

A Communication-Efficient Algorithm for Exponentially Fast Non-Bayesian Learning in Networks

Aritra Mitra, John A. Richards, and Shreyas Sundaram

Abstract—We introduce a simple time-triggered protocol to achieve communication-efficient non-Bayesian learning over a network. Specifically, we consider a scenario where a group of agents interact over a graph with the aim of discerning the true state of the world that generates their joint observation profiles. To address this problem, we propose a novel distributed learning rule wherein agents aggregate neighboring beliefs based on a min-protocol, and the inter-communication intervals grow geometrically at a rate $a \geq 1$. Despite such sparse communication, we show that each agent is still able to rule out every false hypothesis exponentially fast with probability 1, as long as a is finite. For the special case when communication occurs at every time-step, i.e., when $a = 1$, we prove that the asymptotic learning rates resulting from our algorithm are network-structure independent, and a strict improvement over existing rates. In contrast, when $a > 1$, our analysis reveals that the asymptotic learning rates vary across agents, and exhibit a non-trivial dependence on the network topology and the relative entropies of the agents' likelihood models. This motivates us to consider the problem of allocating signal structures to agents to maximize appropriate performance metrics. In certain special cases, we show that the eccentricity centrality and the decay centrality of the underlying graph help identify optimal allocations; for more general cases, we bound the deviation from the optimal allocation as a function of the parameter a , and the diameter of the graph.

I. INTRODUCTION

A typical problem in networked systems involves a global task that needs to be accomplished by a group of entities or agents that are only partially informed about the state of the system. As such, inter-agent communication becomes indispensable for achieving the common goal. Given this premise, it is natural to ask: how frequently must the agents communicate to solve the desired problem? Owing to its practical relevance, the question posed above has received significant recent interest by the control system, information theory and machine learning communities in the context of a variety of problems, namely average consensus [1], optimization [2]–[4], and static parameter estimation [5]. Our goal in this paper is to extend such investigations to the problem of non-Bayesian learning in a network where

the global task involves learning the true state of the world (among a finite set of hypotheses) that explains the private observations of each agent in the network [6]–[11]. Two notable features specific to this problem are as follows. Unlike consensus or distributed optimization, agents are privy to exogenous signals, which, if informative, can enable them to eliminate a subset of the false hypotheses exponentially fast. A related problem where agents receive exogenous signals (measurements) is that of distributed state estimation [12], where the global task entails tracking potentially unstable dynamics. In contrast, the true state remains fixed over time in our setting, thereby simplifying the objective. These attributes play in favor of the problem at hand, motivating us to ask the following questions. (i) Can we design an algorithm that enables each agent to learn the truth with sparse communication schedules (and in fact, even sparser than typically employed for other classes of distributed problems)? (ii) If so, how fast do the agents learn the truth? (iii) Can we quantify the trade-off(s) between sparsity in communication and the rate of learning? We believe that these questions remain largely unexplored. In this context, our **contributions** are as follows.

We develop and analyze a simple time-triggered learning algorithm that builds on our recent work [11], where the data-aggregation step involves a min-protocol as opposed to the consensus-based averaging schemes intrinsic to existing linear [6], [7] and log-linear [8]–[10] learning rules. The basic strategy we employ to achieve communication-efficiency is in line with those in [1], [2], [5], where inter-agent communications become progressively sparser as time evolves. Whereas [1], [2] explore deterministic rules where the inter-communication intervals grow logarithmically and polynomially in time, respectively, [5] analyzes a rule where at each time-step, an agent communicates with its neighbors with a probability that decays to zero sub-linearly. In essence, these approaches establish that as long as the inter-communication intervals do not grow too fast, the global task can still be achieved. We depart from these approaches by allowing the inter-communication intervals to grow much faster: at a geometric rate $a \geq 1$, where the parameter a can be adjusted to control the frequency of communication. We show that our simple time-triggered protocol yields strong guarantees: we prove that even with an arbitrarily large a (leading to a highly sparse communication schedule), each agent is still able to learn the truth exponentially fast with probability 1, provided a is finite. Furthermore, we characterize the dependence of the limiting error exponents on the parameter a , thereby quantifying the trade-offs between

A. Mitra, and S. Sundaram are with the School of Electrical and Computer Engineering at Purdue University. J. A. Richards is with Sandia National Laboratories. Email: {mitra14, sundara2}@purdue.edu, jaricha@sandia.gov. This work was supported in part by NSF CAREER award 1653648, and by a grant from Sandia National Laboratories. Sandia National Laboratories is a multimission laboratory managed and operated by National Technology & Engineering Solutions of Sandia, LLC, a wholly owned subsidiary of Honeywell International Inc., for the U.S. Department of Energy's National Nuclear Security Administration under contract DE-NA0003525. The views expressed in the article do not necessarily represent the views of the U.S. Department of Energy or the United States Government.

communication-efficiency and the speed of learning.

Our analysis subsumes the special case when communication occurs at every time-step, i.e., when $a = 1$; this corresponds to the scenario studied in our previous work [11] which lacked a convergence rate analysis. A significant contribution of this paper is to fill this gap by establishing that when $a = 1$, *the asymptotic learning rates resulting from our proposed algorithm are network-structure independent, and a strict improvement over those existing in the literature.* In contrast, when $a > 1$, we show that the asymptotic learning rates differ from agent to agent, and depend not only on the relative entropies of the agents' signal models, but also on properties of the underlying network. Given this result, we introduce two new measures of the quality of learning, and study the problem of allocating signal structures to agents to maximize such measures. In certain special cases, we show that the eccentricity centrality and the decay centrality of the communication network play key roles in identifying the optimal allocations. For more general cases, we bound the deviation from the optimal allocation as a function of the parameter a , and the diameter of the graph.

II. MODEL AND PROBLEM FORMULATION

Network Model: We consider a group of agents $\mathcal{V} = \{1, \dots, n\}$ that interact with each other over a directed graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ at certain specific time-steps (to be decided by a time-triggered communication schedule). An edge $(i, j) \in \mathcal{E}$ indicates that agent i can directly transmit information to agent j ; the set of all neighbors of agent i is defined as $\mathcal{N}_i = \{j \in \mathcal{V} : (j, i) \in \mathcal{E}\}$. For a strongly-connected graph $\mathcal{G}$, we will use $d(i, j)$ to denote the length of the shortest path from i to j , and $\bar{d}(\mathcal{G})$ to denote the diameter of $\mathcal{G}$.

Observation Model: Let $\Theta = \{\theta_1, \theta_2, \dots, \theta_m\}$ denote m possible states of the world, with each state representing a hypothesis. A specific state $\theta^* \in \Theta$, referred to as the true state of the world, gets realized. Conditional on its realization, at each time-step $t \in \mathbb{N}_+$, every agent $i \in \mathcal{V}$ privately observes a signal $s_{i,t} \in \mathcal{S}_i$, where $\mathcal{S}_i$ denotes the signal space of agent i .¹ The joint observation profile so generated across the network is denoted $s_t = (s_{1,t}, s_{2,t}, \dots, s_{n,t})$, where $s_t \in \mathcal{S}$, and $\mathcal{S} = \mathcal{S}_1 \times \mathcal{S}_2 \times \dots \times \mathcal{S}_n$. Specifically, the signal s_t is generated based on a conditional likelihood function $l(\cdot|\theta^*)$, the i -th marginal of which is denoted $l_i(\cdot|\theta^*)$, and is available to agent i . The signal structure of each agent $i \in \mathcal{V}$ is thus characterized by a family of parameterized marginals $l_i = \{l_i(w_i|\theta) : \theta \in \Theta, w_i \in \mathcal{S}_i\}$. We make certain standard assumptions [6]–[10]: (i) The signal space of each agent i , namely $\mathcal{S}_i$, is finite. (ii) Each agent i has knowledge of its local likelihood functions $\{l_i(\cdot|\theta_p)\}_{p=1}^m$, and it holds that $l_i(w_i|\theta) > 0, \forall w_i \in \mathcal{S}_i$, and $\forall \theta \in \Theta$. (iii) The observation sequence of each agent is described by an i.i.d. random process over time; however, at any given time-step, the observations of different agents may potentially be correlated. (iv) There exists a fixed true state of the world $\theta^* \in \Theta$ (unknown to

the agents) that generates the observations of all the agents. The probability space for our model is denoted $(\Omega, \mathcal{F}, \mathbb{P}^{\theta^*})$, where $\Omega \triangleq \{\omega : \omega = (s_1, s_2, \dots), \forall s_t \in \mathcal{S}, \forall t \in \mathbb{N}_+\}$, $\mathcal{F}$ is the σ -algebra generated by the observation profiles, and $\mathbb{P}^{\theta^*}$ is the probability measure induced by sample paths in Ω . Specifically, $\mathbb{P}^{\theta^*} = \prod_{t=1}^{\infty} l(\cdot|\theta^*)$. We will use the abbreviation

a.s. to indicate almost sure occurrence of an event w.r.t. $\mathbb{P}^{\theta^*}$.

Given the above setting, the goal of each agent in the network is to eventually learn the true state θ^* - a task that may be impossible for any agent to achieve in isolation. Specifically, let $\Theta_i^{\theta^*} \triangleq \{\theta \in \Theta : l_i(w_i|\theta) = l_i(w_i|\theta^*), \forall w_i \in \mathcal{S}_i\}$ represent the set of hypotheses that are *observationally equivalent* to θ^* from the perspective of agent i . An agent i is deemed partially informative about the truth if $|\Theta_i^{\theta^*}| > 1$. Since potentially every agent can be partially informative in the sense described above, inter-agent communications become necessary for each agent to learn the truth.

In this context, our **objectives** in this paper are to develop an understanding of (i) the amount of leeway that the above problem affords in terms of sparsifying inter-agent communications without compromising on the objective of learning the truth; and (ii) the trade-offs between sparse communication and the rate of learning. To this end, we recall the following definition from [11].

Definition 1. (Source agents) An agent i is said to be a source agent for a pair of distinct hypotheses $\theta_p, \theta_q \in \Theta$ if it can distinguish between them, i.e., if $D(l_i(\cdot|\theta_p)||l_i(\cdot|\theta_q)) > 0$, where $D(l_i(\cdot|\theta_p)||l_i(\cdot|\theta_q))$ represents the KL-divergence [13] between the distributions $l_i(\cdot|\theta_p)$ and $l_i(\cdot|\theta_q)$. The set of source agents for pair (θ_p, θ_q) is denoted $\mathcal{S}(\theta_p, \theta_q)$. $\square$

Throughout the rest of the paper, we will use $K_i(\theta_p, \theta_q)$ as a shorthand for $D(l_i(\cdot|\theta_p)||l_i(\cdot|\theta_q))$.

III. A COMMUNICATION-EFFICIENT LEARNING RULE

In this section, we formally introduce a simple time-triggered belief update rule parameterized by a constant $a \in \mathbb{N}_+$ that determines the frequency of communication (to be made more precise below). Every agent i maintains a local belief vector $\pi_{i,t}$, and an actual belief vector $\mu_{i,t}$, each of which are probability distributions over the hypothesis set Θ . These vectors are initialized with $\pi_{i,0}(\theta) > 0, \mu_{i,0}(\theta) > 0, \forall \theta \in \Theta, \forall i \in \mathcal{V}$, and subsequently updated as follows.

- **Update of the local beliefs:** At each time-step $t+1 \in \mathbb{N}_+$, $\pi_{i,t+1}$ is updated via a standard Bayesian rule:

$$\pi_{i,t+1}(\theta) = \frac{l_i(s_{i,t+1}|\theta)\pi_{i,t}(\theta)}{\sum_{p=1}^m l_i(s_{i,t+1}|\theta_p)\pi_{i,t}(\theta_p)}. \quad (1)$$

- **Update of the actual beliefs:** Let $\mathbb{I} = \{t_k\}_{k \in \mathbb{N}_+}$ be a sequence of time-steps satisfying $t_{k+1} - t_k = a^k, \forall k \in \mathbb{N}_+$, with $t_1 = 1$. If $t+1 \in \mathbb{I}$, then $\mu_{i,t+1}$ is updated as

$$\mu_{i,t+1}(\theta) = \frac{\min\{\{\mu_{j,t}(\theta)\}_{j \in \mathcal{N}_i, \pi_{i,t+1}(\theta)}\}}{\sum_{p=1}^m \min\{\{\mu_{j,t}(\theta_p)\}_{j \in \mathcal{N}_i, \pi_{i,t+1}(\theta_p)}\}}. \quad (2)$$

If $t+1 \notin \mathbb{I}$, $\mu_{i,t+1}$ is simply held constant as follows:

$$\mu_{i,t+1}(\theta) = \mu_{i,t}(\theta). \quad (3)$$

¹We use $\mathbb{N}$ and $\mathbb{N}_+$ to represent the set of non-negative integers and positive integers, respectively.

In words, while the local beliefs are updated at every time-step, the actual beliefs are updated only at time-steps that belong to the set $\mathbb{I}$, i.e., an agent $i \in \mathcal{V}$ is allowed to transmit $\mu_{i,t}$ to its out-neighbors, and receive $\mu_{j,t}$ from each in-neighbor j in $\mathcal{G}$ if and only if $t+1 \in \mathbb{I}$. When $a = 1$, the actual beliefs get updated via (2) at every time-step, and we recover the rule in [11]. When $a > 1$, note that the inter-communication intervals grow exponentially at a rate dictated by the parameter a . Prior to analyzing the impact of such sparse communication, a few comments are in order.

First, notice that the data-aggregation rule in (2) is based on a min-protocol, as opposed to any form of “belief-averaging” commonly employed in the distributed learning literature [6]–[10]. Essentially, while the local belief updates (1) capture what an agent can learn by itself, the actual belief updates (2) incorporate information from the rest of the network. Second, we note that the proposed time-triggered protocol is simple, easy to implement, and computationally cheap. At the same time, the exponentially growing intervals afford a much sparser communication schedule relative to related literature. Third, while one can potentially consider extensions of this algorithm that account for asynchronicity, communication failures, delays etc., we focus on the scheme here in order to concretely isolate the trade-off between sparse communication and the asymptotic rates of learning, and highlight how the network structure impacts such rates.

IV. MAIN RESULT AND DISCUSSION

The main result of the paper is as follows (the proof will be provided in the next section).

Theorem 1. *Suppose the communication parameter satisfies $a > 1$, and the following hold: (i) for every pair $\theta_p, \theta_q \in \Theta$, $S(\theta_p, \theta_q) \neq \emptyset$; and (ii) $\mathcal{G}$ is strongly-connected. Then, the distributed learning rule given by (1), (2), (3) guarantees:*

- **(Consistency):** *For each agent $i \in \mathcal{V}$, $\mu_{i,t}(\theta^*) \rightarrow 1$ a.s.*
- **(Asymptotic Rate of Rejection of False Hypotheses):** *The following holds for each $i \in \mathcal{V}$, and $\theta \in \Theta \setminus \{\theta^*\}$:*

$$\liminf_{t \rightarrow \infty} -\frac{\log \mu_{i,t}(\theta)}{t} \geq \max_{v \in S(\theta^*, \theta)} \frac{K_v(\theta^*, \theta)}{a^{(d(v,i)+1)}} \text{ a.s.} \quad (4)$$

□

We immediately obtain the following important corollary.

Corollary 1. *Suppose communication occurs at every time-step, i.e., suppose $a = 1$. Let the conditions in Theorem 1 hold. Then, the proposed learning rule guarantees consistency, and the following holds $\forall i \in \mathcal{V}$, and $\theta \in \Theta \setminus \{\theta^*\}$:*

$$\liminf_{t \rightarrow \infty} -\frac{\log \mu_{i,t}(\theta)}{t} \geq \max_{v \in S(\theta^*, \theta)} K_v(\theta^*, \theta) \text{ a.s.} \quad (5)$$

□

Implications of Theorem 1: We first note that despite its simplicity, the algorithm proposed in Section III provides strong guarantees: eq. (4) indicates that although the inter-communication intervals grow exponentially at an arbitrarily large (but finite) rate a , each agent is still able to eliminate every false hypothesis exponentially fast with probability 1. More interestingly, (4) reveals that unlike existing literature

[6]–[10], the asymptotic learning rates are *agent-specific*, i.e., different agents may discover the truth at different rates.² In particular, note from the RHS of (4) that, when considering the asymptotic rate of rejection of a particular false hypothesis at a given agent i , one needs to account for the attenuated relative entropies of the corresponding source agents, where the attenuation factor scales exponentially with the distances of i from such source agents. It is easy to see from (4) that sparser communication schedules (corresponding to larger a ’s) incur lower learning rates. Moreover, since such rates depend upon the network-structure when $a > 1$, a poor allocation of signal structures to agents can have adverse effects on the learning rates of certain agents.

Implications of Corollary 1: Let us now compare the performance of our algorithm with that of existing “belief-averaging” schemes [6]–[10] when communication occurs at every time-step, i.e., when $a = 1$, which is the standard distributed hypothesis testing setup. In sharp contrast to when $a > 1$, Corollary 1 indicates that the asymptotic learning rates are *network-structure independent*, and *identical* for each agent. Moreover, under the conditions on the observation model and the network structure given in Thm. 1, both linear [6], [7] and log-linear [8]–[10] opinion pooling lead to an asymptotic rate of rejection of the form $\sum_{i \in \mathcal{V}} \nu_i K_i(\theta^*, \theta)$ for each $\theta \in \Theta \setminus \{\theta^*\}$, and the rate is identical for each agent. Here, ν_i represents the eigenvector centrality of agent i , which is strictly positive for a strongly-connected graph. Thus, referring to (5), we conclude that a significant contribution of the algorithm proposed in this paper is that it yields *strictly better* asymptotic learning rates than those existing, for the standard setting when $a = 1$.

V. PROOF OF THE MAIN RESULT

In order to prove Theorem 1, we require a few intermediate results, the first of which is quite standard (see [11]).

Lemma 1. *Consider a false hypothesis $\theta \in \Theta \setminus \{\theta^*\}$, and an agent $i \in S(\theta^*, \theta)$. Suppose $\pi_{i,0}(\theta_p) > 0, \forall \theta_p \in \Theta$. Then, the update rule (1) ensures that (i) $\pi_{i,t}(\theta) \rightarrow 0$ a.s., (ii) $\pi_{i,\infty}(\theta^*) \triangleq \lim_{t \rightarrow \infty} \pi_{i,t}(\theta^*)$ exists a.s. and satisfies $\pi_{i,\infty}(\theta^*) \geq \pi_{i,0}(\theta^*)$, and (iii) the following holds:*

$$\lim_{t \rightarrow \infty} \frac{1}{t} \log \frac{\pi_{i,t}(\theta)}{\pi_{i,t}(\theta^*)} = -K_i(\theta^*, \theta) \text{ a.s.} \quad (6)$$

□

Lemma 2. *Suppose the conditions in Theorem 1 hold, and each agent employs the algorithm given by (1), (2), and (3). Then, there exists a set $\bar{\Omega} \subseteq \Omega$ with the following properties: (i) $\mathbb{P}^{\theta^*}(\bar{\Omega}) = 1$, and (ii) for each $\omega \in \bar{\Omega}$, there exist constants $\eta(\omega) \in (0, 1)$ and $t'(\omega) \in (0, \infty)$ such that*

$$\pi_{i,t}(\theta^*) \geq \eta(\omega), \mu_{i,t}(\theta^*) \geq \eta(\omega), \forall t \geq t'(\omega), \forall i \in \mathcal{V}. \quad (7)$$

□

Proof. Let $\bar{\Omega} \subseteq \Omega$ denote the set of sample paths for which the assertions in Lemma 1 hold for each false hypothesis $\theta \in$

²We use the lower bounds derived in (4), (5) as a proxy when referring to the corresponding asymptotic learning rates.

$\Theta \setminus \{\theta^*\}$. Based on Lemma 1, we note that $\mathbb{P}^{\theta^*}(\bar{\Omega}) = 1$. Thus, it suffices to establish (7) for each sample path $\omega \in \bar{\Omega}$. To this end, fix $\omega \in \bar{\Omega}$. Following similar arguments as in [11], one can find $\eta(\omega) \in (0, 1)$ and $t'(\omega) \in (0, \infty)$, such that $\forall i \in \mathcal{V}$, $\pi_{i,t}(\theta^*) \geq \eta(\omega)$, $\forall t \geq t'(\omega)$, and $\mu_{i,t'}(\theta^*) \geq \eta(\omega)$. We thus focus only on establishing that $\mu_{i,t}(\theta^*) \geq \eta(\omega)$, $\forall t > t'(\omega)$, $\forall i \in \mathcal{V}$. To this end, let $\bar{t}(\omega) > t'(\omega)$ be the first time-step following $t'(\omega)$ that belongs to the set $\mathbb{I}$. Based on (3), notice that $\mu_{i,t}(\theta^*) \geq \eta(\omega)$ for all $t \in [t'(\omega), \bar{t}(\omega))$, and for each $i \in \mathcal{V}$. Based on (2), at time-step $\bar{t}(\omega) \in \mathbb{I}$, $\mu_{i,\bar{t}(\omega)}(\theta^*)$ for an agent $i \in \mathcal{V}$ satisfies:

$$\begin{aligned} \mu_{i,\bar{t}(\omega)}(\theta^*) &\geq \frac{\eta(\omega)}{\sum_{p=1}^m \min\{\{\mu_{j,\bar{t}(\omega)-1}(\theta_p)\}_{j \in \mathcal{N}_i}, \pi_{i,\bar{t}(\omega)}(\theta_p)\}} \\ &\geq \frac{\eta(\omega)}{\sum_{p=1}^m \pi_{i,\bar{t}(\omega)}(\theta_p)} = \eta(\omega), \end{aligned} \quad (8)$$

where the last equality follows from the fact that the local belief vectors generated via (1) are valid probability distributions over Θ at each time-step, and hence $\sum_{p=1}^m \pi_{i,\bar{t}(\omega)}(\theta_p) = 1$. The above argument applies identically to each agent in $\mathcal{V}$. Furthermore, it is easily seen that based on (3), and a similar reasoning as above, identical conclusions can be drawn for each time-step $t > t'(\omega)$, $t \in \mathbb{I}$ when the agents update their actual beliefs based on (2). This readily establishes (7). $\square$

Lemma 3. Consider a false hypothesis $\theta \in \Theta \setminus \{\theta^*\}$ and an agent $v \in \mathcal{S}(\theta^*, \theta)$. Suppose the conditions stated in Theorem 1 hold. Then, the learning rule described by equations (1), (2) and (3) guarantee the following for each agent $i \in \mathcal{V}$:

$$\liminf_{t \rightarrow \infty} -\frac{\log \mu_{i,t}(\theta)}{t} \geq \frac{K_v(\theta^*, \theta)}{a^{(d(v,i)+1)}} \text{ a.s.} \quad (9)$$

$\square$

Proof. Throughout this proof, we use the same notation as in Lemma 2. Fix an $\omega \in \bar{\Omega}$, an agent $v \in \mathcal{S}(\theta^*, \theta)$, and an agent $i \in \mathcal{V}$. Since condition (ii) in Thm. 1 is met, there exists a path of shortest length from v to i ; we shall induct on the length of such a path. First, we consider the base case when $d(v, i) = 0$, i.e., when $i = v$. Fix $\epsilon > 0$, and notice that since $v \in \mathcal{S}(\theta^*, \theta)$, Eq. (6) implies that $\exists t_v(\omega, \theta, \epsilon)$, such that:

$$\pi_{v,t}(\theta) < e^{-(K_v(\theta^*, \theta) - \epsilon)t}, \forall t \geq t_v(\omega, \theta, \epsilon). \quad (10)$$

Since $\omega \in \bar{\Omega}$, Lemma 2 guarantees the existence of a time-step $t'(\omega) < \infty$, and a constant $\eta(\omega) > 0$, such that on ω , $\pi_{i,t}(\theta^*) \geq \eta(\omega)$, $\mu_{i,t}(\theta^*) \geq \eta(\omega)$, $\forall t \geq t'(\omega)$, $\forall i \in \mathcal{V}$. Let $\bar{t}_v(\omega, \theta, \epsilon) = \max\{t'(\omega), t_v(\omega, \theta, \epsilon)\}$. Let $\tilde{t} > \bar{t}_v$ be the first time-step following $\bar{t}_v$ that belongs to $\mathbb{I}$.³ Then, Eq. (2) leads to the following inequalities:

$$\begin{aligned} \mu_{v,\tilde{t}}(\theta) &\stackrel{(a)}{\leq} \frac{\pi_{v,\tilde{t}}(\theta)}{\sum_{p=1}^m \min\{\{\mu_{j,\tilde{t}-1}(\theta_p)\}_{j \in \mathcal{N}_i}, \pi_{v,\tilde{t}}(\theta_p)\}} \\ &\stackrel{(b)}{<} \frac{e^{-(K_v(\theta^*, \theta) - \epsilon)\tilde{t}}}{\eta(\omega)} = C(\omega)e^{-(K_v(\theta^*, \theta) - \epsilon)\tilde{t}}, \end{aligned} \quad (11)$$

³We suppress the dependence of various quantities on the parameters ω, θ , and ϵ for brevity.

where $C(\omega) = \eta(\omega)^{-1}$. Regarding the inequalities in (11), (a) follows directly from (2), whereas (b) follows from (10) and the fact that $\eta(\omega)$ lower bounds the beliefs (both local and actual) of all agents on the true state θ^* . Note that consecutive trigger-points $t_k, t_{k+1} \in \mathbb{I}$ satisfy $t_{k+1} = at_k + 1$. Based on (3), we then have:

$$\mu_{v,t}(\theta) < C(\omega)e^{-(K_v(\theta^*, \theta) - \epsilon)\tilde{t}}, \forall t \in [\tilde{t}, a\tilde{t} + 1]. \quad (12)$$

Observe that the next update of $\mu_{v,t}(\theta)$ takes place at $a\tilde{t} + 1$. Repeating the reasoning used to arrive at (11) yields:

$$\mu_{v,a\tilde{t}+1}(\theta) < C(\omega)e^{-(K_v(\theta^*, \theta) - \epsilon)(a\tilde{t}+1)}. \quad (13)$$

Coupled with the above inequality, (3) once again implies:

$$\mu_{v,t}(\theta) < C(\omega)e^{-(K_v(\theta^*, \theta) - \epsilon)(a\tilde{t}+1)}, \forall t \in [a\tilde{t}+1, a^2\tilde{t}+a+1]. \quad (14)$$

Generalizing the above reasoning, we obtain:

$$\mu_{v,t}(\theta) < C(\omega)e^{-(K_v(\theta^*, \theta) - \epsilon)(a^p\tilde{t} + f(p))}, \quad (15)$$

$\forall t \in [a^p\tilde{t} + f(p), a^{(p+1)}\tilde{t} + af(p) + 1]$, $p \in \mathbb{N}$, where $f(p) = (a^p - 1)/(a - 1)$. We conclude that $\forall t \geq \tilde{t}$:

$$\mu_{v,t}(\theta) < C(\omega)e^{-(K_v(\theta^*, \theta) - \epsilon)(a^{p(t)}\tilde{t} + f(p(t)))}, \quad (16)$$

where

$$p(t) = \lfloor g(t) \rfloor, \quad g(t) = \frac{\log \frac{(a-1)t+1}{(a-1)\tilde{t}+1}}{\log a}. \quad (17)$$

From (16), we obtain that $\forall t \geq \tilde{t}$:

$$-\frac{\log \mu_{v,t}(\theta)}{t} > \frac{(K_v(\theta^*, \theta) - \epsilon)(a^{p(t)}\tilde{t} + f(p(t)))}{t} - \frac{\log C(\omega)}{t}. \quad (18)$$

Let $\alpha_v(\theta, \epsilon) = (K_v(\theta^*, \theta) - \epsilon)$. Then, taking the limit inferior on both sides of the above inequality yields:

$$\begin{aligned} \liminf_{t \rightarrow \infty} -\frac{\log \mu_{v,t}(\theta)}{t} &\geq \alpha_v(\theta, \epsilon) \lim_{t \rightarrow \infty} \frac{1}{t} \left[a^{p(t)}\tilde{t} + \frac{a^{p(t)} - 1}{a - 1} \right] \\ &\geq \frac{\alpha_v(\theta, \epsilon)}{a} \lim_{t \rightarrow \infty} \frac{1}{t} \left[a^{g(t)}(\tilde{t} + \frac{1}{a - 1}) \right] \\ &= \frac{\alpha_v(\theta, \epsilon)}{a}, \end{aligned} \quad (19)$$

where the second inequality follows from the fact that $\lfloor x \rfloor > x - 1$, $\forall x \in \mathbb{R}$, and the final equality results from further simplifications based on (17). Finally, noting that ϵ can be made arbitrarily small in the above inequality establishes (9) for the base case when $d(v, i) = 0$. To proceed, suppose (9) holds for each node $i \in \mathcal{V}$ satisfying $0 \leq d(v, i) \leq q$, where q is a non-negative integer satisfying $q \leq \bar{d}(\mathcal{G}) - 1$ (recall that $\bar{d}(\mathcal{G})$ represents the diameter of the graph $\mathcal{G}$). Let $i \in \mathcal{V}$ be such that $d(v, i) = q + 1$. Thus, there must exist some node $l \in \mathcal{N}_i$ such that $d(v, l) = q$. The induction hypothesis applies to this node l , and hence, we have:

$$\liminf_{t \rightarrow \infty} -\frac{\log \mu_{l,t}(\theta)}{t} \geq \frac{K_v(\theta^*, \theta)}{a^{(q+1)}} \text{ a.s.} \quad (20)$$

Let $\tilde{\Omega} \subseteq \bar{\Omega}$ be the set of sample paths for which the above inequality holds. With $\bar{\Omega}$ defined as before, notice that $\mathbb{P}^{\theta^*}(\tilde{\Omega} \cap \bar{\Omega}) = 1$, since $\tilde{\Omega}$ and $\bar{\Omega}$ each have $\mathbb{P}^{\theta^*}$ -measure 1.

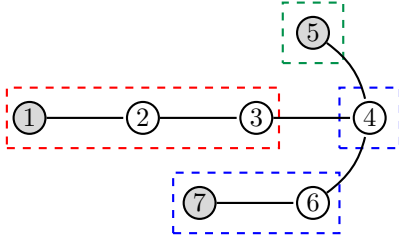


Fig. 1. The figure represents the network for the simulation example in Section VI. Based on the parameters of the model, Theorem 1 implies that the asymptotic rates of rejection of θ_2 for the agents enclosed in the red, blue and green rectangles are dictated by the relative entropies of agents 1, 7 and 5, respectively, illustrating the agent-specific learning rate phenomenon.

Pick an arbitrary sample path $\omega \in \tilde{\Omega} \cap \bar{\Omega}$, and notice that based on arguments identical to the base case, on the sample path ω there exists a time-step $\bar{t}_l$, such that the beliefs of all agents on θ^* are bounded below by $\eta(\omega)$ following $\bar{t}_l$, and

$$\mu_{l,t}(\theta) < e^{-(H_l(\theta^*, \theta) - \epsilon)t}, \forall t \geq \bar{t}_l, \quad (21)$$

where $\epsilon > 0$ is an arbitrary small number, and $H_l(\theta^*, \theta) = K_v(\theta^*, \theta)/a^{(q+1)}$. Proceeding as in the base case, let $\tau > \bar{t}_l$ be the first time-step following $\bar{t}_l$ that belongs to the set $\mathbb{I}$. Noting that $l \in \mathcal{N}_i$, using (2), (21), and similar arguments as those used to arrive at (11), we obtain:

$$\begin{aligned} \mu_{i,\tau}(\theta) &\leq \frac{\mu_{l,\tau-1}(\theta)}{\sum_{p=1}^m \min\{\{\mu_{j,\tau-1}(\theta_p)\}_{j \in \mathcal{N}_i}, \pi_{i,\tau}(\theta_p)\}} \\ &< \frac{e^{-(H_l(\theta^*, \theta) - \epsilon)(\tau-1)}}{\eta(\omega)} = C_l(\omega) e^{-(H_l(\theta^*, \theta) - \epsilon)\tau}, \end{aligned} \quad (22)$$

where $C_l(\omega) = e^{(H_l(\theta^*, \theta) - \epsilon)}/\eta(\omega)$. Repeating the above analysis for each time-step of the form $a^p\tau + f(p)$, $p \in \mathbb{N}_+$, using (3), and arguments similar to the base case yields:

$$\mu_{i,t}(\theta) < C_l(\omega) e^{-(H_l(\theta^*, \theta) - \epsilon)(a^{\bar{p}(t)}\tau + f(\bar{p}(t)))}, \forall t \geq \tau, \quad (23)$$

where

$$\bar{p}(t) = \lfloor \bar{g}(t) \rfloor, \quad \bar{g}(t) = \frac{\log \frac{(a-1)t+1}{(a-1)\tau+1}}{\log a}. \quad (24)$$

Notice that the inequality in (23) resembles that in (16). Thus, repeating the analysis for the base case yields:

$$\liminf_{t \rightarrow \infty} -\frac{\log \mu_{i,t}(\theta)}{t} \geq \frac{H_l(\theta^*, \theta)}{a} - \frac{\epsilon}{a}. \quad (25)$$

The induction step follows by substituting the expression for $H_l(\theta^*, \theta)$ in (25), and recalling that $d(v, i) = q + 1$. $\square$

We are now in position to prove Theorem 1.

Proof. (Theorem 1) Fix a $\theta \in \Theta \setminus \{\theta^*\}$. Based on condition (i) of the Theorem, $\mathcal{S}(\theta^*, \theta)$ is non-empty, and based on condition (ii), there exists a path from each agent $v \in \mathcal{S}(\theta^*, \theta)$ to every agent in $\mathcal{V} \setminus \{v\}$; Eq. (4) then follows from Lemma 3. By definition of a source set, $K_v(\theta^*, \theta) > 0, \forall v \in \mathcal{S}(\theta^*, \theta)$; (4) then implies $\lim_{t \rightarrow \infty} \mu_{i,t}(\theta) = 0$ a.s., $\forall i \in \mathcal{V}$. $\square$

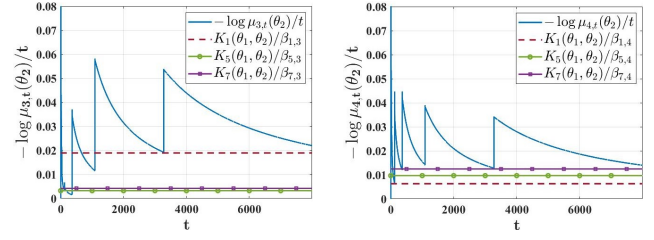


Fig. 2. Plots of the instantaneous rates of decay of the beliefs of agents 3 and 4 on the false hypothesis θ_2 . In the plots, $\beta_{i,j} = a^{(d(i,j)+1)}$.

VI. SIMULATION EXAMPLE

Suppose $\Theta = \{\theta_1, \theta_2\}$, and $\theta^* = \theta_1$. The network of agents is depicted in Fig. 1. The signal space is given by $\mathcal{S}_i = \{1, 2\}, \forall i \in \mathcal{V}$. The agent likelihood models satisfy: $l_i(1|\theta_1) = 0.5, \forall i \in \{1, \dots, 7\}$, $l_1(1|\theta_2) = 0.9$, $l_5(1|\theta_2) = 0.7$, $l_7(1|\theta_2) = 0.85$, and $l_i(1|\theta_2) = 0.5, \forall i \in \{2, 3, 4, 6\}$. Thus, only agents 1, 5 and 7 can distinguish between θ_1 and θ_2 , with their relative entropies satisfying $K_1(\theta_1, \theta_2) > K_7(\theta_1, \theta_2) > K_5(\theta_1, \theta_2) > 0$. Let $a = 3$, and note that $K_7(\theta_1, \theta_2)/a^2 > K_5(\theta_1, \theta_2)/a > K_1(\theta_1, \theta_2)/a^3$. Fig. 2 plots the instantaneous rates of rejection of the false hypothesis θ_2 for agents 3 and 4, resulting from our algorithm. Based on Figs. 1 and 2, we observe: (i) each informative agent dominates the speed of learning of agents that are close to it in $\mathcal{G}$; (ii) the rate of rejection of θ_2 is indeed agent-specific; and (iii) the simulations agree very closely with the lower bounds on the limiting rates of rejection in Thm. 1.

VII. THE IMPACT OF INFORMATION ALLOCATION ON ASYMPTOTIC LEARNING RATES

Thm. 1 indicates that the learning rates of the agents are shaped by a non-trivial interplay between the relative entropies of their signal models and the network structure. In view of this fact, we now aim to analyze how information should be allocated to the agents in order to maximize appropriate performance metrics that are a function of the learning rates. Our investigation is inspired by similar questions in [7]; however, unlike [7], our learning rule leads to learning rates that are agent-dependent when $a > 1$ (see Sec. VI). Thus, the performance metrics we seek to maximize differ from those in [7]. As we shall soon see, while the eigenvector centrality plays a key role in shaping the speed of learning in [7], alternate network centrality measures become important when it comes to the belief dynamics generated by our rule.

Throughout this section, we consider a strongly-connected graph $\mathcal{G}$, and a set of n signal structures $\mathcal{L} = \{l_1, \dots, l_n\}$, where each l_i represents a family of parameterized marginals as defined in Section II. By an allocation of signal structures to agents, we imply a bijection $\psi: \mathcal{L} \rightarrow \mathcal{V}$. Let Ψ represent the set of all possible bijections between $\mathcal{L}$ and $\mathcal{V}$. Our aim is to optimally pick $\psi \in \Psi$ to maximize the performance metrics that we define next. Given a pair of hypotheses $\theta_p, \theta_q \in \Theta$, recall from (4) that based on our learning rule,

$$\rho_i^\psi(\theta_p, \theta_q) \triangleq \max_{v \in \mathcal{S}^\psi(\theta_p, \theta_q)} \frac{K_v^\psi(\theta_p, \theta_q)}{a^{(d(v,i)+1)}} \quad (26)$$

lower bounds the limiting rate at which agent i rules out θ_q , when θ_p is realized as the true state; the superscript ψ reflects the dependence of the corresponding objects on the allocation policy ψ . We now introduce two measures of the quality of learning that are specific to our setting:

$$\rho_{\text{avg}}^{\psi} \triangleq \min_{\theta_p, \theta_q \in \Theta} \frac{1}{n} \sum_{i \in \mathcal{V}} \rho_i^{\psi}(\theta_p, \theta_q), \rho_{\text{min}}^{\psi} \triangleq \min_{\theta_p, \theta_q \in \Theta} \min_{i \in \mathcal{V}} \rho_i^{\psi}(\theta_p, \theta_q). \quad (27)$$

While ρ_{avg}^{ψ} captures the average rate of learning across the network, ρ_{min}^{ψ} focuses on the agent that converges the slowest; since any state in Θ can be realized, these metrics account for the pair of states that is the hardest to distinguish. We seek to maximize ρ_{avg}^{ψ} and ρ_{min}^{ψ} over the set Ψ . Our first result on this topic makes a connection to two popular distance-based network centrality measures, namely, the *eccentricity centrality* ξ_i [14], and the *decay centrality* $\kappa_i(\delta)$ [15], of each agent $i \in \mathcal{V}$, defined as:

$$\xi_i = 1 / (\max_{j \in \mathcal{V} \setminus \{i\}} d(i, j)), \quad \kappa_i(\delta) = \sum_{j \in \mathcal{V} \setminus \{i\}} \delta^{d(i, j)}, \quad (28)$$

where $0 < \delta < 1$ is the decay parameter.

Proposition 1. *Suppose $a > 1$, and that there exists a signal structure $l_u \in \mathcal{L}$ such that for all $\theta_p, \theta_q \in \Theta$,*

$$K_{l_u}(\theta_p, \theta_q) / a^{\bar{d}(\mathcal{G})} > K_{l_w}(\theta_p, \theta_q), \forall l_w \in \mathcal{L} \setminus \{l_u\}. \quad (29)$$

Then, (i) any allocation $\psi \in \Psi$ such that $\psi(l_u) \in \arg\max_{i \in \mathcal{V}} \xi_i$ maximizes ρ_{min}^{ψ} , and (ii) any allocation $\psi \in \Psi$ such that $\psi(l_u) \in \arg\max_{i \in \mathcal{V}} \kappa_i(\frac{1}{a})$ maximizes ρ_{avg}^{ψ} .⁴ $\square$

The intuition behind the above result is as follows. Suppose there exists a signal structure that is sufficiently stronger in its discriminatory power than the others. Then, an agent allocated such a structure will govern the rate of learning of every other agent. To expedite learning, it thus makes sense to allocate such a dominant signal structure to the most central agent in the network, where the specific centrality measure depends on the performance metric. Note that eccentricity centrality and decay centrality have been widely studied in the context of information spread over social and economic networks [16], [17]. While the former bears connections to information cascades [16], the latter facilitates selection of an “implant” node that maximizes diffusion of a certain product or idea over a network [17]. Prop. 1 identifies conditions under which the above centralities have similar implications for the belief dynamics generated by our rule.

While Prop. 1 allows one to identify the optimal allocation by simply computing the appropriate centrality measures, the scenario becomes much more complicated if no additional structure is imposed on the agents’ likelihood models. For such general cases, we provide a coarse upper bound on the sub-optimality of any given allocation.

Proposition 2. *Suppose $a > 1$ and let $\psi_{\alpha}^* \in \Psi$ and $\psi_{\beta}^* \in \Psi$ be allocations that maximize ρ_{min}^{ψ} and ρ_{avg}^{ψ} , respectively.*

⁴Here, the quantity $K_{l_u}(\theta_p, \theta_q)$ should be interpreted differently from $K_u(\theta_p, \theta_q)$; whereas the former indicates a relative entropy associated with the signal structure l_u , the latter indicates a relative entropy associated with agent u once it has been allocated a certain signal structure.

Then, for any allocation $\psi \in \Psi$, the following hold: (i) $\rho_{\text{min}}^{\psi_{\alpha}^} / \rho_{\text{min}}^{\psi} \leq a^{\bar{d}(\mathcal{G})}$; and (ii) $\rho_{\text{avg}}^{\psi_{\beta}^*} / \rho_{\text{avg}}^{\psi} \leq a^{\bar{d}(\mathcal{G})}$. $\square$*

We omit the proofs of Propositions 1 and 2 here due to space constraints; they are available in [18].

VIII. CONCLUSION

We developed and analyzed a simple time-triggered protocol for achieving communication-efficient non-Bayesian learning over a network. Unlike existing approaches, we allowed the inter-communication intervals to grow unbounded over time at an arbitrarily large (but finite) geometric rate $a \geq 1$. We showed that despite such sparse communication, our approach enables each agent to learn the true state exponentially fast with probability 1. We then characterized the limiting error exponents as a function of the parameter a . For the special case when agents interact at every time-step, i.e., when $a = 1$, we proved that our approach yields strictly better asymptotic learning rates than those existing in the literature. Finally, for $a > 1$, we studied the impact of signal allocations on the speed of learning.

REFERENCES

- [1] A. Olshevsky, I. C. Paschalidis, and A. Spiridonoff, “Fully asynchronous push-sum with growing intercommunication intervals,” in *Proc. of the American Control Conference*, 2018, pp. 591–596.
- [2] K. Tsianos, S. Lawlor, and M. G. Rabbat, “Communication/computation tradeoffs in consensus-based distributed optimization,” in *Advances in Neural Info. Proc. systems*, 2012, pp. 1943–1951.
- [3] T. Chen, G. Giannakis, T. Sun, and W. Yin, “Lag: Lazily aggregated gradient for communication-efficient distributed learning,” in *Advances in Neural Info. Proc. Systems*, 2018, pp. 5055–5065.
- [4] G. Lan, S. Lee, and Y. Zhou, “Communication-efficient algorithms for decentralized and stochastic optimization,” *Mathematical Programming*, pp. 1–48, 2017.
- [5] A. K. Sahu, D. Jakovetic, and S. Kar, “Communication optimality trade-offs for distributed estimation,” *arXiv:1801.04050*, 2018.
- [6] A. Jadbabaie, P. Molavi, A. Sandroni, and A. Tahbaz-Salehi, “Non-Bayesian social learning,” *Games and Economic Behavior*, vol. 76, no. 1, pp. 210–225, 2012.
- [7] A. Jadbabaie, P. Molavi, and A. Tahbaz-Salehi, “Information heterogeneity and the speed of learning in social networks,” *Columbia Bus. Sch. Res. Paper*, pp. 13–28, 2013.
- [8] S. Shahrampour, A. Rakhlin, and A. Jadbabaie, “Distributed detection: Finite-time analysis and impact of network topology,” *IEEE Trans. on Autom. Control*, vol. 61, no. 11, pp. 3256–3268, 2016.
- [9] A. Nedic, A. Olshevsky, and C. A. Uribe, “Fast convergence rates for distributed Non-Bayesian learning,” *IEEE Trans. on Autom. Control*, vol. 62, no. 11, pp. 5538–5553, 2017.
- [10] A. Lalitha, T. Javidi, and A. Sarwate, “Social learning and distributed hypothesis testing,” *IEEE Trans. on Info. Theory*, vol. 64, no. 9, 2018.
- [11] A. Mitra, J. A. Richards, and S. Sundaram, “A new approach for distributed hypothesis testing with extensions to Byzantine-resilience,” in *Proc. of the American Control Conference*, 2019, pp. 261–266.
- [12] A. Mitra and S. Sundaram, “Distributed observers for LTI systems,” *IEEE Trans. on Autom. Control*, vol. 63, no. 11, pp. 3689–3704, 2018.
- [13] T. M. Cover and J. A. Thomas, *Elements of information theory*. John Wiley & Sons, 2012.
- [14] P. Hage and F. Harary, “Eccentricity and centrality in networks,” *Social networks*, vol. 17, no. 1, pp. 57–63, 1995.
- [15] N. Tsakas, “On decay centrality,” *The BE Journal of Theoretical Economics*, 2016.
- [16] M. Jallil and M. Perc, “Information cascades in complex networks,” *Journal of Complex Networks*, vol. 5, no. 5, pp. 665–693, 2017.
- [17] K. Chatterjee and B. Dutta, “Credibility and strategic learning in networks,” *Int. Economic Review*, vol. 57, no. 3, pp. 759–786, 2016.
- [18] A. Mitra, J. A. Richards, and S. Sundaram, “A communication-efficient algorithm for exponentially fast non-Bayesian learning in networks,” *arXiv preprint arXiv:1909.01505*, 2019.