the same weights as at the steady state.

⁵⁴That is, there is a directed path from each node to each other node.

Suppose we are at the initial steady state until time t , but in period t and all subsequent periods agents instead use weights $\bar{W}$. These weights give an alternate updating function $\bar{\Phi}$ on the space of covariance matrices. Because the weights $\bar{W}$ are positive and fixed, all coordinates of $\bar{\Phi}$ are increasing, linear functions of all previous period variances and covariances. Explicitly, the diagonal terms are

$$[\bar{\Phi}(\mathbf{V}_t)]_{ii} = (\bar{w}_i^s)^2 \sigma_i^2 + \sum_{j,j' \leq n} \bar{W}_{ij} \bar{W}_{ij'} (\rho^2 V_{jj',t} + 1)$$

and the off-diagonal terms are

$$[\bar{\Phi}(\mathbf{V}_t)]_{ii'} = \sum_{j,j' \leq n} \bar{W}_{ij} \bar{W}_{i'j'} (\rho^2 V_{jj,t'} + 1).$$

So it is sufficient to show the variances $\bar{\Phi}^h(\mathbf{V}_t)$ after applying $\bar{\Phi}$ for h periods Pareto dominate the variances in $\mathbf{V}_t$ for some h .

In period t , the change in weights decreases the covariance $V_{jk,t}$ of k and some other agent j , who l also observes, by $f(\epsilon)$ of order $\Theta(\epsilon)$. By the envelope theorem, the change in weights only increases the variance V_{kk} by $O(\epsilon^2)$. Taking ϵ sufficiently small, we can ignore $O(\epsilon^2)$ terms.

There exists a constant $\delta > 0$ such that all initial weights on observed neighbors are at least δ . Then each coordinate $[\Phi(\mathbf{V})]_{ii}$ is linear with coefficient at least δ^2 on each variance or covariance of agents observed by i .

Because agent l observes k and another agent, agent l 's variance will decrease below its equilibrium level by at least $\delta^2 f(\epsilon)$ in period $t + 1$. Because $\bar{\Phi}$ is increasing in all entries and we are only decreasing covariances, agent l 's variance will also decrease below its initial level by at least $\delta^2 f(\epsilon)$ in all periods $t' > t + 1$.

Because the network is strongly connected and finite, the network has a diameter. After $d + 1$ periods, the variances of all agents have decreased by at least $\delta^{2d+2} f(\epsilon)$ from their initial levels. This gives a Pareto improvement.

Case (2): Consider a naive agent k who observes at least two neighbors. We can write agent k 's period t action as

$$a_{k,t} = w_k^s s_{i,t} + \sum_{j \in N_i} W_{kj} \rho a_{j,t-1}.$$

Define new weights $\bar{W}$ as in the proof of case (1). Because agent k is naive and the summation $\sum_{j \in N_i} W_{kj} \rho a_{j,t-1}$ has at least two terms, she believes the variance of this summation is smaller than its true value. So marginally increasing the weight on $s_{k,t}$ and decreasing the weight on this summation decreases her action variance. This deviation also decreases her covariance with any other agent. The remainder of the proof proceeds as in case (1).

F.7. Proof of Proposition 5. Suppose the social influence

$$\text{SI}(i) = \sum_{j \in N} \sum_{k=1}^{\infty} \left(\rho \widehat{W} \right)_{ji}^k \widehat{w}_i^s = \left[\mathbf{1}' \left(I - \rho \widehat{W} \right)^{-1} - \mathbf{1}' \right]_i \widehat{w}_i^s$$

does not converge for some i . Then in particular, there exists j such that $\sum_{k=0}^{\infty} \left(\rho \widehat{W} \right)_{ji}^k \widehat{w}_i^s$ does not converge. We can write

$$a_{j,t} = \sum_{\tau=0}^{\infty} \sum_{j' \in N} \left(\rho \widehat{W} \right)_{jj'}^{\tau} \widehat{w}_{j'}^s s_{j',t-\tau}.$$

This expression is the sum of

$$\sum_{\tau=0}^{\infty} \left(\rho \widehat{W} \right)_{ji}^{\tau} \widehat{w}_i^s \eta_{i,t-\tau}$$

and independent terms corresponding to signal errors of agents other than i and changes in the state. Because $\sum_{\tau=0}^{\infty} \left(\rho \widehat{W} \right)_{ji}^{\tau} \widehat{w}_i^s$ does not converge, the payoff to action $a_{j,t}$ must therefore be $-\infty$. But we showed in the proof of Proposition 1 that agent j 's equilibrium payoff is at least $-\sigma_j^2$, which gives a contradiction.

Given convergence, the expression for $\text{SI}(i)$ follows from the identity $(I - M)^{-1} = \sum_{k=0}^{\infty} M^k$.

F.8. Proof of Proposition 6. The social signal $r_{i,t}$ is the same for all agents, and we will refer to it as r_t . We can express the social signal as

$$r_t = w_A \sum_{i: \sigma_i = \sigma_A} a_{i,t-1} + w_B \sum_{i: \sigma_i = \sigma_B} a_{i,t-1} \quad (\text{F.5})$$

for some weights w_A and w_B .

We can rewrite the actions $a_{i,t-1}$ for i with signal variance σ_A^2 as

$$a_{i,t-1} = \frac{K}{K + \sigma_A^{-2}} \rho r_{i,t-1} + \frac{\sigma_A - 2}{K + \sigma_A^{-2}} s_{i,t-1},$$

where $K = \rho^2 \kappa_{t-1}^2 + 1$ is the equilibrium variance of ρr_{t-1} about the state θ_{t-1} . The analogous formula holds for agents i with signal variance σ_B^2 .

Substituting the formulas for $a_{i,t-1}$ into equation (F.5) and taking variances,

$$\begin{aligned}
\kappa_t^2 &= \text{Var}(r_t - \theta_{t-1}) \\
&= \text{Var}\left(\frac{nK}{2} \left(\frac{w_A}{K + \sigma_A^{-2}} + \frac{w_B}{K + \sigma_B^{-2}}\right) (r_{t-1} - \theta_{t-1})\right) + \text{Var}\left(w_A \sum_{i \in S_A} \frac{\sigma_A^{-2}}{K + \sigma_A^{-2}} (s_{i,t-1} - \theta_{t-1})\right) \\
&\quad + \text{Var}\left(w_B \sum_{i \in S_B} \frac{\sigma_B^{-2}}{K + \sigma_B^{-2}} (s_{i,t-1} - \theta_{t-1})\right) \\
&= \frac{n^2 K}{4} \left(\frac{w_A}{K + \sigma_A^{-2}} + \frac{w_B}{K + \sigma_B^{-2}}\right)^2 + \frac{n}{2} \cdot \frac{w_A^2 \sigma_A^{-2}}{(K + \sigma_A^{-2})^2} + \frac{n}{2} \cdot \frac{w_B^2 \sigma_B^{-2}}{(K + \sigma_B^{-2})^2}.
\end{aligned}$$

The equilibrium weights $\hat{w}_A$ and $\hat{w}_B$ minimize this expression. Using the fact that $\hat{w}_A + \hat{w}_B = \frac{2}{n}$, we have that $\hat{w}_A$ satisfies

$$\begin{aligned}
(\kappa^2)'_{t+1}(\hat{w}_A) &= \frac{n^2 K}{2} \left(\frac{\hat{w}_A}{K + \sigma_A^{-2}} + \frac{\hat{w}_B}{K + \sigma_B^{-2}}\right) \left(\frac{1}{K + \sigma_A^{-2}} - \frac{1}{K + \sigma_B^{-2}}\right) + \frac{n \hat{w}_A \sigma_A^{-2}}{(K + \sigma_A^{-2})^2} - \frac{n \hat{w}_B \sigma_B^{-2}}{(K + \sigma_B^{-2})^2} \\
&= 0.
\end{aligned}$$

This equation, along with $\hat{w}_A + \hat{w}_B = \frac{2}{n}$, allows us to explicitly solve for $\hat{w}_A$ and $\hat{w}_B$ in terms of k and exogenous variables. In particular, we get that

$$\frac{\hat{w}_A}{\hat{w}_B} = \left(\frac{K + \sigma_A^{-2}}{K + \sigma_B^{-2}}\right) \left(\frac{2 + Kn\sigma_B^2 - K(n-2)\sigma_A^2}{2 + Kn\sigma_A^2 - K(n-2)\sigma_B^2}\right). \quad (\text{F.6})$$

We now turn to analyzing social influences. Recall that

$$\text{SI}(i) = \sum_{j=1}^n \sum_{k=1}^{\infty} \left(\rho \widehat{W}\right)_{ij}^k \hat{w}_i^s. \quad (\text{F.7})$$

On the complete graph, this expression is proportional to the product of the weight placed on agent i by the social signal r_t and agent i 's self-weight w_i^s . Therefore, we compute

$$\frac{\text{SI}(A)}{\text{SI}(B)} = \frac{\hat{w}_A}{\hat{w}_B} \cdot \frac{\frac{\sigma_A^{-2}}{K + \sigma_A^{-2}}}{\frac{\sigma_B^{-2}}{K + \sigma_B^{-2}}}.$$

Substituting from equation (F.6),

$$\frac{\text{SI}(A)}{\text{SI}(B)} = \left(\frac{\sigma_A^{-2}}{\sigma_B^{-2}}\right) \left(\frac{2 + Kn\sigma_B^2 - K(n-2)\sigma_A^2}{2 + Kn\sigma_A^2 - K(n-2)\sigma_B^2}\right).$$

We want to show that the left-hand side is greater than $\frac{\sigma_A^{-2}}{\sigma_B^{-2}}$ whenever $\sigma_A^{-2} > \sigma_B^{-2}$, which is equivalent to showing

$$\frac{2 + Kn\sigma_B^2 - K(n-2)\sigma_A^2}{2 + Kn\sigma_A^2 - K(n-2)\sigma_B^2} > 1$$

whenever $\sigma_A^{-2} > \sigma_B^{-2}$ and this fraction is positive.

To see this, note that the difference between the numerator and denominator of the fraction is

$$\begin{aligned} (Kn\sigma_B^2 - K(n-2)\sigma_A^2) - (Kn\sigma_A^2 - K(n-2)\sigma_B^2) &= 2K(n-1)(\sigma_B^2 - \sigma_A^2) \\ &> 0 \end{aligned}$$

as desired.

APPENDIX G. MODEL WITH A STARTING TIME (ONLINE APPENDIX)

In introducing the model (Section 2), we made the set of time indices $\mathcal{T}$ equal to $\mathbb{Z}$, the set of all integers. Here we study the variant with an initial time period, $t = 0$: thus, we take $\mathcal{T}$ to be $\mathbb{Z}_{\geq 0}$, the non-negative integers. This section shows that there is a unique equilibrium outcome. In large networks, a suitable analogue of Theorem 1 holds, with both aggregation quality and outcomes similar to those obtained there. Similarly, the negative result of Proposition 2 also has a counterpart in this model.

Let θ_0 be drawn according to the stationary distribution of the state process: $\theta_0 \sim \mathcal{N}\left(0, \frac{1}{1-\rho}\right)$. After this, the state random variables θ_t satisfy the AR(1) evolution

$$\theta_{t+1} = \rho\theta_t + \nu_{t+1},$$

where ρ is a constant with $0 < |\rho| < 1$ and $\nu_{t+1} \sim \mathcal{N}(0, \sigma_\nu^2)$ are independent innovations. Actions, payoffs, signals, and observations are the same as in the main model, with the obvious modification that in the initial periods, $t < m$, information sets are smaller as there are not yet prior actions to observe.⁵⁵ To save on notation, we write actions as if agents had an improper prior, understanding that the adjustment for actions taken under the natural prior $\theta_t \sim \mathcal{N}\left(0, \frac{1}{1-\rho}\right)$ is immediate.

In this model, there is a straightforward prediction of behavior. A Nash equilibrium here refers to an equilibrium of the game involving all agents (i, t) for all time indices in $\mathcal{T}$.

Fact 2. In the model with $\mathcal{T} = \mathbb{Z}_{\geq 0}$, there is a unique Nash equilibrium, and it is in linear strategies. The initial generation ($t = 0$) plays a linear strategy based on private signals only. In any period $t > 0$, given linear strategies from prior periods, players' best responses are linear. For time periods $t > m$, we have

$$\mathbf{V}_t = \Phi(\mathbf{V}_{t-1}).$$

This fact follows from the observation that the initial ($t = 0$) generation faces a problem of forming a conditional expectation of a Gaussian state based on Gaussian signals, so their

⁵⁵The actions for $t < 0$ can be set to arbitrary (commonly known) constants.

optimal strategies are linear. From then on, the analysis of Section 3.1 characterizes best-response behavior inductively. Note that for arbitrary environments, the fact does not imply that $\mathbf{V}_t$ must converge.

Our main purpose in this section is to give analogues of the main results on learning in large networks. We use the same definition of an environment—in terms of the distribution of networks and signals—as in Section 4.1. For simplicity, we work with $m = 1$, though the arguments for our positive result extend straightforwardly.

The analogue of Theorem 1 is:

Theorem 3. *Consider the $\mathcal{T} = \mathbb{Z}_{\geq 0}$ model. If an environment satisfies signal diversity, there is $C > 0$ such that asymptotically almost surely $\widehat{\kappa}_{i,t}^2 < C/n$ for all i at all times $t \geq 1$ in the unique Nash equilibrium.*

In particular, this implies that the covariance matrix in each period $t \geq 1$ is very close (in the Euclidean norm) to the good-learning equilibrium from Theorem 1. We sketch the proof, which uses the material we developed in Appendix B. We define A_t as in that proof (Section B.1). Take a $\beta > 0$, to be specified later, and consider

$$\overline{\mathcal{W}} = \mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}} \cup \widetilde{\Phi} \left(\mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}} \right).$$

First, for large enough β , we have that $A_1 \in \overline{\mathcal{W}}$: In the unique Nash equilibrium, at $t = 1$, agents simply take weighted averages of their neighbors' signals, weighted by their precisions. So $A_1 \in \overline{\mathcal{W}}$ by the central limit theorem for β sufficiently large. Second, we use the previously established fact (recall Section B.2.2) that $\widetilde{\Phi}(\overline{\mathcal{W}}) \subset \overline{\mathcal{W}}$ to deduce that $A_t \in \overline{\mathcal{W}}$ at all future times. Finally, we observe that $\overline{\mathcal{W}} \subseteq \mathcal{W}_{\frac{\beta}{n}, \frac{1}{n}}$ by construction.

Without signal diversity, bad learning can occur forever, in the unique equilibrium. The analogue of Proposition 2 is immediate. In graphs with symmetric neighbors, Φ is a contraction when $m = 1$. So iteration of it arrives at the unique fixed point, and thus a learning outcome far from the benchmark.

APPENDIX H. NAIVE AGENTS (ONLINE APPENDIX)

In this section we provide rigorous detail for the analysis given in 5.1. We will describe outcomes with two signal types, σ_A^2 and σ_B^2 .⁵⁶ We use the same random network model as in Section 4.2 and assume each network type contains equal shares of agents with each signal type.

We can define variances

$$V_A^\infty = \frac{\rho^2 \kappa_t^2 + 1 + \sigma_A^{-2}}{(1 + \sigma_A^{-2})^2}, \quad V_B^\infty = \frac{\rho^2 \kappa_t^2 + 1 + \sigma_B^{-2}}{(1 + \sigma_B^{-2})^2} \quad (\text{H.1})$$

⁵⁶The general case, with many signal types, is similar.

where

$$\kappa_t^{-2} = 1 - \frac{1}{(\sigma_A^{-2} + \sigma_B^{-2})} \left(\frac{\sigma_A^{-2}}{1 + \sigma_A^{-2}} + \frac{\sigma_B^{-2}}{1 + \sigma_B^{-2}} \right).$$

Naive agents' equilibrium variances converge to these values.

Proposition 8. Under the assumptions in this subsection:

- (1) There is a unique equilibrium on G_n .
- (2) Given any $\delta > 0$, asymptotically almost surely all agents' equilibrium variances are within δ of V_A^∞ and V_B^∞ .
- (3) There exists $\varepsilon > 0$ such that asymptotically almost surely the ε -aggregation benchmark is not achieved, and when $\sigma_A^2 = \sigma_B^2$ asymptotically almost surely all agents' variances are larger than V^∞ .

Aggregating information well requires a sophisticated response to the correlations in observed actions. Because naive agents completely ignore these correlations, their learning outcomes are poor. In particular their variances are larger than at the equilibria we discussed in the Bayesian case, even when that equilibrium is inefficient ($\sigma_A^2 = \sigma_B^2$).

When signal qualities are homogeneous ($\sigma_A^2 = \sigma_B^2$), we obtain the same limit on any network with enough observations. That is, on any sequence $(G_n)_{n=1}^\infty$ of (deterministic) networks with the minimum degree diverging to ∞ and any sequence of equilibria, the equilibrium action variances of all agents converge to V_A^∞ .

H.1. Proof of Proposition 8. We first check that there is a unique naive equilibrium. As in the Bayesian case, covariances are updated according to equations (3.3):

$$\mathbf{V}_{ii,t} = (w_{i,t}^s)^2 \sigma_i^2 + \sum W_{ik,t} W_{ik',t} (\rho^2 \mathbf{V}_{kk',t-1} + 1) \text{ and } \mathbf{V}_{ij,t} = \sum W_{ik,t} W_{i'k',t} (\rho^2 \mathbf{V}_{kk',t-1} + 1).$$

The weights $W_{ik,t}$ and $w_{i,t}^s$ are now all positive constants that do not depend on $\mathbf{V}_{t-1}$. So differentiating this formula, we find that all partial derivatives are bounded above by $1 - w_{i,t}^s < 1$. So the updating map (which we call Φ^{naive}) is a contraction in the sup norm on $\mathcal{V}$. In particular, there is at most one equilibrium.

The remainder of the proof characterizes the variances of agents at this equilibrium. We first construct a candidate equilibrium with variances converging to V_A^∞ and V_B^∞ , and then we show that for n sufficiently large, there exists an equilibrium nearby in $\mathcal{V}$.

To construct the candidate equilibrium, suppose that each agent observes the same number of neighbors of each signal type. Then there exists an equilibrium $\hat{\mathbf{V}}^{sym}$ where covariances depend only on signal types, i.e., $\hat{\mathbf{V}}^{sym}$ is invariant under permutations of indices that do not change signal types. We now show variances of the two signal types at this equilibrium converge to V_A^∞ and V_B^∞ .

To estimate θ_{t-1} , a naive agent combines observed actions from the previous period with weight proportional to their precisions σ_A^{-2} or σ_B^{-2} . The naive agent incorrectly believes this gives an almost perfect estimate of θ_{t-1} . So the weight on older observations vanishes as $n \rightarrow \infty$. The naive agent then combines this estimate of θ_{t-1} with her private signal, with weights converging to the weights she uses if the estimate is perfect.

Agent i observes $\frac{|N_i|}{2}$ neighbors of each signal type, so her estimate $r_{i,t}^{naive}$ of θ_{t-1} is approximately:

$$r_{i,t}^{naive} = \frac{2}{|N_i|(\sigma_A^{-2} + \sigma_B^{-2})} \left[\sigma_A^{-2} \sum_{j \in N_i, \sigma_j^2 = \sigma_A^2} a_{j,t-1} + \sigma_B^{-2} \sum_{j \in N_i, \sigma_j^2 = \sigma_B^2} a_{j,t-1} \right].$$

The actual variance of this estimate converges to:

$$\text{Var}(r_{i,t}^{naive} - \theta_{t-1}) = \frac{1}{(\sigma_A^{-2} + \sigma_B^{-2})} [\sigma_A^{-4} \text{Cov}_{AA}^\infty + \sigma_B^{-4} \text{Cov}_{BB}^\infty + 2\sigma_A^{-2}\sigma_B^{-2} \text{Cov}_{AB}^\infty] \quad (\text{H.2})$$

where Cov_{AA}^∞ is the covariance of two distinct agents of signal type A and Cov_{BB}^∞ and Cov_{AB}^∞ are defined similarly.

Since agents believe this variance is close to 1, the action of any agent with signal variance σ_A^2 is approximately:

$$a_{i,t} = \frac{r_{i,t}^{naive} + \sigma_A^{-2} s_{i,t}}{1 + \sigma_A^{-2}}.$$

We can then compute the limits of the covariances of two distinct agents of various signal types to be:

$$\text{Cov}_{AA}^\infty = \frac{\rho^2 \kappa_t^2 + 1}{(1 + \sigma_A^{-2})^2}; \quad \text{Cov}_{BB}^\infty = \frac{\rho^2 \kappa_t^2 + 1}{(1 + \sigma_B^{-2})^2}; \quad \text{Cov}_{AB}^\infty = \frac{\rho^2 \kappa_t^2 + 1}{(1 + \sigma_A^{-2})(1 + \sigma_B^{-2})}.$$

Plugging into H.2 we obtain

$$\kappa_t^{-2} = 1 - \frac{1}{(\sigma_A^{-2} + \sigma_B^{-2})} \left(\frac{\sigma_A^{-2}}{1 + \sigma_A^{-2}} + \frac{\sigma_B^{-2}}{1 + \sigma_B^{-2}} \right).$$

Using this formula, we can check that the limits of agent variances in $\widehat{\mathbf{V}}^{sym}$ match equations H.1.

We must check there is an equilibrium near $\widehat{\mathbf{V}}^{sym}$ with high probability. Let $\zeta = 1/n$. Let E be the event that for each agent i , the number of agents observed by i with private signal variance σ_A^2 is within a factor of $[1 - \zeta^2, 1 + \zeta^2]$ of its expected value, and similarly the number of agents observed by i with private signal variance σ_B^2 is within a factor of $[1 - \zeta^2, 1 + \zeta^2]$ of its expected value. This event implies that each agent observes a linear number of neighbors and observes approximately the same number of agents with each signal quality. We can show as

in the proof of Theorem 1 that for n sufficiently large, the event E occurs with probability at least $1 - \zeta$. We condition on E for the remainder of the proof.

Let $\mathcal{V}_\varepsilon$ be the ε -ball around in $\widehat{\mathbf{V}}^{sym}$ the sup norm. We claim that for n sufficiently large, the updating map preserves this ball: $\Phi^{naive}(\mathcal{V}_\varepsilon) \subset \mathcal{V}_\varepsilon$. We have $\Phi^{naive}(\widehat{\mathbf{V}}^{sym}) = \widehat{\mathbf{V}}^{sym}$ up to terms of $O(1/n)$. As we showed in the first paragraph of this proof, the partial derivatives of Φ^{naive} are bounded above by a constant less than one. For n large enough, these facts imply $\Phi^{naive}(\mathcal{V}_\varepsilon) \subset \mathcal{V}_\varepsilon$. We conclude there is an equilibrium in $\mathcal{V}_\varepsilon$ by the Brouwer fixed point theorem.

Finally, we compare the equilibrium variances to the ε -aggregation benchmark and to V^∞ . It is easy to see these variances are worse than the ε -aggregation benchmark for n large for some $\varepsilon > 0$, and therefore by Theorem 1 also asymptotically worse than the Bayesian case when $\sigma_A^2 \neq \sigma_B^2$.

In the case $\sigma_A^2 = \sigma_B^2$, it is sufficient to show that Bayesian agents place more weight on their private signals (since asymptotically action error comes from past changes in the state and not signal errors). Call the private signal variance σ^2 . For Bayesian agents, we showed in Theorem 1 that the weight on the private signal is equal to $\frac{\sigma^{-2}}{\sigma^{-2} + (\rho^2 \text{Cov}^\infty + 1)^{-1}}$ where Cov^∞ solves

$$\text{Cov}^\infty = \frac{(\rho^2 \text{Cov}^\infty + 1)^{-1}}{[\sigma^{-2} + (\rho^2 \text{Cov}^\infty + 1)^{-1}]^2}.$$

For naive agents, the weight on the private signal is equal to $\frac{\sigma^{-2}}{\sigma^{-2} + 1}$, which is smaller since $\text{Cov}^\infty > 0$.

APPENDIX I. SOCIALLY OPTIMAL LEARNING OUTCOMES WITH NON-DIVERSE SIGNALS (ONLINE APPENDIX)

In this section, we show that a social planner can achieve vanishing aggregation errors even when signals are non-diverse. Thus, slower rate of learning at equilibrium with non-diverse signals is a consequence of individual incentives rather than a necessary feature of the environment.

Let G_n be the complete network with n agents. Suppose that $\sigma_i^2 = \sigma^2$ for all i and $m = 1$.

Proposition 9. Let $\varepsilon > 0$. Under the assumptions in this section, for n sufficiently large there exist weights $\mathbf{W}$ and $\mathbf{w}^s$ such that at the corresponding steady state on G_n , the ε -aggregation benchmark is achieved.

Proof. An agent with a social signal equal to θ_{t-1} would place weight $\frac{\sigma^{-2}}{\sigma^{-2} + 1}$ on her private signal and weight $\frac{1}{\sigma^{-2} + 1}$ on her social signal. Let $w_A^s = \frac{\sigma^{-2}}{\sigma^{-2} + 1} + \delta$ and $w_B^s = \frac{\sigma^{-2}}{\sigma^{-2} + 1} - \delta$, where we will take $\delta > 0$ to be small.

Assume that the first $\lfloor n/2 \rfloor$ agents place weight w_A^s on their private signals and weight $1 - w_A^s$ on a common social signal r_t we will define, while the remaining agents place weight

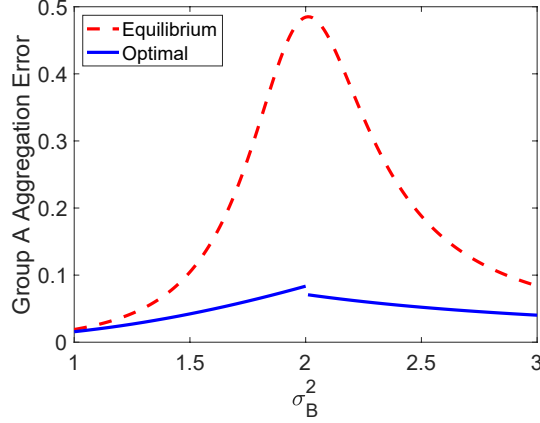


FIGURE I.1. Social planner's optimum and Bayesian learning. The red curve shows equilibrium aggregation errors on a complete graph with $n = 600$ agents, split into two equally-sized groups with private signal variances $\sigma_A^2 = 2$ and σ_B^2 varying. The blue curve plots the aggregation errors when weights are chosen by a social planner the sum of agents' steady-state action variances.

w_B^s on their private signals and weight $1 - w_B^s$ on the social signal r_t . As in the proof of Theorem 2,

$$\begin{aligned} \frac{1}{\lfloor n/2 \rfloor} \sum_{j=1}^{\lfloor n/2 \rfloor} a_{j,t-1} &= w_A^s \theta_{t-1} + (1 - w_A^s) r_{t-1} + O(n^{-1/2}), \\ \frac{1}{\lceil n/2 \rceil} \sum_{j=\lfloor n/2 \rfloor+1}^n a_{j,t-1} &= w_B^s \theta_{t-1} + (1 - w_B^s) r_{t-1} + O(n^{-1/2}). \end{aligned}$$

There is a linear combination of these summations equal to $\theta_{t-1} + O(n^{-1/2})$, and we can take r_t equal to this linear combination. Taking δ sufficiently small and then n sufficiently large, we find that ε -perfect aggregation is achieved. $\square$

In Figure I.1, we consider equilibrium and socially optimal outcomes with $n = 600$. Half of agents are in group A, with signal variance $\sigma_A^2 = 2$, while the other half are in group B, with signal variance σ_B^2 changing. In blue we plot average equilibrium aggregation errors for group A. In green we plot the average aggregation errors of group A when a social planner minimizes the total action variance (of both groups). The weights that each agent puts on her own private signal and the other agents are set to depend only on the groups. Under these socially optimal weights agents learn very well, and heterogeneity in signal variances only has a small impact.