

Experimentation in Networks[†]

By SIMON BOARD AND MORITZ MEYER-TER-VEHN*

We propose a model of strategic experimentation on social networks in which forward-looking agents learn from their own and neighbors' successes. In equilibrium, private discovery is followed by social diffusion. Social learning crowds out own experimentation, so total information decreases with network density; we determine density thresholds below which agents' asymptotic learning is perfect. By contrast, agent welfare is single peaked in network density and achieves a second-best benchmark level at intermediate levels that strike a balance between discovery and diffusion. (JEL D82, D83, D86, O31, O33, Z13)

The discovery and diffusion of innovations are key drivers of long-term economic growth. This is illustrated by the seminal papers of Griliches (1957) and Coleman, Katz, and Menzel (1957), which document the spread of new technologies by farmers and doctors. From the perspective of societal welfare, discovery and diffusion are complements; Mokyr (1992) argues that both are required for sustained economic progress. From an individual strategic perspective, they are substitutes; Grossman and Stiglitz (1980) famously point out that if prices aggregate information efficiently, then individual agents have no incentive to privately generate such information. Economic theory has made large strides in understanding information acquisition and aggregation in centralized settings such as financial markets, auctions, and collective experimentation. These incentives are less well understood in decentralized settings, where information slowly diffuses through society. This paper seeks to reconcile these forces in a parsimonious equilibrium model of experimentation in networks.

The classic paper on this topic, Bala and Goyal (1998), considers myopic non-Bayesian agents who do not reason about the behavior of their neighbors. This short-cuts strategic considerations and allows one to solve the model as a sequence of static decision problems. By contrast, our agents are forward-looking and Bayesian, so they reason about the network and future learning opportunities. To maintain

* Board: UCLA (email: sboard@econ.ucla.edu); Meyer-ter-Vehn: UCLA (email: mtv@econ.ucla.edu). Jeff Ely was the coeditor for this article. We have received useful comments from Andy Atkeson, Alessandro Bonatti, Roberto Corrao, Krishna Dasaratha, Glenn Ellison, Drew Fudenberg, Ben Golub, Yannai Gonczarowski, Johannes Hörner, Matt Jackson, Ben Jones, Emir Kamenica, Andreas Kleiner, Oleg Itskhoki, Alessandro Lizzeri, Gustavo Manso, Phil Reny, and Basil Williams. We also thank seminar audiences at ASU, University of Bonn, Cambridge-INET, CETC 2022, University of Chicago, CNSE 2022, Columbia, Collegio Carlo Alberto, Haas Marketing, Harvard/MIT, LSE, Monash, HKUST, Kellogg, LA Theory Day, Oxford, Penn, Penn State, SAET 2022, Simons Berkeley, University of Arizona, UT Austin, UCL, UCLA, VSET (Bristol/Carlos III/Essex), University of Wisconsin, Yale. We gratefully acknowledge financial support from NSF Grant 2149291.

[†]Go to <https://doi.org/10.1257/aer.20230233> to visit the article page for additional materials and author disclosure statement(s).

tractability, we impose structure on the network and assume agents learn via perfect good news events; this reduces each agent's problem to choosing a deterministic cutoff time, with social learning described by simple ordinary differential equations, opening the gate to a myriad of questions about experimentation in networks.

We use this new approach to study how asymptotic information and welfare depend on network density, as measured by either the degree in regular random networks or by the core size in core-periphery networks. For either measure, we show that agents' asymptotic information decreases monotonically in network density and they eventually learn the truth when the network is sufficiently sparse. By contrast, welfare is single peaked in network density and attains a second-best welfare benchmark when density is intermediate; such networks both encourage generation of information and quickly diffuse the discoveries. Collectively, these results paint a clear picture about learning dynamics, information aggregation, and welfare in networks of forward-looking, Bayesian agents.

In the model, I agents (Iris, John, Kata . . .) are connected by an exogenous network (e.g., clique, tree, core-periphery). They can each exert effort experimenting with a new technology whose state is high or low; effort generates successes at random times if and only if the state is high. Agents learn from their own and neighbors' successes but do not observe neighbors' effort. This simple model captures a number of applications: consider farmers learning about the success of a new crop from neighbors, doctors learning about a new drug from colleagues, or landowners learning the best way to extract shale gas from nearby frackers.

In Section II, we first characterize Iris's best response to arbitrary strategies of others. Observing a success perfectly reveals the high state, so she exerts effort forever after. Before this time, Iris's effort, or "experimentation," is based on her *social learning curve*, i.e., the expected effort of her neighbors. We show that Iris's dynamic experimentation problem is solved by a simple *cutoff strategy*: in the absence of success, Iris stops experimenting at some cutoff time τ_i . An increase in social information crowds out Iris's private experimentation, lowering her cutoff time. Unsuccessful past social learning makes Iris pessimistic, while future social information lowers the information value of her own experimentation.

Next, we illustrate how to generate Iris's social learning curve from others' cutoff times via examples. In the clique network, social learning is fast but shallow. The agents collectively experiment as much as a single agent would by herself. Adding agents speeds up learning but does not raise aggregate information because the density of the network chokes off experimentation prematurely. In the line network, social learning is deep but slow. The agents collectively experiment an infinite amount. Eventually, they learn the state perfectly, but the sparsity of the network constrains the speed of learning.

In Section III, we study the effect of network density on asymptotic information and welfare. Specifically, we consider two canonical classes of networks (regular random networks and core-periphery networks) as $I \rightarrow \infty$. To study aggregate information, define the *asymptotic information* to be the total information created by society; there is *asymptotic learning* if asymptotic information is unbounded, meaning that the agents eventually learn the state. To study welfare, we propose a second-best benchmark that upper bounds equilibrium utility (of the worst-off agent) across all networks. The clique does not attain this benchmark because the

network is too dense and agents do not generate enough information; the line does not attain it either because the network is too sparse and learning is too slow. But we show that both large random networks and large core-periphery networks do attain the benchmark when network density is intermediate.

We first study large regular random networks with degree n^I . This model encompasses sparse trees, where n^I does not depend on I , and dense cliques, where $n^I/I \rightarrow 1$. Theorem 1 completely characterizes asymptotic information and welfare as functions of network density. Asymptotic information falls in network density, and asymptotic learning obtains if density is below a threshold. Specifically, the agents fully learn if the *time-diameter* (the typical time for information to travel between two agents) exceeds a threshold σ^* . Welfare is single peaked in network density and attains the second-best benchmark if $n^I \rightarrow \infty$ and $n^I/I \rightarrow 0$. Intuitively, asymptotic learning requires sparsity to sustain experimentation incentives; high welfare requires intermediate density to ensure both generation and prompt diffusion of information.

To study the role of network position on experimentation incentives, we next turn to core-periphery networks, where K^I *core* agents connect to everyone, while $I - K^I$ *peripheral* agents connect only to all core agents. In equilibrium, core agents have more social information than peripherals, so they experiment less and have higher utility. While core agents experiment little themselves (if at all), they serve an important role as information brokers connecting the peripherals. As $I \rightarrow \infty$, asymptotic learning and welfare exhibit similar properties to large random networks, with core size substituting for the degree. Theorem 2 completely characterizes asymptotic information and welfare as functions of network density. Asymptotic information decreases in network density, and asymptotic learning obtains if K^I remains below a threshold κ^* . Welfare is single peaked in network density and attains the second-best benchmark if K^I exceeds κ^* and $K^I/I \rightarrow 0$.

Our two families of networks differ in their network structure and thus exhibit different social learning dynamics. In large random networks, independent successes are achieved over time by agents scattered throughout the network; for a typical agent, all this social information arrives in a single burst at a fixed time σ . By contrast, in core-periphery networks, independent successes are achieved in the first instant by a small fraction of peripherals; for a typical peripheral agent, this initial burst of social information arrives slowly as it is filtered through the core. The resulting cumulative social learning curves are thus convex for large random networks but concave for core-periphery networks.

Our analysis of large random networks and core-periphery networks points to a fundamental trade-off between social learning and welfare. These goals are often thought to be aligned; Hayek (1945) famously emphasizes the importance of information aggregation for allocative efficiency. However, in our model agents must be incentivized to acquire information, so the fast diffusion required for second-best welfare can lower total information. Indeed, for core-periphery networks, the two goals are mutually exclusive.

Literature.—At the core of the paper is a “perfect good news” model of strategic experimentation with unobserved actions and private payoffs. In the context of a clique, Keller, Rady, and Cripps (2005) study a good-news model with observed

actions and private payoffs; Bonatti and Hörner (2011) consider a good-news model with unobserved actions and public payoffs; and Bonatti and Hörner (2017) consider a bad-news model with unobserved actions and private payoffs. In all of these papers, equilibrium is in mixed strategies. Specifically, in the first two papers, agents gradually phase out their experimentation as the public belief approaches the exit threshold. In our model, agents use simple cutoff strategies; this allows us to go beyond the clique and solve for equilibria in rich classes of networks. We also think that the assumptions of unobserved actions and private payoffs are a natural way to model a network of farmers, doctors or frackers whose externalities are purely informational.

Observational learning in networks was pioneered by Bala and Goyal (1998), who study myopic, non-Bayesian agents and provide conditions on the network under which (i) agents reach a consensus and (ii) the agents learn the state.¹ Subsequent work has generalized these two limited results in models with forward-looking, Bayesian agents who incorporate the future value of information when choosing to experiment. Rosenberg, Solan, and Vieille (2009) consider a very general model that encompasses strategic experimentation in networks and shows that all agents eventually play the same action. Camargo (2014) considers a continuum-agent model with “random sampling” and shows that information aggregates if each action is myopically optimal for a positive measure of agents’ heterogeneous priors. By focusing on good-news learning, we can characterize learning dynamics at each point in time rather than restricting attention to long-run behavior. This is important because agents care about when innovations diffuse and not just if they diffuse; indeed, this consideration underlies the contrast between sparse networks that induce asymptotic learning and the denser networks that maximize welfare.²

Most closely related to our model, Salish (2015) embeds a discrete-time version of Keller, Rady, and Cripps’s (2005) strategic experimentation model in a network. Neighbors observe each others’ actions, which thus signal successes of second neighbors; Salish (2015) sidesteps such signaling by introducing an additional learning channel, whereby successes are automatically transmitted across the network, one link per period. More recently, Manso and Pourbabaee (2023) consider a two-period version of Keller, Rady, and Cripps (2005) on a network. In both of these papers, social learning crowds out private experimentation, but neither of them speaks to our main result of welfare being single peaked in network density.

The complexity of Bayesian updating has led some authors to consider reduced-form models of information acquisition and aggregation. For example, Bramoullé and Kranton (2007) and Galeotti and Goyal (2010) consider a local public goods game where each agent chooses a contribution level and benefits from her neighbors’ contributions. Since our agents optimally choose a deterministic stopping time, we

¹ Sadler (2020b) characterizes outcomes more completely in Bala-Goyal’s model with Brownian learning.

² A parallel literature considers dynamic learning games where private information is initially endowed to agents instead of being learned over time. Gale and Kariv (2003) show that consensus must emerge when agents are Bayesian and myopic. Mossel, Sly, and Tamuz (2015) extend this result to forward-looking agents and also show that agents eventually learn the state if the network is not too connected (e.g., the network is undirected with bounded degree). Another classic literature considers agents who move in sequence, learning from (a subset of) prior agents. Acemoglu et al. (2011) show that society learns the state if signals are unbounded and agents (indirectly) observe an unbounded number of agents. Mossel et al. (2020) unify many of the results in these literatures by looking at steady-state asymptotic behavior.

recover the tractability of the reduced-form models of experimentation in a model of Bayesian learning.

In seeking to characterize learning dynamics in networks, the paper is related to Board and Meyer-ter-Vehn (2021). In that paper, myopic agents sequentially choose to acquire information at a single point in time. Here, forward-looking agents simultaneously choose to acquire information at every point in time. The different models give rise to different economic forces. The forward-looking agents in this paper anticipate the arrival of future social information that crowds out their private experimentation, and the repeated choices give rise to the clean distinction between an experimentation phase and a contagion phase. This paper also focuses on a different question: How do aggregate information and welfare change with network density?

The paper also complements a growing empirical literature on innovation and social learning. Fetter et al. (2020) study the effect of disclosure requirements on the choice of chemical inputs in the shale gas industry. They find that improved disclosure increases interfirm social learning but decreases innovation; this is precisely the trade-off underlying our welfare results. Hodgson (2021) studies the effects of information sharing in a structural model of oil exploration and development. As in our model, forward-looking firms experiment and free ride on others' experimentation, but, as in Bala and Goyal (1998), their beliefs are not Bayesian in that they do not draw inferences about others' beliefs from their action choices or their past exploration rights. In the development literature, Foster and Rosenzweig (1995); Munshi (2004); and Bandiera and Rasul (2006) show that imperfect information is a major barrier to the adoption of new crops and that social learning crowds out farmers' experimentation. These papers focus on the agents' best responses but do not address how information diffuses across a network. Recently, this latter question was taken up by Banerjee et al. (2021) and Beaman et al. (2021) using a non-Bayesian DeGroot model of learning.³ Overall, these literatures lack a simple equilibrium framework with forward-looking Bayesian agents that can be estimated and used for counterfactuals. This paper proposes such a framework.

I. Model

Network.—Agents $\{1, \dots, I\}$ are connected by an undirected network, g , consisting of unordered pairs (i, j) (Iris, John) that represent neighbors who observe each other. The set of Iris's neighbors is $N_i(g)$. The network may be deterministic or random; denote the random network by G with realization g .

Game.—The agents seek to learn about the quality $\theta \in \{L, H\}$ of a new technology. Time is continuous, $t \in [0, \infty)$. At time $t = 0$, the common prior is $\Pr(\theta = H) = p_0$. At each time t , agent i privately chooses effort $A_{i,t} \in [0, 1]$ at flow cost c . This effort results in successes with Poisson arrival rate $A_{i,t} \mathbf{1}\{\theta = H\}$. Agent i observes her own and her neighbors' past successes but not others' actions.

³ There are a variety of other papers that study the impact of social learning on the adoption of new agricultural technology (Besley and Case 1994; Conley and Udry 2010; BenYishay and Mobarak 2019), financial innovations (Banerjee et al. 2013), health interventions (Kremer and Miguel 2007; Dupas 2014), and consumer products (Goolsbee and Klenow 2002; Moretti 2011; Bailey et al. 2022).

If the network is random, she knows G but nothing about the realization g , and G is independent of quality θ .

Payoffs.—Agents receive payoff $x > c$ from their own successes. Payoffs are discounted at rate $r > 0$, so Iris's net present value equals

$$(1) \quad V_i = \max_{\{A_{i,t}\}_{t \geq 0}} E \left[\int_0^\infty e^{-rt} A_{i,t} (x \mathbf{1}\{\theta = H\} - c) dt \right],$$

where the expectation is taken over quality θ , network G , and past observed successes on which $A_{i,t}$ conditions. We solve for weak perfect Bayesian equilibria, where agents who have observed a success infer that $\theta = H$. Since $x > c$, agents set $A_{i,t} \equiv 1$ after observing a success and obtain continuation value $y := (x - c)/r$. To avoid trivialities, we assume that the prior belief exceeds the *single-agent threshold belief*, $p_0 > \underline{p} := c/(x + y)$ (see Section IIA).

Interpretation.—The model assumes that actions are unobservable and generate observable success with a stochastic delay. Jointly, these assumptions lead to the slow diffusion of successes, much like an epidemiological SI model. They also make the analysis tractable from a strategic perspective. Agents do not attempt to trigger their neighbors to experiment and use cutoff strategies rather than the mixed strategies in Keller, Rady, and Cripps's (2005) symmetric equilibrium.

We now map the model to the introductory applications (doctors, farmers, and frackers). First, we discuss how to interpret payoffs, which, in the model, arrive at Poisson frequency when an agent exerts effort in state $\theta = H$.

- **Poisson arrival of payoffs:** Suppose doctors learn about the effectiveness of a new drug. The drug does not work with every patient, but in state $\theta = H$, it works at Poisson frequency. When using the drug, the doctor pays a flow cost c and receives benefit x at Poisson intervals, when the drug is effective.
- **Flow payoff interpretation:** Suppose farmers learn if a new crop works in their climate. While experimenting, they pay a flow cost c , representing the opportunity cost of land. Experimentation takes time, as they must try different inputs (e.g., watering patterns, fertilizer). When the new crop succeeds, they use it thereafter and receive flow payoff $(1 + r)x - c$.
- **Lump-sum payoff interpretation:** Suppose frackers learn whether they can extract natural gas from the ground. They pay flow cost c when experimenting, representing the cost of trying different chemical mixtures. If their exploration succeeds, they receive lump-sum payoff $x + y$, representing the value of the gas.

Under the second and third interpretation, agents only observe any neighbor succeed once; this does not matter since observing one success reveals $\theta = H$ perfectly.

Second, we assume that neighbors' successes are observable. In the applications, this may be because the neighbors see the success directly (e.g., a farmer brings their new crop to market), the neighbors see the agent's payoff from succeeding (e.g., the farmer buys a new truck), or the neighbors see a piece of hard information that someone shares (e.g., the local Bayer representative tells other farmers about

the crop's success). The method of communication affects the identity of neighbors and thus the density of the resulting network. For example, relatively few people will see the farmer's new truck (leading to a sparse network), whereas the Bayer representative may tell all their clients (leading to a dense network).

Third, we assume that neighbors' actions are unobservable. In the applications, this means a farmer does not know whether other farmers are planting the new crop, a doctor does not know whether colleagues are prescribing the drug, and a fracker does not know whether other landowners are actively exploring. We also assume agents do not directly communicate with one another (other than successes). Indeed, an agent has little incentive to reveal her failures, which tend to make her neighbors more pessimistic and lower their experimentation. These assumptions ensure that news spreads slowly over the network.

II. General Analysis

We start with a general analysis of best responses. Section IIA shows agents use cutoff strategies, while Section IIB derives comparative statics.⁴ Section IIC then characterizes equilibrium in three examples, while Section IID establishes general equilibrium existence and discusses equilibrium uniqueness.

A. Best Responses: Cutoff Strategies

In this section, we characterize the best response of a generic agent, Iris, given arbitrary strategies of other agents.

As a benchmark, consider the single-agent experimentation problem, or equivalently, Iris's problem when she has no neighbors. After her first success, she sets $A_{i,t} \equiv 1$. Before that, as long as $A_{i,t} = 1$, her posterior belief evolves according to

$$p_t = P^\theta(t) := \frac{p_0 e^{-t}}{p_0 e^{-t} + 1 - p_0}.$$

Iris thus experiments until time $\bar{\tau}$, when her belief hits the single-agent threshold belief $p_{\bar{\tau}} = \underline{p} = c/(x + y)$. It is also useful to define the *myopic threshold belief* $\bar{p} := c/x$, where Iris would stop if she ignored the future benefit of success, y .

Returning to the general problem where Iris learns from her neighbors $N_i(G)$, write T_i for Iris's first success time and $S_i := \min_{j \in N_i(G)} T_j$ for her neighbors' first success time. After Iris observes a success at $\min\{T_i, S_i\}$, she chooses maximal effort and receives continuation value y . We can thus restrict attention to earlier times and write $\{a_{i,t}^\theta\}_{t \geq 0}$ for her *experimentation*, i.e., her effort before $\min\{T_i, S_i\}$. Define Iris's rate of social learning by

$$(2) \quad b_{i,t} := E^{-i} \left[\sum_{j \in N_i(G)} A_{j,t} \middle| t < S_i \right],$$

⁴The results in Sections IIA and IIB extend far beyond the networks of Section I, for instance, to directed or time-varying networks, or to agents with private information about the network, specifically about their own degree $|N_i(g)|$.

where the expectation E^{-i} is taken over the random network G and others' success times $\{T_j\}_{j \neq i}$, conditional on $\theta = H$ and assuming that no one has observed a success by i . We also define the cumulative of Iris's experimentation $\alpha_{i,t}^\theta := \int_0^t a_{i,s}^\theta ds$ and social learning $\beta_{i,t} := \int_0^t b_{i,s} ds$. We generally refer to both $\{b_{i,t}\}$ and $\{\beta_{i,t}\}$ as Iris's *social learning curve*; when referring specifically to $\beta_{i,t}$, we call it the *cumulative social learning curve*. Since Iris's experimentation is unobservable to others and her own success effectively ends the game for her, Iris takes $\{b_{i,t}\}$ as given. We thus study the best response $\{a_{i,t}^\theta\}$ to $\{b_{i,t}\}$ and drop the i subscript for the rest of the section.

When $\theta = H$, the random time $\min\{T, S\}$ has hazard rate $a_t^\theta + b_t$, implying chance $\exp\{-(\alpha_t^\theta + \beta_t)\}$ of observing no success before t , and a posterior belief equal to

$$p_t = P^\theta(\alpha_t^\theta + \beta_t).$$

Truncating (1) at $\min\{T, S\}$ renders Iris's stochastic control problem deterministic,

$$(3) \quad V = \max_{\{a_t^\theta\}_{t \geq 0}} \int_0^\infty e^{-rt} \left[p_0 e^{-(\alpha_t^\theta + \beta_t)} + (1 - p_0) \right] \left([a_t^\theta(x + y) + b_t y] p_t - a_t^\theta c \right) dt.$$

Intuitively, Iris gets $x + y$ when she succeeds, y when a neighbor succeeds, and effort costs c . The chance of no success by time t is $e^{-(\alpha_t^\theta + \beta_t)}$ when $\theta = H$ and one when $\theta = L$.

Clearly, Iris experiments for beliefs above the myopic threshold, $p_t \geq \bar{p}$. Conversely, equation (3) implies that Iris stops experimenting below the single-agent threshold, $p_t \leq \underline{p}$. For beliefs $p_t \in [\underline{p}, \bar{p}]$, her choice depends on her social learning. We say the prior is *optimistic* if $p_0 > \bar{p}$ and *pessimistic* if $p_0 < \bar{p}$.⁵ An optimistic agent always engages in some experimentation, no matter her social learning curve.

We first claim that Iris uses a *cutoff strategy* in that she experiments maximally until some cutoff time τ and then stops, $a_t^\theta = \mathbf{1}\{t \leq \tau\}$.⁶ Intuitively, it is sub-optimal to stop experimenting at some τ' but then resume it after neighbors' lack of success over $[\tau', \tau'']$. To see why, suppose Iris shirks at time t but works at time $t + \delta$, and consider the effect of front-loading effort ϵ from $t + \delta$ to t . This has two consequences. First, if the effort pays off, i now gets to enjoy the success earlier, raising her value by $r\delta[p_t(x + y) - c]\epsilon$, which is positive in the relevant range of posteriors, $p_t > \underline{p}$. Second, if one of her neighbors succeeds over $[t, t + \delta]$, she ends up working at both t and $t + \delta$, raising her value by $p_t b_t \delta \epsilon (x - c) > 0$. Thus, Iris prefers to front-load experimentation, so she optimally uses a cutoff time τ with cutoff belief $p_\tau \in [\underline{p}, \bar{p}]$, illustrated in Figure 1.

⁵Recall our assumption $p_0 > \underline{p}$; we are pragmatic about calling the boundary case $p_0 = \bar{p}$ optimistic or pessimistic.

⁶Of course, "stopping" is provisional in the sense that Iris starts to work again when she observes one of her neighbors succeed at some $t > \tau$.

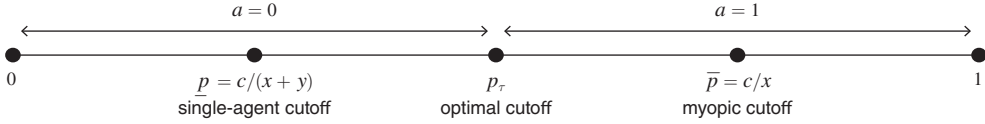


FIGURE 1. BELIEFS

Note: The agent always experiments for posterior beliefs p_t above the myopic cutoff $\bar{p}$ and never below the single-agent cutoff $\underline{p}$.

To characterize the optimal cutoff τ , define Iris's *experimentation incentives* at time t ,

$$(4) \quad \psi_t := p_t \left(x + ry \int_t^\infty e^{-\int_t^s (r+b_u) du} ds \right) - c.$$

To understand (4), suppose that successes from Iris's neighbors arrive at constant rate b , so (4) simplifies to $p_t \left(x + \frac{r}{r+b} y \right) - c$. If she raises the cutoff from t to $t + \delta$, she gains the expected payoff from a success $p_t(x+y)\delta$, forgoes the expected benefit of future social learning $p_t \left(\frac{b}{r+b} y \right) \delta$, and incurs cost $c\delta$. The experimentation incentives are the sum of these three effects. We summarize this discussion as follows.

PROPOSITION 1: *Given social information $\{b_t\}$, the agent's optimal experimentation is given by the cutoff strategy $a_t^\emptyset = \mathbf{1}\{t \leq \tau\}$, where the cutoff time $\tau \in (0, \bar{\tau}]$ uniquely solves $\psi_\tau = 0$ if $\psi_0 > 0$, and $\tau = 0$ if $\psi_0 \leq 0$.*

PROOF:

The proof in Appendix A.A formalizes the front-loading argument and shows that the marginal payoff from experimentation at the cutoff is proportional to ψ_t , which in turn single-crosses from above in t . ■

Proposition 1 reduces the potentially complicated experimentation problem of a forward-looking, Bayesian agent to choosing one number, τ , which is characterized by setting (4) to zero. This tractability allows us to characterize equilibria for rich classes of networks. In contrast to Proposition 1, the seminal papers on strategic experimentation in the clique network, Keller, Rady, and Cripps (2005) and Bonatti and Hörner (2011), both have agents gradually phase out effort in equilibrium. This difference arises because free riding incentives are greater in their models. In Keller, Rady, and Cripps (2005), actions are observable, so Iris's neighbors get pessimistic when her experimentation fails; in Bonatti and Hörner (2011), payoffs are public, so Iris does not want to exert effort if others are about to succeed.

B. Best Responses: Comparative Statics

This section derives two useful comparative statics on Iris's value and her optimal cutoff as a function of social learning.

LEMMA 1: *Higher cumulative social learning $\{\beta_t\}_{t \geq 0}$ raises value V and lowers the optimal cutoff τ .*

PROOF:

Higher $\{\beta_t\}$ constitutes Blackwell-better information and raises V . Experimentation incentives (4) fall both in pre-cutoff learning β_τ , which lowers the cutoff belief $p_\tau = P^\theta(\tau + \beta_\tau)$, and in future learning $\{b_t\}_{t \geq \tau}$. To show that ψ_τ falls in cumulative learning $\{\beta_t\}$, we need to compare the impact of “early” and “late” increases in b_t . Specifically, differentiating time- τ experimentation incentives (4) with respect to time- t social learning, we get⁷

$$(5) \quad -\frac{\partial \psi_\tau}{\partial b_t} = \begin{cases} p_\tau \left(ry \int_\tau^\infty e^{-\int_\tau^s (r+b_u) du} ds + x - c \right) & \text{for } t < \tau \\ p_\tau ry \int_t^\infty e^{-\int_\tau^s (r+b_u) du} ds & \text{for } t > \tau, \end{cases}$$

where the case $t < \tau$ uses $\partial p_\tau / \partial b_t = -p_\tau(1 - p_\tau)$ and $(1 - p_\tau)(x + ry \int_\tau^\infty e^{-\int_\tau^s (r+b_u) du} ds) = x + ry \int_\tau^\infty e^{-\int_\tau^s (r+b_u) du} ds - (\psi_\tau + c)$. Clearly, (5) is positive and falls in t , weakly for $t < \tau$ and discontinuously at $t = \tau$. Thus, earlier learning reduces incentives more, so ψ_τ falls as a function of $\{\beta_t\}$. Since ψ_t strictly single-crosses from above (by the Proof of Proposition 1), the solution τ of $\psi_\tau = 0$ falls in $\{\beta_t\}$. ■

Equation (5) tells us that pre-cutoff learning β_τ crowds out the agent’s experimentation more than post-cutoff learning $\{b_t\}_{t \geq \tau}$. After the cutoff, it crowds out the option value of own experimentation $ry \int_\tau^\infty e^{-\int_\tau^s (r+b_u) du} ds$, as seen in the second line of (5) for $t = \tau$. Before the cutoff, the additional term $x - c$ in the first line of (5) represents the reduced opportunity of achieving a first success at τ .⁸

In online Appendix C.1, we show that Iris’s learning curve $\{\beta_t\}$ rises in other agents’ cutoffs. Together with Lemma 1, this means that τ_i falls in τ_{-i} , so the game has strategic substitutes.⁹

Our next result provides a tool for rich comparisons of equilibrium values. Lemma 1 is of limited value for such comparative statics because the order on social learning $\{\beta_t\}$ is highly incomplete. To obtain a sharper tool, the Proof of Lemma 2 shows that by truncating the integral expression for an agent’s value (3) at τ , we can write the agent’s value as a function of only two variables: the cutoff τ and pre-cutoff social learning β_τ ,

$$(6) \quad V = \frac{p_0 x - c}{r} + e^{-r\tau} \left(p_0 e^{-\beta_\tau - \tau} (x - c) - (1 - p_0) c \frac{r - 1}{r} \right) =: \mathcal{V}(\tau, \beta_\tau).$$

⁷Formally, define $\frac{\partial \psi_\tau}{\partial b_t} = \lim_{\epsilon \rightarrow 0} \frac{1}{\epsilon} [\psi_\tau(\{b_s^{\epsilon}\}_{s \geq 0}) - \psi_\tau(\{b_s\}_{s \geq 0})]$, where $b_s^{\epsilon} := b_s + \mathbf{1}\{s \in [t - \epsilon, t]\}$.

⁸For optimistic agents, $p_0 > \bar{p}$, this asymmetry is stark. A finite amount $\beta_t = \bar{\tau}$ of pre-cutoff learning fully crowds out incentives by inducing $p_t < \bar{p}$, and so $\psi_t < 0$. In contrast, no amount of post-cutoff learning fully crowds out incentives since $\psi_0 > p_0 x - c > 0$ for any $\{\beta_t\}$.

⁹Strategic substitutes owe to our assumption of perfect good news learning. Duffie, Malamud, and Manso (2014) show the possibility of strategic complementarity in a game in which agents acquire imperfect signals and then engage in dynamic bilateral trade in randomly matched pairs.

Before τ , Iris exerts effort anyway, so she does not care about the timing of social learning $\{b_t\}_{t \leq \tau}$. After τ , learning $\{b_t\}_{t \geq \tau}$ matters only via the continuation value $V_\tau = p_\tau y \int_\tau^\infty b_s e^{-\int_\tau^s (r+b_u) du} ds = p_\tau(x+y) - c$,¹⁰ which is a function of (τ, β_τ) since $p_\tau = P^\emptyset(\tau + \beta_\tau)$.

LEMMA 2: *For any social learning curve $\{\beta_t\}$ with $\psi_0 \geq 0$ and optimal cutoff τ , the agent's value is given by (6). The function $\mathcal{V}(\tau, \beta_\tau)$ falls in both arguments, with $\partial_\tau \mathcal{V} < \partial_\beta \mathcal{V} < 0$.*

PROOF:

See Appendix A.B.

The fact that $\mathcal{V}(\tau, \beta_\tau)$ falls in β_τ may sound counterintuitive. It occurs because we fix the optimal stopping time τ , as characterized by $\psi_\tau = 0$. A rise in pre-cutoff learning β_τ must be compensated by a fall in post-cutoff learning $\{b_t\}_{t \geq \tau}$ in order to keep τ constant. More strongly, since pre-cutoff learning has a discontinuously larger effect on ψ_τ than post-cutoff learning by (5), we must reduce the latter by a larger amount to compensate. In contrast to (5), the effect of social learning on value, $\partial V / \partial b_t$, is continuous in t , so the combination of a small raise of $b_{\tau-\epsilon}$ and a large drop of $b_{\tau+\epsilon}$ decreases value.

Lemma 2 assumes $\psi_0 \geq 0$, so the agent's stopping time is characterized by $\psi_\tau = 0$.¹¹ This assumption is satisfied in the random networks in Section IIIB where all agents exert some effort, and for peripheral agents in the core-periphery networks in Section IIIC.

Lemma 2 is key to compare equilibrium welfare across agents and networks since τ and β_τ can be characterized in many cases. For example, suppose one agent optimally shirks $\tau = 0$, while another optimally works $\tau' > 0$. Since $\beta_\tau = 0$ and $\beta_{\tau'} \geq 0$, the shirker has higher utility than the worker, $\mathcal{V}(0,0) > \mathcal{V}(\tau', \beta_{\tau'})$.

C. Equilibrium: Examples

So far, we studied Iris's best response τ_i as a function of reduced-form social learning curves $\{\beta_{i,t}\}$. To close the model in equilibrium, we must study how individual cutoffs $\tau_{-i} = \{\tau_j\}_{j \neq i}$ aggregate into $\{\beta_{i,t}\}$, as illustrated in Figure 2. Here, we demonstrate this aggregation in three canonical example networks, foreshadowing the more general analysis in Section III.

Example 1 (Clique): Assume that all I agents observe each other. We claim there is a unique equilibrium in which all agents equally divide the single-agent experimentation between them; that is, $\tau_i = \bar{\tau}/I$ for all agents i , where $\bar{\tau}$ solves $P^\emptyset(\bar{\tau}) = \underline{p}$. The resulting social learning curve is illustrated in Figure 3 (panel A).

¹⁰The first equality leverages the fact that all learning after τ is social $\{b_t\}_{t \geq \tau}$, and the second leverages the indifference condition $\psi_\tau = 0$.

¹¹Otherwise, if the agent receives too much social information and $\psi_0 < 0$, her value equals $V = V_0 = p_0 y \int_0^\infty b_s e^{-(rs+\beta_s)} ds = p_0(x+y) - c - \psi_0 = \mathcal{V}(0,0) - \psi_0 > \mathcal{V}(0,0)$.

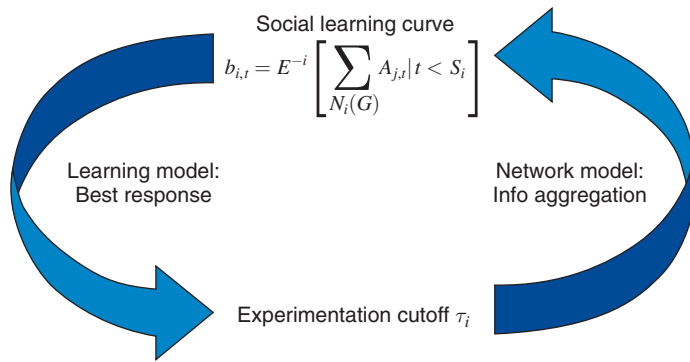


FIGURE 2. EQUILIBRIUM ANALYSIS

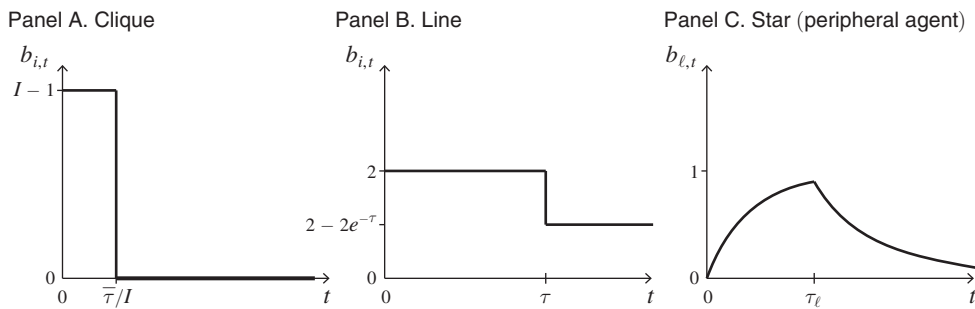


FIGURE 3. SOCIAL LEARNING CURVES

Note: This picture illustrates the rate of social learning b_{it} as defined in equation (2) for Examples 1–3, as described in the text.

As I rises, aggregate information is constant, while welfare rises as learning accelerates and agents share the cost of experimentation.

We prove our claim in two steps. First, the agents collectively experiment as much as one isolated agent, $\sum_i \tau_i = \bar{\tau}$. This is because any agent who experiments the longest expects no social information after her cutoff, $b_{i,s} = 0$ for $s > \tau_i$. Hence, she faces the first-order condition of the single-agent problem, $P^\theta(\sum_j \tau_j)(x + y) - c = 0 = \underline{p}(x + y) - c$. Second, the agents split total experimentation evenly, $\tau_j = \bar{\tau}/I$. This is because all agents are indifferent when p_t reaches $\underline{p}$ and prefer to front-load experimentation, so they all experiment until $\bar{\tau}/I$.¹²

¹²The uniqueness of equilibrium is notable since public good problems with linear costs feature a continuum of equilibria. Bramoullé and Kranton (2007) select an equilibrium via a stability criterion, while Galeotti and Goyal (2010) select via a network formation game; we resolve this indeterminacy through impatience. In experimentation papers there are also asymmetric equilibria (e.g., Keller, Rady, and Cripps 2005; Bonatti and Hörner 2011). As discussed after Proposition 1, free-riding incentives are weaker in our paper, leading putative asymmetric equilibria to unravel.

Example 2 (Line): Consider the following (infinite) network:¹³

$$\dots - j' - i - j - k - \dots$$

In the unique symmetric equilibrium, private discovery in an initial *experimentation phase* of length τ is followed by social diffusion in a *contagion phase*. For example, suppose Kata succeeds in the experimentation phase, while Iris and John and John' do not. After τ , Kata's success means that Kata and John continue to work, while Iris shirks. Eventually, John also succeeds, and Iris resumes work.

Let a_t be i 's expectation of j 's effort conditional on i not seeing a success,

$$(7) \quad a_t := E^{-i}[A_{j,t} | t < S_i] = 1 - \Pr^{-i}(t < T_k | t < T_j) \mathbf{1}\{\tau < t\}.$$

The second equality uses that, in the absence of successes by i and j , John works at times after τ if and only if Kata has succeeded. Further,

$$(8) \quad \Pr^{-i}(t < T_k | t < T_j) = \frac{\Pr^{-i}(t < T_j, T_k)}{\Pr^{-i}(t < T_j)} = \frac{e^{-\tau - \alpha_t}}{e^{-\alpha_t}} = e^{-\tau}.$$

The denominator is the chance that John's cumulative experimentation $\alpha_t = \int_0^t a_s ds$ fails to yield a success. For the numerator, the hazard rate of the success time $\min\{T_j, T_k\}$ equals two in the experimentation phase $t \leq \tau$; in the contagion phase $t > \tau$, the lack of success by i, j, k implies $A_{j,t} = 0$, so the hazard rate drops to $E^{-i}[A_{k,t} | t < T_j, T_k] = E^{-j}[A_{k,t} | t < T_k] = a_t$.

Substituting (8) into (7) yields $a_t = 1 - e^{-\tau} \mathbf{1}\{\tau < t\}$, which is constant at $t > \tau$. While the unconditional probability that Kata has succeeded, and hence, John works, rises over time, this positive effect is exactly offset by conditioning on the bad news event that John has not succeeded yet, $t < T_j$.

Since Iris has two neighbors, her social learning curve equals $b_{i,t} \equiv 2(1 - e^{-\tau} \mathbf{1}\{\tau < t\})$, as illustrated in Figure 3 (panel B). Using (4), the equilibrium stopping time τ solves

$$(9) \quad \psi_\tau = P^\emptyset(3\tau) \left(x + \frac{r}{r + 2(1 - e^{-\tau})} y \right) - c = 0.$$

Example 2' (Tree): Generalizing Example 2, we consider an (infinite) *tree* where everyone has n neighbors. Iris's expectation of neighbor John's effort in (7) now considers the event that none of his $n - 1$ other neighbors k has succeeded. In turn, $1 - a_t$ in (8) becomes

$$(10) \quad \Pr^{-i}(t < T_k, \forall k \in N_j | t < T_j) = \frac{\Pr^{-i}(t < T_j, T_k, \forall k \in N_j)}{\Pr^{-i}(t < T_j)} = \frac{e^{-\tau - (n-1)\alpha_t}}{e^{-\alpha_t}}.$$

¹³This example has infinite agents, but we can approximate it with a sequence of finite random networks that generate circles of exploding size. These finite networks admit unique, symmetric equilibria that converge to the symmetric equilibrium described here (see Board and Meyer-ter-Vehn 2024).

The denominator is the same as (8). For the numerator, the hazard rate of the success time $\min\{T_j, T_k, T_{k'}, \dots\}$ equals n in the experimentation phase $t \leq \tau$; in the contagion phase $t > \tau$, the lack of success by $k, k', \dots$ implies $A_{j,t} = 0$, so the hazard rate drops to a_t for each of j 's neighbors k other than i .

Simplifying and differentiating (10), Iris's belief at $t \geq \tau$ follows the ODE

$$(11) \quad \dot{a} = (n-2)a(1-a)$$

with initial condition $a_\tau = 1 - e^{-(n-1)\tau}$ given by the probability that one of John's $n-1$ other neighbors succeeded in the experimentation phase.¹⁴

For the line, $n = 2$, we recover the time-invariant beliefs $a_t \equiv 1 - e^{-\tau}$ for $t \geq \tau$. For $n \geq 3$, Iris's belief rises over time because the good news from John's expected inflow of information outweighs the bad news from his observed lack of success. The net effect is captured by the factor $(n-2)$ in (11): the more neighbors John has, the faster he observes a success, and the faster Iris's rate of social learning increases.

Example 3 (Star): The star consists of one core agent (Kata, k) and L peripheral agents (Lili, ℓ) and undirected links between k and each ℓ . In any equilibrium, peripherals use a common cutoff τ_ℓ (by Proposition 3). Moreover, Kata learns faster than the peripherals and so experiments less herself, $\tau_k < \tau_\ell$ (by Lemma 5). Indeed, for $p_0 < \bar{p}$ and large L , Kata does not experiment herself, $\tau_k = 0$.¹⁵

When $\tau_k = 0$, information is generated by peripherals but flows via Kata, who serves as the information broker. If a peripheral succeeds before τ_ℓ , Kata sees this and starts to work; her eventual success then triggers all other peripherals to work. The resulting social learning curve for peripheral agent Lili $b_{\ell,t}$ undergoes two phases, illustrated in Figure 3 (panel C). Up to time τ_ℓ , it increases because of other peripherals' experimentation. After τ_ℓ , no more additional information is created, and $b_{\ell,t}$ falls as the information filters through Kata and Lili becomes pessimistic about Kata having seen a success. These dynamics are analogous to water that flows into a reservoir while the peripherals experiment and slowly drains out through a bottleneck as Kata conveys the information.

D. Equilibrium: Existence and Uniqueness

We round off our preliminary analysis by establishing equilibrium existence and limited uniqueness results.

PROPOSITION 2: *Equilibrium exists.*

¹⁴We can rewrite (11) as $\frac{d}{dt} \log \frac{a_t}{1-a_t} = n-2$ and solve in closed form for $a_t = 1 / (1 + \exp\{-(n-2)(t+\gamma)\})$, with constant γ determined by the initial condition $a_\tau = 1 - e^{-(n-1)\tau}$.

¹⁵This is analogous to the equilibrium where peripheral agents provide all of the public good in Bramoullé and Kranton's (2007) static, reduced-form model of experimentation (right panel of their Figure 1b).

PROOF:

We argue that the best-response mapping in cutoff vectors $\{\tau_j^*\} : [0, \bar{\tau}]^I \rightarrow [0, \bar{\tau}]^I$ is continuous, which implies equilibrium existence by Brouwer's fixed point theorem. First note that i 's social learning curve $\{b_{i,t}\}_{t \geq 0}$, defined in equation (2), is pointwise continuous in $\{\tau_j\}_{j \neq i}$ for all $t \neq \tau_j$. Then, Lebesgue's dominated convergence theorem implies that incentives $\psi_{i,t}$ in (4) are also continuous in $\{\tau_j\}_{j \neq i}$ for all t . Finally, since $\psi_{i,t}$ strictly single-crosses in t (see the Proof of Proposition 1), its root $\tau_i^*(\{\tau_j\}_{j \neq i})$ is also continuous. ■

Uniqueness is more difficult. We are not aware of networks with multiple equilibria but can only prove uniqueness under strong assumptions. For a deterministic network g and agents $i \neq j$, define $g_{i \leftrightarrow j}$ to be the same network when switching i and j .¹⁶ For a random network G , define $G_{i \leftrightarrow j}$ by $\Pr(G_{i \leftrightarrow j} = g) = \Pr(G = g_{i \leftrightarrow j})$ for all g . In analogy to sequences of random variables, we say that i and j are *exchangeable* in G if and only if $G_{i \leftrightarrow j} = G$; network G is exchangeable if $G_{i \leftrightarrow j} = G$ for any pair of agents i, j .

PROPOSITION 3: *If i and j are exchangeable in G , then in any equilibrium $\tau_i = \tau_j$.*

PROOF:

See online Appendix C.2.

For intuition, consider a deterministic network where i and j are not connected. By contradiction, assume $\tau_j < \tau_i$. Since i 's additional learning over $[\tau_j, \tau_i]$ is more immediate to i than to j , who only benefits indirectly via some other agent k , we can argue that $\min\{T_i, S_i\}$ is smaller than $\min\{T_j, S_j\}$. This greater chance of learning the state depresses i 's experimentation incentives below j 's, leading to the contradiction that $\tau_j > \tau_i$.

COROLLARY 1: *Exchangeable networks have a unique equilibrium, characterized by a cutoff $\tau \in (0, \bar{\tau}]$, such that $\tau_i = \tau$ for all i .*

PROOF:

Proposition 3 implies that all agents must share the same cutoff τ . Uniqueness follows from strategic substitutes: when agents $j \neq i$ raise their cutoff τ_{-i} , i 's social learning rises (see online Appendix C.1), which in turn lowers own experimentation τ_i by Lemma 1. ■

Exchangeability is so demanding that only two deterministic networks satisfy it, the clique and the empty network. If $j \in N_i(g)$, exchangeability implies $k \in N_i(g)$ and $j \in N_k(g)$ for all $k \neq i, j$, so g is the clique. A weaker notion of symmetry is vertex-transitivity: for each i, j , there is a graph automorphism of g that

¹⁶Formally, given g , we can define $g_{i \leftrightarrow j}$ by three types of links. First, links involving i and j : $(i, j) \in g_{i \leftrightarrow j}$ if and only if $(i, j) \in g$. Second, links involving one third party: $(i, k) \in g_{i \leftrightarrow j}$ if and only if $(j, k) \in g$, and $(j, k) \in g_{i \leftrightarrow j}$ if and only if $(i, k) \in g$. Third, links between third parties: $(k, \ell) \in g_{i \leftrightarrow j}$ if and only if $(k, \ell) \in g$.

maps i to j . By the Proof of Corollary 1, vertex-transitive networks have exactly one symmetric equilibrium, but we do not know whether there are additional, asymmetric equilibria.¹⁷

With this said, many natural classes of random networks, such as those studied in Section IIIB, are exchangeable, and Corollary 1 applies.¹⁸ Moreover, Proposition 3 is useful beyond exchangeable networks; for instance, equilibria in core-periphery networks in Section IIIC are characterized by one cutoff τ_k for all core agents and another cutoff τ_ℓ for all peripherals.

III. Density of Links

We now turn to the main question of the paper: How do learning and welfare depend on network density? Section IIIA introduces some terminology and a second-best benchmark for welfare. Section IIIB studies large random networks. Section IIIC studies large core-periphery networks. The results for asymptotic learning and welfare in these two sections run parallel to one another, but the learning dynamics differ.

A. Bounds on Learning and Welfare

First consider aggregate information. We study large networks via sequences of networks $\{G^I\}_{I \in \mathbb{N}}$ and write $\beta^I := \min_j \beta_{j,\infty}^I$ for the social information of the least informed agent. If G^I admits multiple equilibria, we consider the infimum values of β^I . We define *asymptotic information* as $\beta = \liminf_{I \rightarrow \infty} \beta^I$. There is *asymptotic learning* if $\beta = \infty$, so all agents eventually learn the state.

Next consider welfare. Iris's value is trivially bounded above by the value of learning the state perfectly immediately, $V_i < p_0 y$. Another, less obvious, upper bound on agents' value comes from the fact that for i to socially learn, some other agent $j \neq i$ must generate that social information. By Lemma 2 and equation (6), this implies that $\min_j V_j < \mathcal{V}(0,0) = p_0(x+y) - c$.¹⁹ This motivates an upper bound on Rawlsian welfare that we refer to as the *welfare benchmark*,

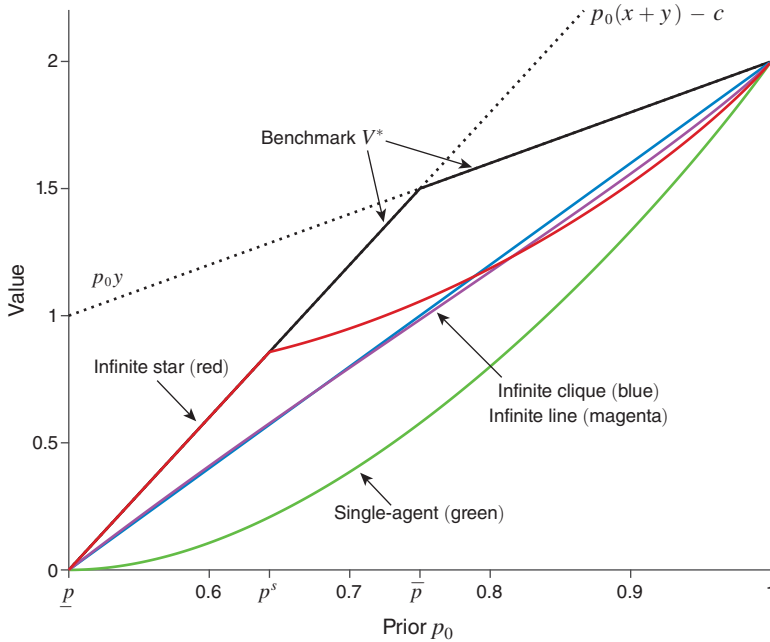
$$V^* := \min\{p_0 y, p_0(x+y) - c\},$$

illustrated in Figure 4 as a function of p_0 . Given a sequence of networks $\{G^I\}_{I \in \mathbb{N}}$, let $V^I := \min_j V_j^I$ be the expected welfare of the worst-off agent. If G^I admits multiple equilibria, we consider the infimum values of V^I . Define *asymptotic welfare* as $V = \liminf_{I \rightarrow \infty} V^I$.

¹⁷ A case in point is the four-agent ring $\dots \leftrightarrow i \leftrightarrow j \leftrightarrow k \leftrightarrow \ell \leftrightarrow i \leftrightarrow \dots$; it is vertex transitive but not exchangeable since i (but not j) is linked to ℓ , and so $g_i \neq g_{i \leftrightarrow j}$. However, i, k are exchangeable, as are j, ℓ ; any equilibrium must thus have $\tau_i = \tau_k$ and $\tau_j = \tau_\ell$. The two equilibrium conditions for these two cutoffs are easy to derive, but we do not know how to prove that $\tau_i = \tau_j$.

¹⁸ Another natural class of exchangeable networks fixes an arbitrary deterministic or random graph and then assigns agents to nodes at random. Agents thus know the global structure of the network but not their own location. We do not know if this class of networks captures all exchangeable networks.

¹⁹ This upper bound relies on agents using equilibrium strategies. Consider a sequence of clique networks in which agents use symmetric cutoffs τ^I that vanish individually, $\tau^I \rightarrow 0$, but explode in aggregate, $I\tau^I \rightarrow \infty$. Agents' payoffs approach $p_0 y$, which exceeds $p_0(x+y) - c$ for pessimistic priors $p_0 < \bar{p}$. However, in equilibrium, such agents have a strict incentive to shirk for large I (see Example 1).

FIGURE 4. WELFARE AS A FUNCTION OF THE PRIOR p_0

Notes: The figure shows the benchmark upper bound V^* and the single-agent lower bound. The benchmark V^* is piecewise linear with a downward kink at the myopic cutoff, $\bar{p}$. The figure also shows welfare in our three examples. The blue line shows the infinite clique (Example 1) in which welfare is $\mathcal{V}(0, \bar{\tau})$, where $\bar{\tau}$ is the single-agent experimentation. The magenta line shows the infinite line (Example 2) in which welfare is $\mathcal{V}(\tau, 2\tau)$, where τ satisfies (9). The red line shows the infinite star (Example 3); in which welfare equals the benchmark $\mathcal{V}(0, 0)$ for $p \leq p^s$ and otherwise $\mathcal{V}(\tau, \tau)$, where τ satisfies (12). The figure assumes the benefit and cost of experimentation are $x = 4$ and $c = 3$ and the interest rate is $r = 1/2$.

While asymptotic learning and the welfare benchmark are both driven by social learning $\{\beta_t\}$, asymptotic learning focuses on the long run, while welfare incorporates discounting. For optimistic priors $p_0 \geq \bar{p}$, the welfare benchmark requires agents learn the state immediately, so clearly they also learn asymptotically. For pessimistic priors $p_0 < \bar{p}$, asymptotic learning and the welfare benchmark are opposing goals. Recall that the value function $\mathcal{V}(\tau, \beta)$ from Lemma 2 falls in τ . Thus, the welfare benchmark $V^* = \mathcal{V}(0, 0)$ requires individual experimentation to vanish $\max_{1 \leq i \leq I} \tau_i^I \rightarrow 0$, while asymptotic learning requires aggregate information to diverge, $\sum_{i=1}^I \tau_i^I \rightarrow \infty$. For core-periphery networks, we will see that these two conditions are mutually exclusive.

Our main results, Theorems 1 and 2, show that sequences of random networks and core-periphery networks can attain asymptotic learning $\beta = \infty$ when network density is small, and the welfare benchmark V^* when network density is intermediate.²⁰

We illustrate these concepts with the three examples from Section IIC.

²⁰In both cases, the proof of the theorem characterizes unique limit points for $\{\beta^I\}$ and $\{V^I\}$, so the $\liminf$ equals the ordinary limit. In the large random networks, equilibria are unique, so taking the infimum over equilibrium values is moot. In finite core-periphery networks, we do not know whether equilibrium is unique, but the unique characterization of the limit points does not rely on taking the infimum over equilibria and rather applies for any equilibrium selection.

Example 1, continued (Clique): In a finite network, the agents share the single-agent experimentation, $\tau_i = \bar{\tau}/I$. As $I \rightarrow \infty$, individual experimentation vanishes, and all learning is social with asymptotic information $\beta = \bar{\tau}$. Agents receive all their social information before stopping, $\beta_\tau = \bar{\tau}$, so asymptotic welfare equals $\mathcal{V}(0, \bar{\tau})$. More concretely, agents' beliefs instantly jump to 1 (if there is a success) or drop to $\underline{p}$ (if there is no success). The payoff to the former is y , so the equilibrium values converge to $\mathcal{V}(0, \bar{\tau}) = (p_0 - \underline{p})y / (1 - \underline{p}) < V^*$, as illustrated in Figure 4. The speed of diffusion in the clique chokes off discovery and means that agents neither asymptotically learn nor obtain the welfare benchmark.

Example 2, continued (Line): In this infinite network, each agent experiments for time $\tau > 0$, where τ solves (9). Asymptotic learning obtains since social learning $\beta_t = 2[\tau + (1 - e^{-\tau})(t - \tau)]$ is unbounded. However, agents learn too slowly, and they do not attain the welfare benchmark. Specifically, each agent experiments for τ and learns τ from each neighbor before stopping, so $\beta_\tau = 2\tau$ and welfare equals $\mathcal{V}(\tau, 2\tau) < V^*$.

Example 3, continued (Star): When there is a large number of peripherals, the core agent Kata shirks, $\tau_k = 0$. The peripherals thus do all the experimentation and have the lowest information and welfare, so we focus on them. We show in Section IIIC that agents asymptotically learn if and only if $p_0 \geq p^s$. The threshold p^s is defined so that peripheral agents barely work at $t = 0$ if they think Kata will instantly learn the state and choose $b_{k,t} \equiv 1$ thereafter,

$$(12) \quad \psi_{\ell,0} = p^s \left(x + \frac{r}{r+1} y \right) - c = 0.$$

Note that $p^s < \bar{p}$, so agents asymptotically learn if they have an optimistic prior.

The welfare result is exactly the opposite. Agents attain the welfare benchmark if and only if $p_0 \leq p^s$. For a high prior, $p_0 > p^s$, the peripheral agents experiment in the limit, $\tau_\ell > 0$, meaning Kata instantly learns. Thus, a peripheral agent learns $\beta_{\tau_\ell} = \tau_\ell$ before stopping and has value $\mathcal{V}(\tau_\ell, \tau_\ell) < \mathcal{V}(0,0) = V^*$.²¹ For a low prior, $p_0 \leq p^s$, the peripherals stop experimenting in the limit, $\tau_\ell = 0$, and since $\beta_{\tau_\ell} < \tau_\ell$, their value converges to its upper bound $V^* = \mathcal{V}(0,0)$.²² Thus, asymptotic learning and the welfare benchmark are not only distinct concepts but in fact mutually exclusive (for generic priors $p_0 \neq p^s$).

B. Large Random Networks

We first study large random networks. This is a tractable and canonical class of networks that can capture realistic contagion dynamics. For simplicity, we focus on regular networks, where agents all have the same number of neighbors. This class

²¹The latter equality presumes $p_0 \leq \bar{p}$. For $p_0 > \bar{p}$, the welfare benchmark $V^* = p_0 y$ requires immediate perfect social information, which is clearly impossible with a single neighbor.

²²It may appear paradoxical that the welfare upper bound V^* is achieved for low but not high prior beliefs p_0 . This is because V^* itself rises as function of p_0 and is hence a more demanding benchmark for high p_0 .

is rich enough to encompass the clique and trees, as in Sadler (2020a) and Board and Meyer-ter-Vehn (2021). After our main result, we discuss which insights generalize to nondegenerate degree distributions.

We construct a *regular random network* as follows. Each of the I agents has $\hat{n}^I \geq 2$ link stubs. We randomly draw pairs of stubs and connect them into undirected links. We then prune self-links (from i to i), multilinks (from i to j), and, if $\hat{n}^I I$ is odd, the single leftover stub. We assume that agents observe nothing about the network realization, not even their own degree; omitting such information seems innocuous since agents asymptotically know their degree (see Lemma 3).

By construction, the random network is exchangeable, so Corollary 1 implies there is a unique equilibrium. Denote the symmetric cutoff by τ^I and agents' value by V^I . We consider sequences of such networks with degrees $\{\hat{n}^I\}$ and assume existence of the limits $\nu := \lim \hat{n}^I \geq 2$,²³ $\lambda := \lim \hat{n}^I / \log I$, and $\hat{\rho} := \lim \hat{n}^I / I$, possibly equal to ∞ .

Let N^I be the number of realized links of a random agent. Some stubs may fail to form links, so N^I is random with expectation $n^I := E[N^I] < \hat{n}^I$. We now argue that we can ignore this complication as $I \rightarrow \infty$.

LEMMA 3: *As the network grows large, $I \rightarrow \infty$:*

- (a) *Realized degree: $N^I / n^I \xrightarrow{D} 1$.*
- (b) *Expected degree: $n^I / I \rightarrow 1 - e^{-\hat{\rho}}$. If $\hat{\rho} = 0$, then $n^I / \hat{n}^I \rightarrow 1$.*
- (c) *Social information at the cutoff time: $\lim \beta_{\tau^I}^I = \lim n^I \tau^I$.*

PROOF:

See online Appendix C.3.

Part (a) means agents know their realized degree N^I in the limit. Part (b) means we can ignore the distinction between stubs and links when $\hat{\rho} = 0$. And part (c) means agents do not update N^I during the experimentation phase, a consequence of part (a).

We next introduce three relevant regions of limit network density:

- (i) *Sparse Networks.*—Here, agents have a bounded number of links, with $\nu := \lim \hat{n}^I = \lim n^I \in \{2, 3, \dots\}$. Such networks approximate trees in the following local sense: for any agent i and any $r \in \mathbb{N}$, the probability that i has ν first neighbors, $\nu(\nu - 1)$ second neighbors, ..., $\nu(\nu - 1)^{r-1}$ neighbors at distance r , and all these agents are distinct converges to one as $I \rightarrow \infty$. Board and Meyer-ter-Vehn (2024) show that contagion dynamics and equilibrium converge to those of infinite trees in Example 2'.

²³The restriction $\hat{n}^I \geq 2$ ensures that the component of a typical agent has size proportional to I .

- (ii) *Intermediate Networks*.—Here, agents' links are of order $\log I$, with $\lambda := \lim \hat{n}^I / \log I = \lim n^I / \log I \in (0, \infty)$. In such networks, information spreads across the network in finite time, as in Milgram's (1967) six degrees of separation. Indeed, Lemma 4 shows that the inverse $1/\lambda$ measures the network's *time-diameter*, i.e., the time for information to travel between two random agents in the network.²⁴
- (iii) *Dense Networks*.—Here, agents are connected to a fixed proportion of other agents $\rho := \lim n^I / I = 1 - e^{-\hat{\rho}} \in [0, 1]$. Agents are at most two links apart, and we approximate the clique from Example 1 when $\rho = 1$.

The set of network densities is the union $\{\nu | \nu \in \mathbb{N}\} \cup \{\lambda \cdot \log I | \lambda \in [0, \infty]\} \cup \{\rho \cdot I | \rho \in [0, 1]\}$, endowed with its natural order, after identifying $\infty \cdot \log I$ with $0 \cdot I$.²⁵

We now define the threshold density for asymptotic learning. For pessimistic priors, $p_0 < \bar{p}$, let $\sigma^* \in [0, \infty)$ be such that perfectly learning the state at time σ^* renders an agent indifferent about experimenting at $t = 0$. Using (4),

$$(13) \quad \psi_0 = p_0(x + (1 - e^{-r\sigma^*})y) - c = 0.$$

Here, $e^{-r\sigma^*}y$ is the post-experimentation continuation value. For optimistic priors, $p_0 \geq \bar{p}$, set $\sigma^* = 0$.

THEOREM 1: *In random networks $\{\hat{n}^I\}$ as $I \rightarrow \infty$.*

- (a) *Asymptotic information β is a decreasing function of network density. It attains asymptotic learning $\beta = \infty$ if and only if $\lambda \leq 1/\sigma^*$, and it strictly falls when $\lambda \geq 1/\sigma^*$.²⁶*
- (b) *Welfare V is a single-peaked function of network density. It strictly rises when $\nu < \infty$, attains the benchmark V^* if and only if $\nu = \infty$ and $\rho = 0$, and then strictly falls when $\rho > 0$.*

PROOF:

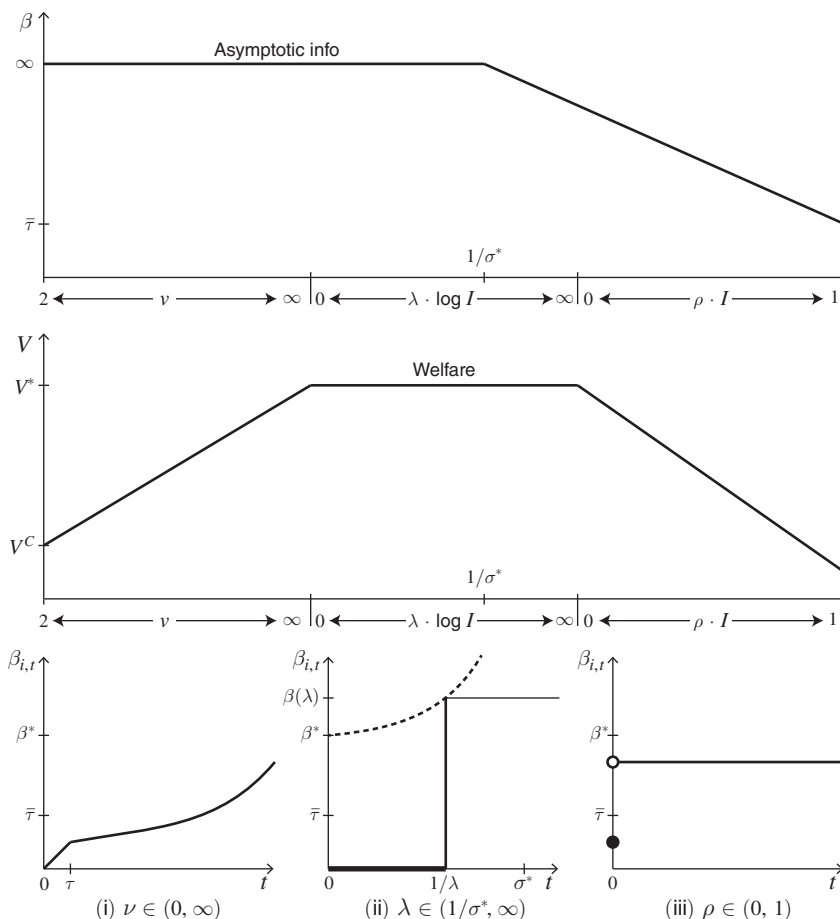
See Appendix B.A.

Asymptotic learning requires sparse networks. Intuitively, denser networks accelerate diffusion, crowd out discovery, and undermine learning in the long run.

²⁴This time-diameter $1/\lambda = \lim(\log I / n^I)$ is smaller than the typical diameter estimate for large random networks $\lim(\log I / \log n^I)$. The smaller diameter reflects a faster contagion process: contagion in our model does not travel one link in every discrete time period; rather, each link transmits continuously with rate one. Much like compound interest, this allows nodes infected at $t' \in [t, t+1]$ to begin transmitting immediately instead of having to wait until $t+1$.

²⁵This order treats many sequences of networks as equally dense. For instance, $n^I = \log \log I$ or $n^I = (\log I)^{1/2}$ both correspond to $\nu = \infty, \lambda = 0$. Theorem 1 shows that asymptotic information β and welfare V of a sequence of networks $\{n^I\}$ only depend on its limit density (ν, λ, ρ) .

²⁶For optimistic priors $p_0 \geq \bar{p}$, where $1/\sigma^* = \infty$, asymptotic information is perfect if and only if $\rho = 0$ and strictly falls in $\rho > 0$.

FIGURE 5. LARGE RANDOM NETWORKS FOR PESSIMISTIC PRIORS, $p_0 \leq \bar{p}$

Notes: The top panel shows asymptotic information β as a function of network density, as described in Theorem 2(a). The middle panel shows welfare V as a function of network density, as described in Theorem 2(b). The bottom panel shows the cumulative social learning curves of a typical agent in three canonical cases, as discussed in the text. Note, σ^* is defined by (13), $\beta(\lambda)$ by (15), and β^* by (16).

Welfare attains the benchmark when network density is intermediate. Intuitively, welfare discounts the future and so relies on both information generation and its quick dissemination.

Theorem 1 goes beyond traditional threshold theorems (see, e.g., Jackson 2010, Section 4.2.2). First, we solve for the exact threshold for asymptotic learning within the logarithmic range, $\lambda \leq 1/\sigma^*$. Second, we characterize learning and welfare for network densities where the upper bounds are not attained.

Figure 5 illustrates Theorem 1 for $p_0 < \bar{p}$. The top and middle panels sketch asymptotic information β and welfare V as functions of network density. The bottom panel illustrates the underlying cumulative social learning curves $\{\beta_t\}$ for our three regions of network density.²⁷ We discuss Figure 5 in order of increasing network density.

²⁷We illustrate cumulative social learning $\{\beta_t\}$ since the rate $\{b_t\}$ (in Figure 3) fails to exist for $\nu = \infty$.

We begin with sparse networks, $\nu < \infty$. As $I \rightarrow \infty$, these networks approximate trees, with independent information across Iris's neighbors. Cumulative social learning in the contagion phase $\{\beta_t\}_{t \geq \tau}$, illustrated in Figure 5(i), is convex, with rate $b_t = n a_t$ described by (11). This convexity reflects the fact that an agent has ν first-degree neighbors, $\nu(\nu - 1)$ second-degree neighbors, $\nu(\nu - 1)^2$ third-degree neighbors, etc., so contagion accelerates over time. Each agent experiments for a bounded time $\tau > 0$, which ensures asymptotic learning, while welfare falls short of the benchmark, $\mathcal{V}(\tau, \nu\tau) < \mathcal{V}(0, 0)$. The proof shows that $\tau = \tau(\nu)$ falls in ν and $\mathcal{V}(\tau, \nu\tau)$ rises in ν .

Next, we characterize diffusion in intermediate and dense networks, $\nu = \infty$. As illustrated in Figure 5(ii) and Figure 5(iii), the cumulative social learning curve $\{\beta_t\}$ is a step function with a single step at time σ . That is, agents observe the first success at time σ , or never. In analogy to epidemiological contagion processes, we also say that agents get "exposed" at σ . For intermediate networks, we have $\sigma = \sigma(\lambda) > 0$; for dense networks, we have $\sigma = 0$.

To state the underlying result, consider any sequence of cutoffs $\{\tau^I\}$ (not necessarily equilibrium) with limit $\sigma := \lim_{I \rightarrow \infty} \frac{-\log \tau^I}{n^I} \in [0, \infty]$. Let S be the binary random time, with $\Pr(S = \sigma) = 1 - e^{-\lim_{I \rightarrow \infty} I\tau^I}$ and $\Pr(S = \infty) = e^{-\lim_{I \rightarrow \infty} I\tau^I}$,²⁸ and index i 's exposure time S_i^I by I .

LEMMA 4: *Assume $\nu = \infty$. As $I \rightarrow \infty$, i gets exposed at time σ or never, $S_i^I \xrightarrow{D} S$, and learns all generated information, $\lim_{I \rightarrow \infty} I\tau^I = \beta$.*

PROOF:

The full proof is in Appendix B.B. For an intuition, suppose Iris's neighbors are a negligible share of the population, $\rho = 0$. At τ^I , the chance at least one agent has succeeded is $1 - e^{-I\tau^I}$, and there are approximately $n^I I\tau^I$ exposed agents. The contagion then grows exponentially at rate n^I (which itself diverges $n^I \rightarrow \infty$), so there are approximately $n^I I\tau^I e^{n^I t}$ exposed agents at time t , and, heuristically, everyone is exposed when $n^I I\tau^I e^{n^I t} = I$, or $t = \frac{-\log(n^I I\tau^I)}{n^I} \rightarrow \sigma$.²⁹ This argument slightly overstates exposures because of double-counting. But this problem scales with the share of exposed agents, and we only need the argument as long as this share is negligible; once a fixed share of the population is exposed, all agents are exposed immediately since $n^I \rightarrow \infty$. The proof uses Chernoff bounds to make these arguments rigorous. ■

To sharpen Lemma 4, note that the time-diameter upper-bounds the exposure time

$$(14) \quad \sigma = \lim_{I \rightarrow \infty} \frac{-\log \tau^I}{n^I} = \lim_{I \rightarrow \infty} \frac{\log I - \log I\tau^I}{n^I} = \frac{1}{\lambda} - \lim_{I \rightarrow \infty} \frac{\log I\tau^I}{n^I}.$$

²⁸Note that others' information equals total information, $\lim_{I \rightarrow \infty} (I - 1)\tau^I = \lim_{I \rightarrow \infty} I\tau^I$, both when $\tau^I \rightarrow 0$ and when τ^I is bounded away from zero (so both limits are infinite).

²⁹Recalling footnote 24, here, we see the difference between typical discrete-time contagion models where exposed agents grow like $e^{(\log n^I)t}$ and our continuous-time model with the faster rate $e^{n^I t}$.

With finite aggregate information, $\beta = \lim I\tau^I < \infty$, they coincide $\sigma = 1/\lambda$, but with $\beta = \infty$, a diverging number of agents succeed during experimentation, so exposure can happen earlier, $\sigma \leq 1/\lambda$.

With this characterization of social learning curves for any cutoffs τ^I , we now return to the question of equilibrium. For intermediate networks, the Proof of Theorem 1 shows that pre-cutoff learning must vanish, $(n^I + 1)\tau^I \rightarrow 0$.³⁰ Welfare thus converges to the benchmark $\mathcal{V}(\tau^I, n^I\tau^I) \rightarrow \mathcal{V}(0, 0) = V^*$.

Turning to asymptotic information, the indifference condition (4) at $t = 0$ when anticipating information β at exposure time σ becomes

$$(15) \quad p_0\left(x + \left[1 - e^{-r\sigma}(1 - e^{-\beta})\right]y\right) - c = 0.$$

For low-density intermediate networks with $\lambda \in (0, 1/\sigma^*)$, an exposure time equal to the time-diameter $\sigma = 1/\lambda > \sigma^*$ would render experimentation incentives (15) positive for any β . That is, if the delay exceeded σ^* , no amount of information could fully crowd out own experimentation. Equilibrium must therefore feature $\sigma(\lambda) = \sigma^* < 1/\lambda$, implying infinite information, $\beta(\lambda) = \infty$, by (14), so (15) becomes (13). Thus, in this range, the exposure time and information are independent of network density.

For high-density intermediate networks $\lambda \in (1/\sigma^*, \infty)$, perfect learning $\beta = \infty$ would render incentives (15) negative for any $\sigma(\lambda) \leq 1/\lambda < \sigma^*$. That is, learning is so fast that perfect information at $\sigma(\lambda)$ would choke off experimentation entirely. Instead, equilibrium must feature finite information, $\beta(\lambda) < \infty$, implying $\sigma(\lambda) = 1/\lambda$ by (14). The resulting $(\sigma(\lambda), \beta(\lambda)) = (1/\lambda, \beta(\lambda))$ are described by initial indifference (15) and illustrated by the dashed line in Figure 5(ii). As density λ rises from $1/\sigma^*$ to ∞ , the exposure time $\sigma(\lambda)$ falls from σ^* to zero, and asymptotic information $\beta(\lambda)$ falls from $\beta(\lambda) = \infty$ at $\lambda = 1/\sigma^*$, as captured by (13), to $\beta^* = \beta(\infty)$ defined by

$$(16) \quad p_0(x + e^{-\beta^*}y) = c.$$

Intuitively, the higher incentives due to earlier learning are compensated by less information $\beta(\lambda)$ in order to maintain indifference.

Finally, consider dense networks, where agents are connected to a fixed proportion $\rho \in (0, 1)$ of others. Learning is immediate, as seen in Figure 5(iii). Such networks are analogous to the clique. With total information β , agents learn $\rho\beta$ before stopping and $(1 - \rho)\beta$ immediately after stopping. The indifference condition,

$$P^\emptyset(\rho\beta)(x + e^{-(1-\rho)\beta}y) = c,$$

³⁰Intuitively, with $\nu = \infty$ and $\rho = 0$, the ratio of second neighbors to first neighbors diverges, so nonzero pre-cutoff learning implies immediate, perfect post-cutoff learning and chokes off experimentation for $p_0 < \bar{p}$.

then determines total information β . The solution β falls in ρ , and since learning is immediate, welfare also falls in ρ . As $\rho \rightarrow 1$, we approach the clique, with asymptotic information $\beta \rightarrow \bar{\tau}$ and welfare $V \rightarrow \mathcal{V}(0, \bar{\tau})$.

Theorem 1 is stated for regular, undirected networks. The analysis immediately extends to regular directed networks. Networks with nondegenerate degree distributions introduce an alternative possibility for asymptotic learning to fail. An agent may be isolated, or more generally, the size of her limit component may be finite. This arises with positive probability in Erdos-Renyi networks with bounded expected degree n^l ; asymptotic learning then requires intermediate network density with $\nu = \infty$ and $\lambda \leq 1/\sigma^*$. We study networks with heterogeneous, finite degrees in Board and Meyer-ter-Vehn (2024).³¹

C. Core-Periphery Networks

In this section we study core-periphery networks. Theorem 2 shows that asymptotic information falls with network density, while welfare is single peaked, echoing Theorem 1 for random networks. This analysis serves three purposes. First, core-periphery networks are of intrinsic interest. They are used to describe financial markets (e.g., Li and Schürhoff 2019) and can arise endogenously in network formation models (Galeotti and Goyal 2010). Second, core-periphery networks allow us to examine the role of network position for information generation. Third, core-periphery networks have a different neighborhood structure, with relatively few first neighbors in the core slowly transmitting the information generated by the more numerous peripherals. As a result, social learning curves are then concave rather than convex in the contagion phase.

A *core-periphery network* is an undirected, deterministic network that consists of K core agents and $L = I - K$ peripheral agents. The core agents k are connected to everyone. The peripheral agents ℓ are only connected to core agents. See Figure 6 for an illustration. When $K = 1$, we have the star from Example 3.

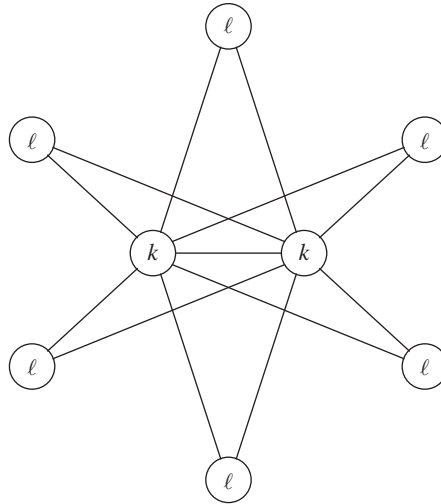
LEMMA 5: *Any equilibrium in a core-periphery network is characterized by two cutoffs, τ_k for all agents in the core and τ_ℓ for all peripherals. Core agents work less, $\tau_k < \tau_\ell$, and have higher values, $V_k > V_\ell$.³²*

PROOF:

By exchangeability and Proposition 3, equilibrium is characterized by cutoffs (τ_k, τ_ℓ) . Core agents k observe all successes immediately, so they have greater total information than peripherals who observe some successes with delay, $\beta_{k,t} + \min\{t, \tau_k\} > \beta_{\ell,t} + \min\{t, \tau_\ell\}$ for all $t > 0$. Lemma 8 in online Appendix C.2

³¹Networks with nondegenerate degree distributions challenge our assumption that agents do not observe their own degree. The role of this assumption is to guarantee symmetry, so equilibrium is unique and characterized by a single cutoff by Corollary 1. If instead agents observe their degree, equilibrium is characterized by a multidimensional fixed point, which makes it difficult to derive comparative statics.

³²Lemma 5 and its proof apply as stated to nested split graphs (Koenig, Tessone, and Zenou 2014). In such networks, high-degree agents work less and have higher values than low-degree agents in any equilibrium.

FIGURE 6. CORE-PERIPHERY NETWORK WITH $k = 2$ CORE AGENTS AND $\ell = 6$ PERIPHERALS

implies $\tau_k < \tau_\ell$.³³ Since peripherals experiment more, core agents have greater social learning, $\beta_{k,t} > \beta_{\ell,t}$ for all $t > 0$, so $V_k > V_\ell$ by Lemma 1. ■

We now characterize equilibrium cutoffs. Core agents k observe all successes immediately, so their social learning follows $b_{k,t} \equiv (K - 1) \mathbf{1}\{t \leq \tau_k\} + L \mathbf{1}\{t \leq \tau_\ell\}$. Experimentation incentives (4) are given by

$$(17) \quad \psi_{k,\tau_k} = P^\emptyset(I\tau_k) \left(x + y \left[1 - \left[1 - e^{-(r+L)(\tau_\ell - \tau_k)} \right] \frac{L}{r+L} \right] \right) - c,$$

where the opportunity cost is the continuation value from having L peripherals experiment over $[\tau_k, \tau_\ell]$. In equilibrium, $\psi_{k,\tau_k} \leq 0$ with equality if $\tau_k > 0$.

Peripheral agents ℓ only observe the successes of core agents, so their social learning $b_{\ell,t}$ equals K before τ_k and then drops to Ka_t , where $a_t := \Pr^{-\ell}(T_{\ell'} < t \text{ for at least one } \ell' \neq \ell | t < T_k \text{ for all } k)$ is the conditional probability that some other peripheral agent has succeeded by t and hence, the core agents are working. This follows

$$(18) \quad \frac{\dot{a}}{1-a} = (L-1) \mathbf{1}\{t \leq \tau_\ell\} - Ka = \begin{cases} L-1-Ka, & t \in (\tau_k, \tau_\ell) \\ -Ka, & t > \tau_\ell \end{cases}$$

with boundary condition $a_{\tau_k} = 1 - e^{-(L-1)\tau_k}$, as shown in online Appendix C.4. Before τ_ℓ , social learning a_t rises because of experimentation by the other $L-1$ peripherals, tempered by the lack of success by the K core agents. After τ_ℓ , only the

³³There is a subtlety here. Lemma 1 tells us that more social learning leads to less experimentation, but this is insufficient to conclude that core agents experiment less. For example, consider the star network and assume peripherals do not experiment; the core agent then has no social information but the same amount of total information as peripherals. In the online Appendix Lemma 8 adapts the arguments from Lemma 1 to show that greater total learning (including self-learning) implies less experimentation.

latter effect remains, so learning $b_{\ell,t} = Ka_t$ slows down. Using equation (4), peripherals' cutoff $\tau_\ell > 0$ then solves

$$\psi_{\ell,\tau_\ell} = P^\emptyset \left(K \left[\tau_k + \int_{\tau_k}^{\tau_\ell} a_s ds \right] + \tau_\ell \right) \left(x + ry \int_{\tau_\ell}^{\infty} e^{-\int_{\tau_\ell}^t (r+Ka_s) ds} dt \right) - c = 0.$$

For fixed $I < \infty$, we do not know whether the equilibrium cutoffs (τ_k, τ_ℓ) are unique, but the Proof of Theorem 2 shows that any equilibrium converges to the same limit.

In order to cleanly characterize how social information and welfare depend on the network density, we consider sequences of core-periphery networks with core sizes $\{K^I\}_{I \in \mathbb{N}}$. We assume the following two limits exist. Let $\kappa := \lim K^I \in \mathbb{N} \cup \{\infty\}$ be the limit of absolute core size, and $\rho := \lim K^I/I \in [0, 1]$ the limit of relative core size, as a proportion of the population. The set of network densities is the union $\{\kappa | \kappa \in \mathbb{N}\} \cup \{\rho \cdot I | \rho \in [0, 1]\}$ endowed with its natural order.

We now define a threshold on core size that is critical for both asymptotic learning and welfare. For pessimistic priors $p_0 < \bar{p}$, define $\kappa^* \in (0, \infty)$ such that learning from κ^* core agents who experiment forever, $b_{\ell,t} \equiv \kappa^*$, renders a peripheral agent indifferent about experimenting at $t = 0$,

$$(19) \quad \psi_{\ell,0} = p_0 \left(x + \frac{r}{r + \kappa^*} y \right) - c = 0.$$

For optimistic priors, $p_0 \geq \bar{p}$, set $\kappa^* = \infty$.

THEOREM 2: *In core-periphery networks $\{K^I\}$ and any equilibria $\{\tau_k^I, \tau_\ell^I\}$ as $I \rightarrow \infty$,*

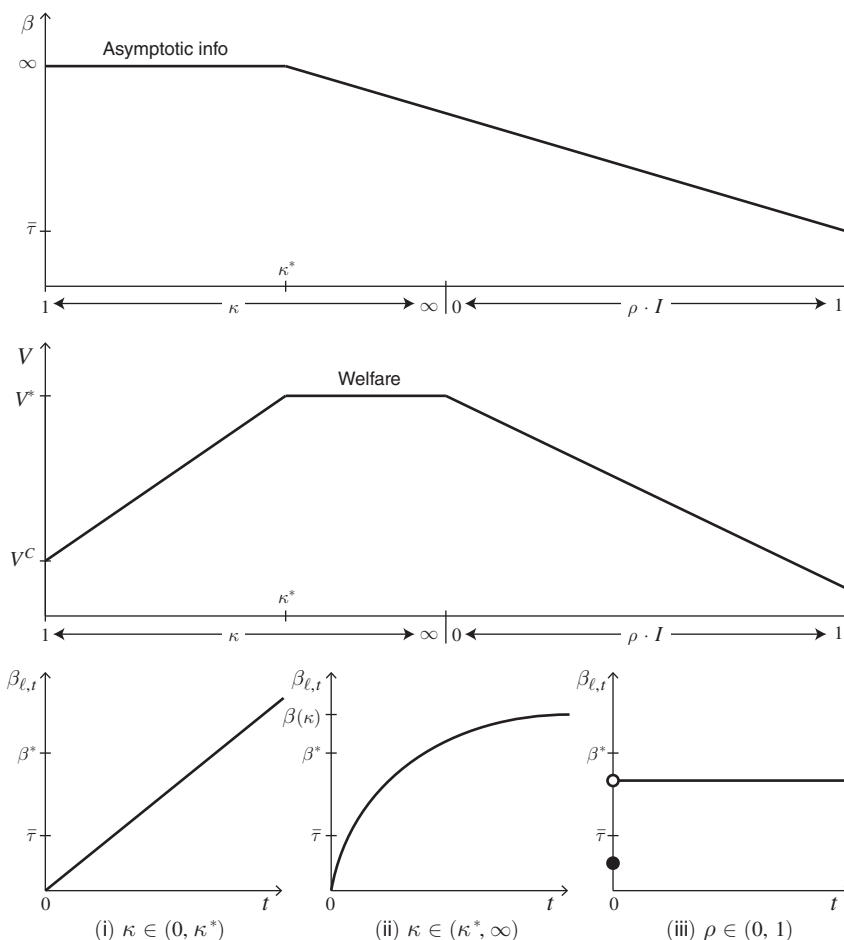
- (a) *Asymptotic information β is a decreasing function of network density. It attains asymptotic learning $\beta = \infty$ if and only if $\kappa \leq \kappa^*$ and strictly falls when $\kappa \geq \kappa^*$.³⁴*
- (b) *Welfare V is a single-peaked function of network density. It strictly rises when $\kappa \leq \kappa^*$, it attains the benchmark V^* if and only if $\kappa \in [\kappa^*, \infty]$ and $\rho = 0$, and strictly falls when $\rho > 0$.*

PROOF:

See online Appendix C.5.

As $I \rightarrow \infty$, asymptotic learning is achieved for sufficiently small core size; welfare attains the benchmark for intermediate core size. Figure 7 illustrates Theorem 2 for $p_0 < \bar{p}$. The top and middle panels sketch asymptotic information β and welfare V as functions of core size. The bottom panel illustrates three typical social learning

³⁴For optimistic priors $p_0 \geq \bar{p}$, where $\kappa^* = \infty$, asymptotic information is perfect if and only if $\rho = 0$ and strictly falls in $\rho > 0$.

FIGURE 7. CORE-PERIPHERY FOR PESSIMISTIC PRIORS, $p_0 \leq \bar{p}$.

Notes: The top panel shows asymptotic information β as a function of network density, as described in Theorem 2(a). The middle panel shows welfare V as a function of network density, as described in Theorem 2(b). The bottom panel shows the learning curves of a peripheral agent in three regions of network density, as discussed in the text. Note, β^* is defined by (16), κ^* by (19), and $\beta(\kappa)$ by (20).

curves $\{\beta_{\ell,t}\}$. While asymptotic learning and second-best welfare may a priori seem to be related goals, Theorem 2 shows that for pessimistic priors, they are generically mutually exclusive. Asymptotic learning requires a small core size $\kappa \leq \kappa^*$, while second-best welfare requires a large core size $\kappa \geq \kappa^*$.³⁵

As with random networks, there are three regions of network density with qualitatively different social learning dynamics. First, consider a small core $\kappa < \kappa^*$, as illustrated in Figure 7(i). The exploding number of peripherals experiment for a bounded time interval, $\tau_\ell > 0$, and collectively create an exploding amount

³⁵The pessimistic prior assumption is important. For optimistic priors, $p_0 \geq \bar{p}$, it is easier to motivate agents to experiment. Our welfare benchmark requires asymptotic learning, and both of these goals are obtained if $\kappa = \infty$ and $\rho = 0$. While this is a single point $0 \cdot I$ in our density order, there are many sequences that satisfy both conditions, e.g., $K^I = \log I$, $K^I = \log \log I$, $K^I = \sqrt{I}$.

of information in an instant. This crowds out experimentation by core agents. Peripherals choose to experiment since the flow of social information is restricted by the small core size. Formally, $\beta_{\ell,t} = \kappa t$, so equation (19) implies $\psi_{\ell,0} > 0$ given that $\kappa < \kappa^*$. Asymptotic learning obtains, but since each peripheral generates a nonvanishing amount of information, welfare falls short of the benchmark $\mathcal{V}(\tau_\ell, \kappa \tau_\ell) < V^*$.

Second, consider an intermediate core $\kappa \in (\kappa^*, \infty)$, as illustrated in Figure 7(ii).³⁶ With this core size, perfect information from peripherals would crowd out peripherals' experimentation incentives. In equilibrium, peripheral agents lower their cut-offs, limiting their total information $\beta(\kappa) = \lim (I - K^I - 1) \tau_\ell^I < \infty$. The level of $\beta(\kappa)$ is determined by peripheral agents' indifference condition at $t = 0$,

$$(20) \quad \psi_{\ell,0} = p_0 \left(x + ry \int_0^\infty e^{-(rt+\beta_{\ell,t})} dt \right) - c = 0,$$

where ℓ 's social learning curve satisfies $1 - e^{-\beta_{\ell,t}} = (1 - e^{-\beta(\kappa)})(1 - e^{-\kappa t})$. Intuitively, ℓ learns the state if some peripheral learned it and a core agent succeeds. As in the star, $b_{\ell,t}$ falls over time as agents grow pessimistic about the chance that one of them succeeded. Asymptotic learning fails, but agents do obtain the welfare benchmark, $\mathcal{V}(0,0)$, as social pre-cutoff learning $\kappa \tau_\ell^I$ vanishes. For large κ , the core transmits information increasingly fast, reinforcing the crowding out and reducing asymptotic information. When $\kappa = \infty$ but $\rho = 0$, $\beta = \beta^*$ solves (16), so $\beta_{\ell,t}$ jumps to β^* and remains constant thereafter.

Third, consider a large core $\rho \in (0, 1]$, as illustrated by Figure 7(iii). Now core agents generate a nonvanishing share of total information. Social learning becomes immediate, $\beta_{\ell,t} = \beta_{k,t} = \beta$ for all $t > 0$, and core agents' indifference condition becomes

$$P^\emptyset(\beta_{k,\tau})(x + ye^{-(\beta-\beta_{k,\tau})}) - c = 0,$$

with pre-cutoff learning $\beta_{k,\tau} = \lim I \tau_k^I$. This equation together with the analogous, but more involved expression for peripherals' pre-cutoff learning $\lim \beta_{\ell,\tau_\ell}^I$ pin down aggregate information β , which falls in ρ . As $\rho \rightarrow 1$, we approach the clique, with asymptotic information $\beta \rightarrow \bar{\tau}$ and welfare $V \rightarrow \mathcal{V}(0, \bar{\tau})$.

These results are reassuringly parallel to the ones for random networks in Section IIIB. In both cases, asymptotic information decreases in density, while welfare is single peaked. The underlying economic forces share similarities that transcend these two cases. For example, the general tension between learning and welfare for pessimistic priors $p_0 < \bar{p}$ is apparent from the welfare benchmark, $V^* = \mathcal{V}(0,0)$, which requires individual experimentation to vanish, undermining asymptotic learning. However, significant differences arise from the higher ratio of second neighbors to first neighbors. First, the contrast between asymptotic information and welfare is starker. With $p_0 < \bar{p}$, asymptotic learning and second-best welfare are mutually exclusive under core-periphery networks; this stems from the

³⁶For optimistic priors, $p_0 \geq \bar{p}$, this region is empty. But, as pointed out in footnote 35, there are many networks with $\kappa = \infty$ and $\rho = 0$.

small diameter together with the discreteness of the core size. By contrast, large random networks with intermediate network density $\lambda \in [0, 1/\sigma^*]$ accommodate both goals, as information aggregation occurs a long way from the typical agent and the learning time σ adjusts continuously to its equilibrium level. Second, social learning slows down over time in core-periphery networks with a finite core, as the information trickles through the core. By contrast, social learning speeds up over time in random networks, as the number of indirect neighbors grows exponentially with path length.

In our core-periphery network, peripheral agents connect to all core agents. In financial applications, one might assume instead that peripheral broker-dealers each connect to a single hub agent, in a fully connected core. This network is sparser than our core-periphery networks in that the information flow between a typical pair of peripherals ℓ, ℓ' must pass through their associated pair of core agents k, k' instead of any core agent. This bounds peripherals' social learning $b_{\ell,t} \leq t$, which in turn bounds their welfare below V^* when $\kappa^* > 1$.

IV. Conclusion

In Mokyr's (1992, p. 176) study of the history of innovation, he writes that, in addition to financial incentives,

decentralization was equally important because it meant that search and experimentation were carried out by many independent units, possibly over and over again. This duplication of effort was not the most cost effective way of engaging in technological process [...] But this system minimizes the probability of a technological opportunity being missed.

This paper has studied such a decentralized society and showed that welfare is single peaked in network density. Centralized societies quickly spread information and minimize the wasteful duplication of effort, but innovations are more likely to be missed (i.e., society fails to aggregate information). Thus, our results formalize the general concern that the rise of interconnectedness (e.g., social media) may crowd out original thought and opinion formation and lead to less informed societal outcomes.

When it comes to a particular application, when is a network too dense? Indeed, a farmer might be part of a scarce network (if they learn about successful crops from neighbors) or a dense network (if they learn from Bayer representatives). Our analysis provides a detail-free thought experiment for the critical network density. If a farmer is happy to experiment even if she knows some other farmers have already succeeded, then the network is "sparse," asymptotic learning obtains, and institutions that enhance diffusion tend to raise welfare (e.g., the founding of agricultural universities in the nineteenth century). If she instead waits to learn via diffusion, the network is "dense," asymptotic learning fails, and further raising density may lower welfare (e.g., global research networks). Of course, our model abstracts from practical considerations that are important for policy recommendations. For example, the first agent to succeed may obtain higher profits, or Bayer may subsidize early adoption of the crop. To inform policies, one would also wish to solve for socially optimal experimentation patterns.

This paper focuses on the role of networks in facilitating social learning. Another possibility is that $I \rightarrow \infty$ agents are connected in a clique network and observe others' successes with a fixed delay $\sigma > 0$. Fix $p_0 < \bar{p}$ and define σ^* as in equation (13). When $\sigma > \sigma^*$, initial experimentation incentives are positive, so $\tau := \lim \tau^I > 0$ solves $P^\emptyset(\tau)(x + [1 - e^{-r(\sigma-\tau)}]y) = c$. Agents learn perfectly at σ , but welfare is below second-best $\mathcal{V}(\tau, 0) < V^*$. Conversely, when $\sigma < \sigma^*$, perfect learning at σ would eliminate experimentation incentives for finite I , so total information $\beta < \infty$ solves (15). Since $\tau^I \approx \beta/I \rightarrow 0$, welfare is second-best $\mathcal{V}(0, 0) = V^*$. As in core-periphery networks, perfect learning and the welfare benchmark are generically incompatible. In contrast, in large random networks the learning time σ is endogenous and equals σ^* for a wide range of intermediate network densities.

Finally, we return to the classic paper of Bala and Goyal (1998) for a more detailed comparison. Two key long-run predictions carry over: individual beliefs converge due to the martingale convergence theorem, and, in connected networks, agents' beliefs converge to consensus. However, two important differences arise. First, our forward-looking agents experiment strategically, and so they can achieve asymptotic learning and/or second-best welfare even for pessimistic initial beliefs, $p_0 < \bar{p}$. The second type of difference stems from more technical modeling assumptions. For example, asymptotic learning fails in Bala and Goyal's famous "royal family" example, in which a small clique of "royals" are observed by everyone, while a directed line of "peasants" are only observed by their neighbor. Intuitively, if all royals receive strong negative signals in period 1, everyone switches to the safe action forever after. By contrast, our agents asymptotically learn because some peasants succeed during the initial experimentation phase, which then spreads to everybody else. The key is that perfect good-news learning (or continuous-time experimentation with repeated imperfect Poisson signals) generates unbounded signals, which can overturn any imperfect social information.

APPENDIX A. PROOFS FROM SECTION II

A. Proof of Proposition 1 (Cutoff Strategies)

To formalize the discussion surrounding the statement of Proposition 1, we first introduce a shorthand for the time- t discount factor on time- s payoffs in the event that i observes no success over $[t, s]$ from (3),

$$\Lambda_{t,s} := e^{-r(s-t)} \left(p_t e^{-\int_t^s (a_u^\emptyset + b_u) du} + 1 - p_t \right) = e^{-\int_t^s r + p_u(a_u^\emptyset + b_u) du},$$

where the second expression integrates by parts and is convenient for taking derivatives. Next, we generalize Iris's value of optimal experimentation (3) to time- t continuation payoffs from arbitrary, integrable experimentation,

$$(21) \quad \Pi_t = \Pi_t(\{a_s^\emptyset\}) = \int_t^\infty \Lambda_{t,s} (p_s(a_s^\emptyset(x+y) + b_s y) - a_s^\emptyset c) ds,$$

and recall from footnote 7 the derivative with respect to “experimentation just before t ,” $\frac{\partial \Pi_0}{\partial a_t^\theta} := \lim_{\epsilon \rightarrow 0} \frac{1}{\epsilon} \left(\Pi_0(\{a_s^{t,\epsilon}\}) - \Pi_0(\{a_s^\theta\}) \right)$, where $a_s^{t,\epsilon} := a_s^\theta + \mathbf{1}\{s \in [t-\epsilon, t]\}$.

We will show that front-loading incentives are positive, equal to

$$(22) \quad -\frac{d}{dt} \frac{\partial \Pi_0}{\partial a_t^\theta} = \Lambda_{0,t} (r(p_t(x+y) - c) + p_t b_t(x-c)).$$

The term $r(p_t(x+y) - c)$ is the time value of front-loading own experimentation from $t + \delta$ to t , while $p_t b_t(x-c)$ captures the value of additional experimentation that arises from neighbors succeeding in $[t, t + \delta]$. Since (22) is positive, agents maximally front-load effort, so cutoff strategies are optimal.³⁷

To establish the second derivative (22), we first derive convenient expressions for Π and its various first derivatives. Truncating (3) for Π_0 at time t , we get

$$(23) \quad \Pi_0 = \int_0^t \Lambda_{0,s} (p_s(a_s^\theta(x+y) + b_s y) - a_s^\theta c) ds + \Lambda_{0,t} \Pi_t.$$

To compute $\partial \Pi_t / \partial p_t$, Bayes’ rule implies $\Lambda_{t,s} p_s = e^{-r(s-t)} p_t e^{-\int_s^t (a_u^\theta + b_u) du}$, and we rewrite (21)

$$(24) \quad \begin{aligned} \Pi_t = & \int_t^\infty e^{-r(s-t)} \left(p_t e^{-\int_t^s (a_u^\theta + b_u) du} [a_s^\theta(x+y) + b_s y] \right. \\ & \left. - [p_t e^{-\int_t^s (a_u^\theta + b_u) du} + 1 - p_t] a_s^\theta c \right) ds. \end{aligned}$$

Writing $\mathbf{a}_t := \int_t^\infty e^{-r(s-t)} (a_s^\theta c) ds$ with time derivative $\dot{\mathbf{a}}_t = r\mathbf{a}_t - a_t^\theta c$, $\Pi_t + \mathbf{a}_t$ is a linear function of the posterior belief p_t , and so

$$(25) \quad \frac{\partial \Pi_t}{\partial p_t} = \frac{1}{p_t} (\Pi_t + \mathbf{a}_t).$$

To compute $\partial \Pi_0 / \partial a_t^\theta$, define the derivative of the posterior belief $p_t = P^\theta(\alpha_t + \beta_t)$ with respect to “experimentation just before t ,”

$$\frac{\partial p_t(\{a_s^\theta\}_{s \geq 0})}{\partial a_t^\theta} = \lim_{\epsilon \rightarrow 0} \frac{1}{\epsilon} \left[p_t(\{a_s^{t,\epsilon}\}_{s \geq 0}) - p_t(\{a_s^\theta\}_{s \geq 0}) \right] = -p_t(1 - p_t),$$

³⁷To formally show the optimality of cutoff strategies, start with arbitrary cumulative experimentation $\{\tilde{\alpha}_t^\theta\}$. First, omit any experimentation after $\underline{t} := \inf\{t: p_t \leq \underline{p}\}$, i.e., consider $\alpha_t^\theta := \min\{\tilde{\alpha}_t^\theta, \bar{\alpha}\}$, where $\bar{\alpha} = \tilde{\alpha}_{\underline{t}}^\theta$. By (21), $\Pi_0(\{a_s^\theta\}) \geq \Pi_0(\{\tilde{a}_s^\theta\})$. Second, we argue $\Pi_0(\mathbf{1}\{s \leq \bar{\alpha}\}) \geq \Pi_0(\{a_s^\theta\})$. Write $T(\alpha) := \inf\{t: \alpha_t^\theta = \alpha\} \in [\alpha, \infty]$ for the inverse of $\{a_s^\theta\}$. Then, from low to high values of α , front-load experimentation $a_{T(\alpha)}^\theta$ from $T(\alpha)$ to α . Denote $\{a_s^{\theta,\alpha}\}$ the strategy after having front-loaded $[0, \alpha]$; i.e., $a_s^{\theta,\alpha}$ equals 1 for $s \in [0, \alpha]$, 0 for $s \in (\alpha, T(\alpha)]$, and a_s^θ for $s > T(\alpha)$; for $\alpha = \bar{\alpha}$, we obtain the cutoff strategy $a_s^{\theta,\bar{\alpha}} = \mathbf{1}\{s \leq \bar{\alpha}\}$. Since every infinitesimal front-loading of $a_{T(\alpha)}^\theta$ from $T(\alpha)$ to α raises profits, we get

$$\Pi_0(\mathbf{1}\{s \leq \bar{\alpha}\}) - \Pi_0(\{a_s^\theta\}) = \int_0^{\bar{\alpha}} \left(a_{T(\alpha)}^\theta \int_\alpha^{T(\alpha)} \left[-\frac{d}{dt} \frac{\partial \Pi_0}{\partial a_t^\theta}(\{a_s^{\theta,\alpha}\}) \right] dt \right) d\alpha \geq 0.$$

where $a_s^{t,\epsilon} := a_s^\emptyset + \mathbf{1}\{s \in [t-\epsilon, t]\}$; also, $\partial\Lambda_{0,t}/\partial a_t^\emptyset = -p_t\Lambda_{0,t}$. Similarly, differentiating payoff (23) with respect to $a_t^\emptyset$ and using (25),

$$(26) \quad \begin{aligned} \frac{\partial\Pi_0}{\partial a_t^\emptyset} &= \Lambda_{0,t} \left(p_t(x+y) - c + \frac{\partial\Pi_t}{\partial p_t} \frac{\partial p_t}{\partial a_t^\emptyset} - p_t\Pi_t \right) \\ &= \Lambda_{0,t} (p_t(x+y) - c - (1-p_t)\mathbf{a}_t - \Pi_t). \end{aligned}$$

Turning to the time derivatives, we first note $\dot{p}_t = -(a_t^\emptyset + b_t)p_t(1-p_t)$, $\partial\Lambda_{t,s}/\partial t = [r + p_t(a_t^\emptyset + b_t)]\Lambda_{t,s}$, and $\partial\Lambda_{0,t}/\partial t = -[r + p_t(a_t^\emptyset + b_t)]\Lambda_{0,t}$, differentiate (21)

$$\dot{\Pi}_t = -[p_t(a_t^\emptyset(x+y) + b_ty) - a_t^\emptyset c] + [r + p_t(a_t^\emptyset + b_t)]\Pi_t,$$

and then differentiate (26) to get (22),

$$\begin{aligned} \Lambda_{0,t}^{-1} \frac{d}{dt} \frac{\partial\Pi_0}{\partial a_t^\emptyset} &= -[r + p_t(a_t^\emptyset + b_t)][p_t(x+y) - c - (1-p_t)\mathbf{a}_t - \Pi_t] \\ &\quad - p_t(1-p_t)(a_t^\emptyset + b_t)(x+y + \mathbf{a}_t) - (1-p_t)(r\mathbf{a}_t - a_t^\emptyset c) \\ &\quad + [p_t(a_t^\emptyset(x+y) + b_ty) - a_t^\emptyset c] - [r + p_t(a_t^\emptyset + b_t)]\Pi_t \\ &= -r[p_t(x+y) - c] - p_tb_t(x-c). \end{aligned}$$

Having established that cutoff strategies are optimal, we now show that the optimal cutoff is the unique solution of $\psi_\tau = 0$, if $\psi_0 \geq 0$. For cutoff strategies $a_s^\emptyset = \mathbf{1}\{s \leq t\}$, we have $\mathbf{a}_t = \int_t^\infty e^{-r(s-t)}(a_s^\emptyset c)ds = 0$, and (24) simplifies to

$$\Pi_t = p_ty \int_t^\infty e^{-r(s-t)} b_s e^{-\int_t^s b_u du} ds = p_ty \left(1 - r \int_t^\infty e^{-\int_t^s (r+b_u) du} ds \right),$$

where the last equality uses integration by parts. Then (26) simplifies to

$$(27) \quad \begin{aligned} \Lambda_{0,t}^{-1} \frac{\partial\Pi(\mathbf{1}\{s \leq t\})}{\partial a_t^\emptyset} &= p_t(x+y) - c - \Pi_t \\ &= p_t \left(x + ry \int_t^\infty e^{-\int_t^s (r+b_u) du} ds \right) - c = \psi_t. \end{aligned}$$

Since $\Lambda_{0,t}^{-1} > 0$ and $\partial\Pi(\mathbf{1}\{s \leq t\})/\partial a_t^\emptyset$ falls in t , ψ_t strictly single-crosses from above. ■

For future reference, we summarize some properties of experimentation incentives

$$(28) \quad \psi_\tau(\{\beta_t\}) = P^\theta(\tau + \beta_\tau) \left(x + ry \int_\tau^\infty e^{-r(s-\tau) - (\beta_s - \beta_\tau)} ds \right) - c.$$

First, note that while (4) and (5) express ψ instead as a function of the social learning rate $\{b_t\}$, the definition and most properties of ψ extend to any increasing (not necessarily continuous or positive) cumulative social learning curve β_t .

LEMMA 6: *Properties of incentives $\psi_\tau(\{\beta_t\})$ in cutoff τ and social learning $\{\beta_t\}$.*

- (a) $\psi_\tau(\{\beta_t\})$ falls in $\{\beta_t\}$ and thus also in $\{b_t\}$ with partial derivative given in (5).
- (b) $\psi_\tau(\{\beta_t\})$ strictly single-crosses from above in τ and is equi-Lipschitz continuous in τ for all uniformly bounded $\{b_t\}$.
- (c) The root τ of $\psi_\tau = 0$ falls in $\{\beta_t\}$ and strictly falls in $\{b_t\}$.

B. Proof of Lemma 2 (Characterization of $\mathcal{V}$)

We first derive (6),

$$\begin{aligned} V &= \left(p_0 \int_0^\tau e^{-\int_0^t (r+b_s+1) ds} [x + (b_t + 1)y - c] dt \right) - \left[(1 - p_0) \int_0^\tau e^{-rt} c dt \right] \\ &\quad + e^{-r\tau} [p_0 e^{-\beta_\tau - \tau} + (1 - p_0)] V_\tau \\ &= p_0 y (1 - e^{-\beta_\tau - (r+1)\tau}) - (1 - p_0) c \frac{1 - e^{-r\tau}}{r} + e^{-r\tau} (p_0 e^{-\beta_\tau - \tau} [x + y - c] \\ &\quad - [1 - p_0] c) \\ &= \frac{p_0 x - c}{r} + e^{-r\tau} [p_0 e^{-\beta_\tau - \tau} (x - c) - (1 - p_0) c \frac{r - 1}{r}]. \end{aligned}$$

The first line conditions on θ at time 0 and truncates flow payoffs at $t = \tau$. The second line evaluates the first integral using $x - c = ry$, and the last term using $p_0 e^{-\beta_\tau - \tau} + (1 - p_0) = p_0 e^{-\beta_\tau - \tau} / p_\tau$ by Bayes' rule, and $V_\tau = p_\tau y \int_\tau^\infty b_t e^{-\int_t^\infty (r+b_s) ds} dt = p_\tau (x + y) - c$ (using $\psi_\tau = 0$). The last line uses $y = (x - c)/r$ and reorders terms.

The monotonicity in β_τ is immediate from (6). To see the monotonicity in τ , note that the first term in (6) is the payoff from experimenting forever. Thus, the second term is the option value of stopping earlier, which must be positive. Then

$$\begin{aligned} (29) \quad \partial_\tau \mathcal{V} &= -r e^{-r\tau} \left[p_0 e^{-\beta_\tau - \tau} (x - c) - (1 - p_0) c \frac{r - 1}{r} \right] - e^{-r\tau} p_0 e^{-\beta_\tau - \tau} (x - c) \\ &< -e^{-r\tau} p_0 e^{-\beta_\tau - \tau} (x - c) = \partial_\beta \mathcal{V} < 0. \blacksquare \end{aligned}$$

APPENDIX B. PROOFS FROM SECTION III

A. Proof of Theorem 1 (Large Random Networks)

1. *Trees:* $\nu < \infty$.—We wish to show that V rises in ν for finite degrees $\nu < \infty$. We characterized the heuristic equilibrium in the infinite regular tree in Example 2'; write $\tau^{(\nu)}$ for the equilibrium cutoff and $a_t^{(\nu)}(\tau)$ for a random neighbor's expected experimentation, when all agents use the same arbitrary cutoff τ . Board and Meyer-ter-Vehn (2024) shows that equilibrium in the large random networks converges to these heuristics.

We first show that $\tau^{(\nu)}$ falls in ν . For a given cutoff $\tau > 0$, we have $a_t^{(\nu+1)}(\tau) > a_t^{(\nu)}(\tau)$ for all $t \geq \tau$. To see this, at the cutoff we have $a_\tau^{(\nu+1)}(\tau) = 1 - e^{-\nu\tau} > 1 - e^{-(\nu-1)\tau} = a_\tau^{(\nu)}(\tau)$, and this ranking prevails for $t > \tau$ since the RHS of (11) rises in ν . Additionally, $a_t^{(\nu)}(\tau)$ rises in τ , strictly for $\tau < t$. By Lemma 6, for any $\tau \leq \tau^{(\nu+1)}$,

$$\begin{aligned} 0 &= \psi_{\tau^{(\nu+1)}}\left(\left\{(\nu+1)a_t^{(\nu+1)}(\tau^{(\nu+1)})\right\}\right) \leq \psi_\tau\left(\left\{(\nu+1)a_t^{(\nu+1)}(\tau)\right\}\right) \\ &< \psi_\tau\left(\left\{\nu a_t^{(\nu)}(\tau)\right\}\right), \end{aligned}$$

so in equilibrium we must have $\tau^{(\nu)} > \tau^{(\nu+1)}$, as desired.

Below, we argue more strongly that

$$(30) \quad (\nu+1)\tau^{(\nu)} \geq (\nu+2)\tau^{(\nu+1)}.$$

It follows that equilibrium values are increasing in degree. By Lemma 2,

$$\begin{aligned} V^{(\nu+1)} &= \mathcal{V}\left(\tau^{(\nu+1)}, (\nu+2)\tau^{(\nu+1)} - \tau^{(\nu+1)}\right) \\ &> \mathcal{V}\left(\tau^{(\nu+1)}, (\nu+1)\tau^{(\nu)} - \tau^{(\nu+1)}\right) > \mathcal{V}\left(\tau^{(\nu)}, \nu\tau^{(\nu)}\right) = V^{(\nu)}, \end{aligned}$$

where the first inequality uses (30), and the second that adding $\tau^{(\nu)} - \tau^{(\nu+1)} > 0$ to the first argument of $\mathcal{V}$ and subtracting it from the second argument decreases $\mathcal{V}$.

To see (30), assume by contradiction that $\frac{\tau^{(\nu+1)}}{\tau^{(\nu)}} > \frac{\nu+1}{\nu+2} > \frac{\nu-1}{\nu}$. Then $a_{\tau^{(\nu)}}^{(\nu)} = 1 - e^{-(\nu-1)\tau^{(\nu)}} < 1 - e^{-\nu\tau^{(\nu+1)}} = a_{\tau^{(\nu+1)}}^{(\nu+1)}$, and since the ODE (11) rises in ν , the inequality $a_{\tau^{(\nu)}+\delta}^{(\nu)} < a_{\tau^{(\nu+1)}+\delta}^{(\nu+1)}$ is preserved for all $\delta > 0$. We thus get

$$\begin{aligned} 0 &= \psi_{\tau^{(\nu)}}\left(\left\{\nu a_t^{(\nu)}(\tau^{(\nu)})\right\}\right) \\ &= P^\theta\left([\nu+1]\tau^{(\nu)}\right)\left(x + ry \int_{\delta=0}^{\infty} \exp\left\{-r\delta - \nu\alpha_{\tau^{(\nu)}+\delta}^{(\nu)}\right\}d\delta\right) - c \\ &> P^\theta\left([\nu+2]\tau^{(\nu+1)}\right)\left(x + ry \int_{\delta=0}^{\infty} \exp\left\{-r\delta - (\nu+1)\alpha_{\tau^{(\nu+1)}+\delta}^{(\nu+1)}\right\}d\delta\right) - c \\ &= \psi_{\tau^{(\nu+1)}}\left(\left\{(\nu+1)a_t^{(\nu+1)}(\tau^{(\nu+1)})\right\}\right), \end{aligned}$$

contradicting the fact that $\tau^{(\nu+1)}$ is the equilibrium cutoff. This contradiction establishes (30), and so $V^{(\nu+1)} > V^{(\nu)}$.

2. Exploding Number of Neighbors: $\nu = \infty$.—In this case, equilibrium cutoffs must vanish, $\tau^I \rightarrow 0$; otherwise, the posterior belief at the cutoff $P^\theta((n^I + 1)\tau^I) \rightarrow 0$, choking off experimentation incentives. Lemma 4 characterizes the limit of social learning curves as a burst of social information of size β at time σ , i.e., step functions $\beta_t = \beta \mathbf{1}\{t \geq \sigma\}$; for $\sigma = 0$, we split the burst into pre-cutoff information $\beta_\tau := \lim n^I \tau^I = \lim (n^I + 1)\tau^I$ and post-cutoff information $\beta - \beta_\tau$. The equilibrium indifference condition becomes

$$(31) \quad \begin{aligned} \psi_\tau &= \lim \psi_{\tau^I} = P^\theta(\beta_\tau) \left(x + ry \int_0^\infty e^{-rt - (\beta_t - \beta_\tau)} dt \right) - c \\ &= P^\theta(\beta_\tau) \left(x + \left[1 - e^{-r\sigma} (1 - e^{-(\beta - \beta_\tau)}) y \right] \right) - c = 0. \end{aligned}$$

To solve for $\beta_\tau \leq \bar{\tau}, \beta, \sigma$, we complement (31) with the simple observation that

$$(32) \quad \frac{\beta_\tau}{\beta} = \frac{\lim n^I \tau^I}{\lim I \tau^I} = \lim \frac{n^I}{I} = \rho,$$

and two conditions linking the learning time $\sigma = \lim \frac{-\log \tau^I}{n^I}$ to pre-cutoff learning β_τ and total learning β . First, bounded total learning implies the learning time equals the network's time-diameter:

$$(33) \quad \text{If } \beta = \lim I \tau^I < \infty, \text{ then } \sigma = \lim \frac{\log(I \tau^I) - \log \tau^I}{n^I} = \frac{1}{\lambda}.$$

Second, nonvanishing pre-cutoff learning implies immediate learning:

$$(34) \quad \text{If } \beta_\tau = \lim n^I \tau^I > 0, \text{ then } \sigma = \lim \frac{\log(n^I \tau^I) - \log \tau^I}{n^I} = 0.$$

Case 1, $\rho = 0$ and $p_0 > \bar{p}$: The optimistic prior together with the equilibrium condition (31) require nonvanishing pre-cutoff learning, $\beta_\tau > 0$, and so by (34) immediate learning, $\sigma = 0$. Since $\rho = 0$, (32) implies perfect learning $\beta = \infty$. Perfect immediate learning, $\beta = \infty, \sigma = 0$, in turn implies the welfare benchmark $V = p_0 y = V^*$.

Case 2, $\rho = 0$ and $p_0 \leq \bar{p}$: We first observe $\beta_\tau = 0$. Otherwise, if $\beta_\tau > 0$, the proof for Case 1 implies $\beta = \infty$ and $\sigma = 0$, and so experimentation incentives $\psi_\tau < p_0 x - c \leq 0$, contradicting equilibrium. This implies the welfare benchmark, $\lim \mathcal{V}(\tau^I, n^I \tau^I) = \mathcal{V}(0, 0) = V^*$.

Turning to asymptotic information, we now show that $\beta = \infty$ if and only if $\lambda \leq 1/\sigma^*$, and β strictly decreases in density above this threshold. First assume $\lambda \leq 1/\sigma^*$.³⁸ If by contradiction, learning was imperfect $\beta < \infty$, social learning

³⁸For $p_0 = \bar{p}$, we have $\sigma^* = 0$, so this condition is always satisfied.

happens too late, at $\sigma = 1/\lambda \geq \sigma^*$ by (33), so experimentation incentives are strictly positive,

$$\psi_\tau = p_0(x + [1 - e^{-r\sigma}(1 - e^{-\beta})]y) - c > p_0(x + [1 - e^{-r\sigma^*}]y) - c = 0,$$

contradicting equilibrium.

Next assume $\lambda > 1/\sigma^*$. This induces early social learning $\sigma \leq 1/\lambda < \sigma^*$, and equilibrium indifference $p_0(x + [1 - e^{-r\sigma}(1 - e^{-\beta})]y) - c = \psi_\tau = 0 = p_0(x + [1 - e^{-r\sigma^*}]y) - c$ requires $\beta < \infty$. Moreover, $\beta = \beta(\sigma)$ rises in σ and so falls in $\lambda = 1/\sigma$.

Case 3, $\rho > 0$: Then $\beta_\tau = \rho\beta > 0$; (34) then implies $\sigma = 0$, so (31) becomes $P^\emptyset(\rho\beta)(x + e^{-(1-\rho)\beta}y) = c$. Since pre-cutoff learning lowers experimentation incentives more than post-cutoff learning by (5), total learning $\beta = \beta(\rho)$ falls in ρ . Success is observed either immediately, with probability $p_0(1 - e^{-\beta})$, or never; so welfare $V = p_0(1 - e^{-\beta})y$ also falls in ρ . ■

B. Proof of Lemma 4 (Degenerate Exposure Time)

With probability $e^{-I\tau^I}$, no agent succeeds by τ^I , and so $S_i^I = \infty$; from here on, we condition on the complementary event that at least one agent succeeds during experimentation, triggering a contagion process. For now, we also restrict attention to $\lim \hat{n}^I/I = 0$, so that $\hat{n}^I/n^I \rightarrow 1$ by Lemma 3(a). This allows us to work with $\hat{n}^I$ for finite I but switch to n^I in the limit where $\sigma := \lim_{I \rightarrow \infty} \frac{-\log \tau^I}{n^I}$. We discuss the case $\lim \hat{n}^I/I > 0$ later.

The overarching proof strategy is to separate the “geographical”/network aspects of the contagion process from its timing. Specifically, we realize the randomness of the network G^I as agents succeed. To emphasize the analogy to epidemiological SI contagion processes, we refer to agents who have succeeded as *infected*. When k agents are infected, let X_k^I be the random number of *exposed agents*, i.e., that have observed a success but have yet to succeed themselves. Clearly, $X_k^I \leq \hat{n}^I k$; a (relative) *exposure gap*, $\Gamma_k^I := \frac{\hat{n}^I k - X_k^I}{\hat{n}^I k} > 0$, opens up after an exposed j agent succeeds because the exposing agent i already succeeded and cannot be reexposed, or a stub of a succeeding agent connects to an already exposed agent. For $\epsilon > 0$, write $E^I(\epsilon) := \{\Gamma_k^I < 3\epsilon \text{ for all } k \leq \epsilon I/\hat{n}^I\}$ for the event that the gap process remains bounded in early stages of the contagion.

LEMMA 7: For any $\epsilon > 0$, $\lim_{I \rightarrow \infty} \Pr^{-i}(\mathcal{E}^I(\epsilon)) = 1$.³⁹

³⁹Throughout this proof, we use the standard probability measure $\Pr^{-i}$ over the network G and others' success times T_{-i} , conditional on observing no success from i , which is the relevant measure for the time of i 's first observed success S_i^I . As $I \rightarrow \infty$, this coincides with the notationally simpler (but less meaningful) measure $\Pr$ over the network G and all agents' success times $\{T_j\}$. We omit the dependence of $\Pr^{-i}$ on I for notational simplicity.

We postpone the Proof of Lemma 7; the idea is that with $X_k^I \leq \epsilon I$ exposed agents, ϵ small, and $\hat{n}^I$ large, most stubs expose new agents.

For small ϵ , Lemma 7 means that after the approximately $\tau^I I$ initial infections in the experimentation phase, the contagion process resembles a collection of tree networks emanating from these “seeds” at exponential rate $\hat{n}^I$. We now argue that as $\hat{n}^I \rightarrow \infty$, this contagion process reaches a negligible fraction of all agents at any $\underline{t} < \sigma = \lim_{\hat{n}^I} \frac{-\log \tau^I}{\hat{n}^I}$ but approximately all agents at any $\bar{t} > \sigma$.

Specifically, write T_k^I for the k^{th} infection time, and K^I for the (random) number of infected agents at τ^I . Also define inter-arrival times in the contagion phase $\Delta_k^I := T_{k+1}^I - T_k^I$ for $k > K^I$ and $\Delta_k^I := T_{k+1}^I - \tau^I$ for $k = K^I$. The proof idea is to apply Chernoff bounds to $T_k^I - \tau^I = \sum_{\ell=K^I}^{k-1} \Delta_\ell^I$. Toward this goal, note that conditional on the realization of the “geographical exposure process” $\{X_k^I\}_{k \in [K^I, \epsilon I / \hat{n}^I]}$, inter-arrival times Δ_k^I are independent with arrival rate X_k^I . Conditional on $\mathcal{E}^I(\epsilon)$, we have $X_k^I \in [(1 - 3\epsilon) \hat{n}^I k, \hat{n}^I k]$ for $k \leq \epsilon I / \hat{n}^I$, and so

$$(35) \quad E^{-i} \left[e^{-\xi \Delta_k^I} | \mathcal{E}^I(\epsilon) \right] \leq \frac{\hat{n}^I k}{\hat{n}^I k + \xi} \quad \text{for all } \xi \geq 0,$$

$$(36) \quad E^{-i} \left[e^{\xi \Delta_k^I} | \mathcal{E}^I(\epsilon) \right] \leq \frac{(1 - 3\epsilon) \hat{n}^I k}{(1 - 3\epsilon) \hat{n}^I k - \xi} \quad \text{for all } \xi \in [0, (1 - 3\epsilon) \hat{n}^I k].$$

We now derive upper and lower bounds for the k^{th} success time T_k^I in the contagion phase $k \in [K^I, \epsilon I / \hat{n}^I]$; in the limit $I \rightarrow \infty$ these bounds are then shown to imply vanishing chances of getting exposed before σ and after σ , respectively. The upper bound is as follows:

$$\begin{aligned} (37) \quad \Pr^{-i} \left(T_k^I \leq \tau^I + \delta | \mathcal{E}^I(\epsilon), K^I = k^I \right) &= \Pr^{-i} \left(\sum_{\ell=K^I}^{k-1} \Delta_\ell^I \leq \delta | \mathcal{E}^I(\epsilon) \right) \\ &\leq \inf_{\xi \geq 0} e^{\xi \delta} \prod_{\ell=K^I}^{k-1} E^{-i} \left[e^{-\xi \Delta_\ell^I} | \mathcal{E}^I(\epsilon) \right] \\ &\leq \inf_{\xi \geq 0} \exp \left\{ \xi \delta - \sum_{\ell=K^I}^{k-1} \left[\log(\hat{n}^I \ell + \xi) - \log(\hat{n}^I \ell) \right] \right\} \\ &\leq \inf_{\xi \in [0, \hat{n}^I]} \exp \left\{ \xi \delta - \sum_{\ell=K^I}^{k-1} \frac{\xi}{\hat{n}^I} \left[\log(\hat{n}^I(\ell + 1)) - \log(\hat{n}^I \ell) \right] \right\} \\ &= \inf_{\xi \in [0, \hat{n}^I]} \exp \left\{ \xi \left(\delta - \frac{\log k - \log K^I}{\hat{n}^I} \right) \right\} \end{aligned}$$

The first equality drops the τ^I to focus on time since the cutoff, the first inequality is a Chernoff bound, the second uses (35), the third uses the concavity of the logarithm, and the final equality collapses the telescopic sum.

Next, we argue that for fixed $\epsilon > 0$ and the integer floor $k = \lfloor \epsilon I / \hat{n}^I \rfloor$, the fraction on the RHS of (37) (which approximates the time for the contagion process to reach k agents) converges to $\sigma = \lim_{\hat{n}^I \rightarrow \infty} \frac{-\log \tau^I}{\hat{n}^I}$:

$$(38) \quad \frac{\log \lfloor \epsilon I / \hat{n}^I \rfloor - \log K^I}{\hat{n}^I} \xrightarrow{D} \sigma.$$

For $\bar{\beta} = \lim I \tau^I < \infty$, this follows because K^I is almost surely bounded above, so as $\hat{n}^I \rightarrow \infty$, all terms other than $\frac{\log I}{\hat{n}^I}$ vanish, and $\lim_{\hat{n}^I \rightarrow \infty} \frac{\log I}{\hat{n}^I} = \lim_{\hat{n}^I \rightarrow \infty} \frac{\log I - \log \bar{\beta}}{\hat{n}^I} = \lim_{\hat{n}^I \rightarrow \infty} \frac{-\log \tau^I}{\hat{n}^I} = \sigma$. For $\bar{\beta} = \infty$, it follows because, by the law of large numbers, $\frac{K^I}{I \tau^I} \xrightarrow{D} 1$; equivalently, $\log K^I - \log I - \log \tau^I \xrightarrow{D} 0$, so the LHS of (38) becomes $\frac{-\log \tau^I}{\hat{n}^I}$, whose limit is σ .

Exposing any positive fraction $\epsilon > 0$ of nodes requires infecting at least $\epsilon I / \hat{n}^I$ agents, and the chance of this at any time $\underline{t} < \sigma$ vanishes,

$$\begin{aligned} \lim_{I \rightarrow \infty} \Pr^{-i}(T_{\lfloor \epsilon I / \hat{n}^I \rfloor}^I \leq \tau^I + \underline{t}) &= \lim_{I \rightarrow \infty} \Pr^{-i}(T_{\lfloor \epsilon I / \hat{n}^I \rfloor}^I \leq \tau^I + \underline{t} | \mathcal{E}^I(\epsilon)) \\ &\leq \inf_{\xi \geq 0} \exp\{\xi(\underline{t} - \sigma)\} = 0. \end{aligned}$$

The equality uses Lemma 7, and the inequality (37) and (38). Then $\Pr^{-i}(T_{\lfloor \epsilon I / \hat{n}^I \rfloor}^I \leq \underline{t}) \rightarrow 0$.

Finally, for any population share $\epsilon > 0$, the probability that a given agent i has been exposed by time $\underline{t}$ is bounded above by the sum of that share ϵ and the probability that more than share ϵ has been exposed by time $\underline{t}$, $\Pr^{-i}(S_i^I \leq \underline{t}) \leq \epsilon + \Pr^{-i}(T_{\lfloor \epsilon I / \hat{n}^I \rfloor}^I \leq \underline{t})$. Since this inequality holds for any $\epsilon > 0$, we have

$$(39) \quad \lim_{I \rightarrow \infty} \Pr^{-i}(S_i^I \leq \underline{t}) \leq \lim_{\epsilon \rightarrow 0} \lim_{I \rightarrow \infty} [\epsilon + \Pr^{-i}(T_{\lfloor \epsilon I / \hat{n}^I \rfloor}^I \leq \underline{t})] = 0.$$

Turning to the lower bound for T_k^I , using the same steps as for (37), but with (36) substituting for (35) for the second inequality,

$$\begin{aligned} \Pr^{-i}(T_k^I \geq \tau^I + \delta | \mathcal{E}^I(\epsilon), K^I) &\leq \inf_{\xi \geq 0} e^{-\xi \delta} \prod_{\ell=K^I}^{k-1} E^{-i}[e^{\xi \Delta_\ell^I} | \mathcal{E}^I(\epsilon)] \\ &\leq \inf_{\xi \in [0, (1-3\epsilon)\hat{n}^I]} \exp\left\{-\xi \delta + \sum_{\ell=K^I}^{k-1} [\log([1-3\epsilon]\hat{n}^I \ell) - \log([1-3\epsilon]\hat{n}^I \ell - \xi)]\right\} \\ &\leq \inf_{\xi \in [0, (1-3\epsilon)\hat{n}^I]} \exp\left\{-\xi \delta + \sum_{\ell=K^I}^{k-1} \frac{\xi}{(1-3\epsilon)\hat{n}^I} (\log([1-3\epsilon]\hat{n}^I \ell) \right. \\ &\quad \left. - \log([1-3\epsilon]\hat{n}^I [\ell-1]))\right\} \end{aligned}$$

$$= \inf_{\xi \in [0, (1-3\epsilon)\hat{n}^I]} \exp \left\{ -\xi \left[\delta - \frac{\log(k-1) - \log(K^I-1)}{(1-3\epsilon)\hat{n}^I} \right] \right\}.$$

As for the upper bound, for $k = \lfloor \epsilon I / \hat{n}^I \rfloor$, the fraction on the RHS converges, $\frac{\log(\epsilon I / \hat{n}^I - 1) - \log(K^I - 1)}{(1-3\epsilon)\hat{n}^I} \xrightarrow{D} \sigma / (1-3\epsilon)$, so for any $\bar{\delta} > \sigma / (1-3\epsilon)$ in the limit,

$$\lim_{I \rightarrow \infty} \Pr^{-i} \left(T_{\lfloor \epsilon I / \hat{n}^I \rfloor} \geq \tau^I + \bar{\delta} | \mathcal{E}^I(\epsilon) \right) \leq \inf_{\xi \geq 0} \exp \left\{ -\xi \left(\bar{\delta} - \frac{\sigma}{1-3\epsilon} \right) \right\} = 0.$$

Conditional on $\mathcal{E}^I(\epsilon)$, $\lfloor \epsilon I / \hat{n}^I \rfloor$ infections guarantee $\epsilon(1-3\epsilon)I$ exposures by $\tau^I + \bar{\delta}$. For small $\epsilon' > 0$, approximately $\epsilon'\epsilon(1-3\epsilon)I$ of these get infected by $\tau^I + \bar{\delta} + \epsilon'$, generating approximately $\hat{n}^I \epsilon' \epsilon (1-3\epsilon)I$ new exposure possibilities; that is, an exploding number $\hat{n}^I \epsilon' \epsilon (1-3\epsilon) \rightarrow \infty$ for every agent. Now, for any $\bar{t} > \sigma$, we choose $\epsilon, \epsilon' > 0$ small enough and I large enough that $\tau^I + \bar{\delta} + \epsilon' < \bar{t}$ for $\bar{\delta} := \sigma / (1-3\epsilon) + \epsilon' > \sigma$. As $I \rightarrow \infty$, all remaining nodes get exposed before $\tau^I + \bar{\delta} + \epsilon'$ and thus before $\bar{t}$ with probability

$$(40) \quad \lim_{I \rightarrow \infty} \Pr^{-i} (S_i^I \leq \bar{t}) = 1.$$

Jointly, (39) and (40) for any $\underline{t} < \sigma < \bar{t}$ establish Lemma 4.

The case $\lim \hat{n}^I / I > 0$.—So far, we assumed $\lim \hat{n}^I / I = 0$ so that $\lim n^I / \hat{n}^I = 1$. Otherwise, we have $\rho = \lim n^I / I = 1 - \exp(-\lim \hat{n}^I / I) > 0$, implying $\beta_\tau = \rho\beta > 0$, and so the desired learning time equals $\sigma = 0$ by (34). To see that learning is indeed immediate, note that the first infection exposes fraction $\rho > 0$ of nodes. The paragraph preceding (40) then implies that everyone is exposed immediately thereafter. ■

PROOF OF LEMMA 7:

We will construct $p(\epsilon) < 1$ such that for large I and any $k \leq \epsilon I / \hat{n}^I$, the chance of a large exposure gap is bounded above via

$$(41) \quad \Pr^{-i} (\Gamma_k^I > 3\epsilon) < p(\epsilon)^{\hat{n}^k}.$$

Since $\mathcal{E}^I(\epsilon)$ is the complement of the union of these events over $k \geq 1$, Boole's inequality implies $1 - \Pr^{-i}(\mathcal{E}^I(\epsilon)) \leq \sum_{k=1}^{\infty} p(\epsilon)^{\hat{n}^k} = p(\epsilon)^{\hat{n}^I} / (1 - p(\epsilon)^{\hat{n}^I}) \rightarrow 0$, which implies (41).

We construct $p(\epsilon)$ and show (41) with the help of Chernoff bounds. The increment $X_k^I - X_{k-1}^I$ counts the newly exposed agents at the k^{th} infection. If j was exposed himself, he exposes $\hat{n}^I - 1$ others and is himself deducted from X_k^I ; if j was not exposed, he exposes $\hat{n}^I$ others. Each agent exposed by j was already exposed with

probability at most $k\hat{n}^I/I$. Thus, writing Y_ν for iid binary random variables with $\Pr(Y_\nu = 1) = k\hat{n}^I/I$, and $Y_\nu = 0$ else, we can upper-bound the absolute exposure gap

$$(42) \quad \hat{n}^I k \Gamma_k^I = \hat{n}^I k - X_k^I = \sum_{\ell=1}^k [\hat{n}^I - (X_\ell^I - X_{\ell-1}^I)] \stackrel{D}{\leq} 2k + \sum_{\nu=1}^{k\hat{n}^I} Y_\nu.$$

Now define $p(\epsilon) := \inf_{\xi \geq 0} \left(\frac{E[e^{Y_\nu \xi}]}{e^{2\epsilon \xi}} \right)$. We have $p(\epsilon) < 1$ since $E[Y_\nu] = k\hat{n}^I/I < \epsilon$, and so $\frac{E[e^{Y_\nu \xi}]}{e^{2\epsilon \xi}} \approx \frac{1 + E[Y_\nu]\xi}{1 + 2\epsilon \xi} \leq \frac{1 + \epsilon \xi}{1 + 2\epsilon \xi} < 1$ for small $\xi > 0$. For I large, such that $\epsilon \hat{n}^I > 2$, we then get the following Chernoff upper bound for the RHS of (42):

$$\begin{aligned} \Pr\left(2k + \sum_{\nu=1}^{k\hat{n}^I} Y_\nu > 3\epsilon \hat{n}^I k\right) &\leq \Pr\left(\sum_{\nu=1}^{k\hat{n}^I} Y_\nu > 2\epsilon \hat{n}^I k\right) \\ &\leq \inf_{\xi \geq 0} \left(\frac{E[e^{Y_\nu \xi}]}{e^{2\epsilon \xi}} \right)^{\hat{n}^I k} = p(\epsilon)^{\hat{n}^I k}, \end{aligned}$$

which, together with (42), implies (41) and hence Lemma 7. ■

REFERENCES

- Acemoglu, Daron, Munther A. Dahleh, Ilan Lobel, and Asuman Ozdaglar.** 2011. “Bayesian Learning in Social Networks.” *Review of Economic Studies* 78 (4): 1201–36.
- Bailey, Michael, Drew Johnston, Theresa Kuchler, Johannes Stroebel, and Arlene Wong.** 2022. “Peer Effects in Product Adoption.” *American Economic Journal: Applied Economics* 14 (3): 488–526.
- Bala, Venkatesh, and Sanjeev Goyal.** 1998. “Learning from Neighbours.” *Review of Economic Studies* 65 (3): 595–621.
- Bandiera, Oriana, and Imran Rasul.** 2006. “Social Networks and Technology Adoption in Northern Mozambique.” *Economic Journal* 116 (514): 869–902.
- Banerjee, Abhijit, Arun G. Chandrasekhar, Esther Duflo, and Matthew O. Jackson.** 2013. “The Diffusion of Microfinance.” *Science* 341 (6144): 363.
- Banerjee, Abhijit, Emily Breza, Arun G. Chandrasekhar, and Markus Mobius.** 2021. “Naive Learning with Uninformed Agents.” *American Economic Review* 111 (11): 3540–74.
- Beaman, Lori, Ariel BenYishay, Jeremy Magruder, and Ahmed Mushfiq Mobarak.** 2021. “Can Network Theory-based Targeting Increase Technology Adoption?” *American Economic Review* 111 (6): 1918–43.
- BenYishay, Ariel, and A. Mushfiq Mobarak.** 2019. “Social Learning and Incentives for Experimentation and Communication.” *Review of Economic Studies* 86 (3): 976–1009.
- Besley, Timothy, and Anne Case.** 1994. “Diffusion as a Learning Process: Evidence from HYV Cotton.” Princeton University, Wilson School of Public and International Affairs Discussion Paper 174.
- Board, Simon, and Moritz Meyer-ter-Vehn.** 2021. “Learning Dynamics in Social Networks.” *Econometrica* 89 (6): 2601–35.
- Board, Simon, and Moritz Meyer-ter-Vehn.** 2024. “Simple Contagion Dynamics in Tree-Like Networks.” https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4908830.
- Bonatti, Alessandro, and Johannes Hörner.** 2011. “Collaborating.” *American Economic Review* 101 (2): 632–63.
- Bonatti, Alessandro, and Johannes Hörner.** 2017. “Learning to Disagree in a Game of Experimentation.” *Journal of Economic Theory* 169: 234–69.

- Bramoullé, Yann, and Rachel Kranton.** 2007. "Public Goods in Networks." *Journal of Economic Theory* 135 (1): 478–94.
- Camargo, Braz.** 2014. "Learning in Society." *Games and Economic Behavior* 87: 381–96.
- Coleman, James, Elihu Katz, and Herbert Menzel.** 1957. "The Diffusion of an Innovation among Physicians." *Sociometry* 20 (4): 253–70.
- Conley, Timothy G., and Christopher R. Udry.** 2010. "Learning about a New Technology: Pineapple in Ghana." *American Economic Review* 100 (1): 35–69.
- Duffie, Darrell, Semyon Malamud, and Gustavo Manso.** 2014. "Information Percolation with Equilibrium Search Dynamics." *Journal of Economic Theory* 153: 1–32.
- Dupas, Pascaline.** 2014. "Short-run Subsidies and Long-run Adoption of New Health Products: Evidence from a Field Experiment." *Econometrica* 82 (1): 197–228.
- Fetter, T. Robert, Andrew L. Steck, Christopher Timmins, and Douglas Wrenn.** 2020. "Learning by Viewing? Social Learning, Regulatory Disclosure, and Firm Productivity in Shale Gas." NBER Working Paper 25401.
- Foster, Andrew D., and Mark R. Rosenzweig.** 1995. "Learning by Doing and Learning from Others: Human Capital and Technical Change in Agriculture." *Journal of Political Economy* 103 (6): 1176–1209.
- Gale, Douglas, and Shachar Kariv.** 2003. "Bayesian Learning in Social Networks." *Games and Economic Behavior* 45 (2): 329–46.
- Galeotti, Andrea, and Sanjeev Goyal.** 2010. "The Law of the Few." *American Economic Review* 100 (4): 1468–92.
- Goolsbee, Austan, and Peter J. Klenow.** 2002. "Evidence on Learning and Network Externalities in the Diffusion of Home Computers." *Journal of Law and Economics* 45 (2): 317–43.
- Griliches, Zvi.** 1957. "Hybrid Corn: An Exploration in the Economics of Technological Change." *Econometrica* 25: 501–22.
- Grossman, Sanford J., and Joseph E. Stiglitz.** 1980. "On the Impossibility of Informationally Efficient Markets." *American Economic Review* 70 (3): 393–408.
- Hayek, F. A.** 1945. "The Use of Knowledge in Society." *American Economic Review* 35 (4): 519–30.
- Hodgson, Charles.** 2021. "Information Externalities, Free Riding, and Optimal Exploration in the UK Oil Industry." Unpublished.
- Jackson, Matthew O.** 2010. *Social and Economic Networks*. Princeton, NJ: Princeton University Press.
- Keller, Godfrey, Sven Rady, and Martin Cripps.** 2005. "Strategic Experimentation with Exponential Bandits." *Econometrica* 73 (1): 39–68.
- Koenig, Michael, Claudio Tessone, and Yves Zenou.** 2014. "Nestedness in Networks: A Theoretical Model and Some Applications." *Theoretical Economics* 9 (3): 695–752.
- Kremer, Michael, and Edward Miguel.** 2007. "The Illusion of Sustainability." *Quarterly Journal of Economics* 122 (3): 1007–65.
- Li, Dan, and Norman Schürhoff.** 2019. "Dealer Networks." *Journal of Finance* 74 (1): 91–144.
- Manso, Gustavo, and Farzad Pourbabaee.** 2023. "The Impact of Connectivity on the Production and Diffusion of Knowledge." <https://arxiv.org/abs/2202.00729>.
- Milgram, Stanley.** 1967. "The Small World Problem." *Psychology Today* 2 (1): 60–67.
- Mokyr, Joel.** 1992. *The Lever of Riches: Technological Creativity and Economic Progress*. Oxford: Oxford University Press.
- Moretti, Enrico.** 2011. "Social Learning and Peer Effects in Consumption: Evidence from Movie Sales." *Review of Economic Studies* 78 (1): 356–93.
- Mossel, Elchanan, Manuel Mueller-Frank, Allan Sly, and Omer Tamuz.** 2020. "Social Learning Equilibria." *Econometrica* 88 (3): 1235–67.
- Mossel, Elchanan, Allan Sly, and Omer Tamuz.** 2015. "Strategic Learning and the Topology of Social Networks." *Econometrica* 83 (5): 1755–94.
- Munshi, Kaivan.** 2004. "Social Learning in a Heterogeneous Population: Technology Diffusion in the Indian Green Revolution." *Journal of Development Economics* 73 (1): 185–213.
- Rosenberg, Dinah, Eilon Solan, and Nicolas Vieille.** 2009. "Informational Externalities and Emergence of Consensus." *Games and Economic Behavior* 66 (2): 979–94.
- Sadler, Evan.** 2020a. "Diffusion Games." *American Economic Review* 110 (1): 225–70.
- Sadler, Evan.** 2020b. "Innovation Adoption and Collective Experimentation." *Games and Economic Behavior* 120: 121–31.
- Salish, Mirjam.** 2015. "Learning Faster or More Precisely? Strategic Experimentation in Networks." https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2690997.