

Distributed online convex optimization over jointly connected digraphs

David Mateos-Núñez Jorge Cortés

Abstract—This paper considers networked online convex optimization scenarios from a regret analysis perspective. At each round, each agent in the network commits to a decision and incurs in a local cost given by functions that are revealed over time and whose unknown evolution model might be adversarially adaptive to the agent’s behavior. The goal of each agent is to incur a cumulative cost over time with respect to the sum of local functions across the network that is competitive with the best single centralized decision in hindsight. To achieve this, agents cooperate with each other using local averaging over time-varying weight-balanced digraphs as well as subgradient descent on the local cost functions revealed in the previous round. We propose a class of coordination algorithms that generalize distributed online subgradient descent and saddle-point dynamics, allowing proportional-integral (and higher-order) feedback on the disagreement among neighboring agents. We show that our algorithm design achieves logarithmic agent regret (when local objectives are strongly convex), or square-root agent regret (when local objectives are convex) in scenarios where the communication graphs are jointly connected. Simulations in a medical diagnosis application illustrate our results.

Index Terms—distributed optimization; online optimization; regret analysis; jointly connected digraphs.

I. INTRODUCTION

Networked multi-agent systems are being increasingly deployed in scenarios where information is dynamic and increasingly revealed over time. Given the limited resources available to the network combined with the automatic and distributed data-collection, such scenarios bring to the forefront the need for optimizing network behavior in a distributed fashion and in real time. Motivated by these observations, we consider here a group of N agents seeking to solve a sequential decision problem over a time horizon T defined by the objective

$$\sum_{t=1}^T \sum_{i=1}^N f_t^i(x),$$

where each component function $f_t^i : \mathbb{R}^d \rightarrow \mathbb{R}$ becomes available to agent $i \in \{1, \dots, N\}$, and only to it, after having made its decision at time $t \in \{1, \dots, T\}$. Such networked online optimization problems arise in regression, classification, and other estimation problems in machine learning, where the functions $\{f_t^i\}$ measure the fitness of some model parameters, represented by the global decision vector $x \in \mathbb{R}^d$, with

respect to data sets that are incrementally revealed over time and become available in a distributed way (like in sensor networks or label-feedback systems). These problems naturally lend themselves to distributed algorithmic solutions because the information is distributed across the agents and making it centrally available might be costly (e.g., communication overhead, latency, and message drops), undesirable (e.g., privacy and security considerations), or poorly scalable (e.g., data sets which are large and change with time). In these scenarios, each agent would like to compute a provisional estimate of the global optimizer without waiting until all the data has been collected over time across the network. This represents a departure from standard static distributed optimization problems to a distributed online optimization framework that accounts for the adverse scenario of local decisions being evaluated against information available to the agents only after the decisions have been made.

As an example application, consider a network of hospitals that gather data over time about patients that might need some specific medical procedure. Each hospital has to assess the adequacy of the procedure upon prior observations and local communication with neighboring hospitals in the group. At each time instant, the estimated parameters of the predicting model governing the decisions of each hospital are evaluated against all the data most recently made available to the network, while each piece of this newly collected data is used locally by each hospital to improve its model to inform future decisions. In abstract terms, the goal is that each agent of the network performs in temporal average nearly as well as the best decision computed in hindsight had all patient data been centrally available. This notion is called *agent regret* and it is a performance measure for algorithm design and analysis in the framework just described of distributed online optimization.

Literature review: Distributed optimization problems are pervasive in distributed and parallel computation [1], [2], and distributed convex optimization constitutes a rich subfamily with many applications to multi-agent systems. This has motivated a growing body of work, see e.g., [3], [4], [5], [6], [7], [8], on the synthesis of distributed algorithms with asymptotic convergence guarantees to solve a variety of networked problems, including data fusion, network formation, and resource allocation. Our work here generalizes a family of distributed saddle-point subgradient algorithms [9], [7] that enjoy asymptotic convergence with constant stepsizes and robust asymptotic behavior in the presence of noise [10]. Online learning, on the other hand, performs sequential decision making given historical observations on the loss incurred by previous deci-

A preliminary version of this work appeared at the 2014 International Symposium on Mathematical Theory of Networks and Systems. This work was supported by NSF Award CMMI-1300272.

The authors are with the Department of Mechanical and Aerospace Engineering, University of California, San Diego, {dmateosn, cortes}@ucsd.edu.

sions, even when the loss functions are adversarially adaptive to the behavior of the decision maker. Interestingly, in online convex optimization [11], [12], [13], [14], it is doable to be competitive with the best single decision in hindsight. These works show how the regret, i.e., the difference between the cumulative cost over time and the cost of the best single decision in hindsight, is sublinear in the time horizon. Online convex optimization has applications to information theory [12], game theory [15], supervised online machine learning [13], online advertisement placement, and portfolio selection [14]. Algorithmic approaches include online gradient descent [11], online Newton step [16], follow-the-approximate-leader [16], and online alternating directions [17]. A few recent works have explored the combination of distributed and online convex optimization. The work [18] proposes distributed online strategies that rely on the computation and maintenance of spanning trees for global vector-sum operations and work under suitable statistical assumptions on the sequence of objectives. [19] studies decentralized online convex programming for groups of agents whose interaction topology is a chain. The works [20], [21] study agent regret without any statistical assumptions on the sequence of objectives. [20] introduces distributed online projected subgradient descent and shows square-root regret (for convex cost functions) and logarithmic regret (for strongly convex cost functions). However, the analysis critically relies on a projection step onto a compact set at each time step (which automatically guarantees the uniform boundedness of the estimates), and therefore excludes the unconstrained case (given the non-compactness of the whole state space). In contrast, [21] introduces distributed online dual averaging and shows square-root regret (for convex cost functions) using a general regularized projection that admits both unconstrained and constrained optimization, but the logarithmic bound is not established. Both works only consider static and strongly-connected interaction digraphs.

Statement of contributions: We consider a network of agents that communicate over a jointly connected sequence of time-dependent, weight-balanced digraphs. This means that the successive unions of consecutive digraphs over periods of time of a given length are strongly connected. The network is involved in an online unconstrained convex optimization scenario where no model is assumed about the evolution of the local objectives available to the agents. We propose a class of distributed coordination algorithms and study the associated agent regret in the optimization of the sum of the local cost functions across the network. Our algorithm design combines subgradient descent on the local objectives revealed in the previous round and proportional-integral (and/or higher-order) distributed feedback on the disagreement among neighboring agents. Assuming bounded subgradients of the local cost functions, we establish logarithmic agent regret bounds under local strong convexity and square-root agent regret under convexity plus a mild geometric condition. We also characterize the dependence of the regret bounds on the network parameters. Our technical approach uses the concept of network regret, that captures the performance of the sequence of collective estimates across the group of agents. The derivation of the

sublinear regret bounds results from three main steps: the study of the difference between network and agent regret; the analysis of the cumulative disagreement of the online estimates via the input-to-state stability property of a generalized Laplacian consensus dynamics; and the uniform boundedness of the online estimates (and auxiliary variables) when the set of local optimizers is uniformly bounded. With respect to previous work, the contributions advance the current state of the art because of the consideration of unconstrained formulations of the online optimization problem, which makes the discussion valid for regression and classification and raises major technical challenges to ensure the uniform boundedness of estimates; the synthesis of a novel family of coordination algorithms that generalize distributed online subgradient descent and saddle-point dynamics; and the development of regret guarantees under jointly connected interaction digraphs. Our novel analysis framework modularizes the main technical ingredients (the disagreement evolution via linear decoupling and input-to-state stability; the boundedness of estimates and auxiliary states through marginalizing the role of disagreement and learning rates; and the role played by network topology and the convexity properties) and extends and integrates techniques from distributed optimization (e.g., Lyapunov techniques for consensus under joint connectivity) and online optimization (e.g., Doubling Trick bounding techniques for square-root regret). We illustrate our results in a medical diagnosis example.

II. PRELIMINARIES

Here we introduce notational conventions and basic notions.

Linear algebra: We denote by $\mathbb{R}^n$ the n -dimensional Euclidean space, by $I_n \in \mathbb{R}^{n \times n}$ the identity matrix, and by $\mathbf{1}_n \in \mathbb{R}^n$ the column vector of all ones. For simplicity, we often use $(v_1, \dots, v_N)$ to represent the column vector $[v_1^\top, \dots, v_N^\top]^\top$. We denote by $\|\cdot\|_2$ the Euclidean norm and by $\mathcal{B}(x, \epsilon) := \{y \in \mathbb{R}^n : \|y - x\|_2 < \epsilon\}$ and $\bar{\mathcal{B}}(x, \epsilon)$ the open and closed balls, respectively, centered at x of radius ϵ . Given $w \in \mathbb{R}^n \setminus \{0\}$ and $c \in [0, 1]$, we let

$$\mathcal{F}_c(w) := \{v \in \mathbb{R}^n : v^\top w \geq c \|v\|_2 \|w\|_2\}$$

denote the convex cone of vectors in $\mathbb{R}^n$ whose angle with w has a cosine lower bounded by c . A matrix $A \in \mathbb{R}^{n \times n}$ is diagonalizable if it can be written as $A = S_A D_A S_A^{-1}$, where $D_A \in \mathbb{R}^{n \times n}$ is a diagonal matrix (whose entries are the eigenvalues of A), and $S_A \in \mathbb{R}^{n \times n}$ is an invertible matrix (whose columns are the corresponding eigenvectors). If the eigenvalues of A are real, we label them in increasing order from the minimum to the maximum as $\lambda_{\min}(A) = \lambda_1(A), \dots, \lambda_n(A) = \lambda_{\max}(A)$. For $B \in \mathbb{R}^{n \times m}$, we use $\|B\|_2 := \sigma_{\max}(B)$ for the largest singular value of B and $\kappa(B) := \|B\|_2 \|B^{-1}\|_2 = \sigma_{\max}(B) / \sigma_{\min}(B)$ for the condition number of B . The Kronecker product of $B \in \mathbb{R}^{n \times m}$ and $C \in \mathbb{R}^{p \times q}$ is denoted by $B \otimes C \in \mathbb{R}^{np \times mq}$.

Convex functions: Given a convex set $\mathcal{C} \subseteq \mathbb{R}^n$, a function $f : \mathcal{C} \rightarrow \mathbb{R}$ is convex if $f(\alpha x + (1-\alpha)y) \leq \alpha f(x) + (1-\alpha)f(y)$ for all $\alpha \in [0, 1]$ and $x, y \in \mathcal{C}$. A vector $\xi_x \in \mathbb{R}^n$ is a subgradient of f at $x \in \mathcal{C}$ if $f(y) - f(x) \geq \xi_x^\top (y - x)$, for all $y \in \mathcal{C}$. We denote by $\partial f(x)$ the set of all such subgradients. The

characterization in [22, Lemma 3.1.6] asserts that a function $f : \mathcal{C} \rightarrow \mathbb{R}$ is convex if and only if $\partial f(x)$ is nonempty for each $x \in \mathcal{C}$. Equivalently, f is convex if $\partial f(x)$ is nonempty and for each $x \in \mathcal{C}$ and $\xi_x \in \partial f(x)$,

$$f(y) - f(x) \geq \xi_x^\top (y - x) + \frac{p(x,y)}{2} \|y - x\|_2^2,$$

for all $y \in \mathcal{C}$, where $p : \mathcal{C} \times \mathcal{C} \rightarrow \mathbb{R}_{\geq 0}$ is the modulus of strong convexity (whose value may be 0). For $p > 0$, a function f is p -strongly convex on $\mathcal{C}$ if $p(x,y) = p$ for all $x, y \in \mathcal{C}$. Equivalently, f is p -strongly convex on $\mathcal{C}$ if

$$(\xi_y - \xi_x)^\top (y - x) \geq p \|y - x\|_2^2,$$

for each $\xi_x \in \partial f(x)$, $\xi_y \in \partial f(y)$, for all $x, y \in \mathcal{C}$. For convenience, we denote by $\operatorname{argmin}(f)$ the set of minimizers of a convex function f in its domain. For $\beta \in [0, 1]$, a convex function $f : \mathbb{R}^n \rightarrow \mathbb{R}$ with $\operatorname{argmin}(f) \neq \emptyset$ is β -central on $\mathcal{Z} \subseteq \mathbb{R}^n \setminus \operatorname{argmin}(f)$ if for each $x \in \mathcal{Z}$, there exists $y \in \operatorname{argmin}(f)$ such that $-\partial f(x) \subset \mathcal{F}_\beta(y - x)$, i.e.,

$$-\xi_x^\top (y - x) \geq \beta \|\xi_x\|_2 \|y - x\|_2,$$

for all $\xi_x \in \partial f(x)$. Note that any convex function $f : \mathbb{R}^n \rightarrow \mathbb{R}$ with a nonempty set of minimizers is at least 0-central on $\mathbb{R}^n \setminus \operatorname{argmin}(f)$. Finally, a convex function f has H -bounded subgradient sets if there exists $H \in \mathbb{R}_{>0}$ such that $\|\xi_x\|_2 \leq H$ for all $\xi_x \in \partial f(x)$ and $x \in \mathbb{R}^n$.

Graph theory: We review basic notions from graph theory following [23]. A (weighted) digraph $\mathcal{G} := (\mathcal{I}, \mathcal{E}, \mathbf{A})$ is a triplet where $\mathcal{I} := \{1, \dots, N\}$ is the vertex set, $\mathcal{E} \subseteq \mathcal{I} \times \mathcal{I}$ is the edge set, and $\mathbf{A} \in \mathbb{R}_{\geq 0}^{N \times N}$ is the weighted adjacency matrix with the property that $a_{ij} := A_{ij} > 0$ if and only if $(i, j) \in \mathcal{E}$. The complete graph is the digraph with edge set $\mathcal{I} \times \mathcal{I}$. Given $\mathcal{G}_1 = (\mathcal{I}, \mathcal{E}_1, \mathbf{A}_1)$ and $\mathcal{G}_2 = (\mathcal{I}, \mathcal{E}_2, \mathbf{A}_2)$, their union is the digraph $\mathcal{G}_1 \cup \mathcal{G}_2 = (\mathcal{I}, \mathcal{E}_1 \cup \mathcal{E}_2, \mathbf{A}_1 + \mathbf{A}_2)$. A path is an ordered sequence of vertices such that any pair of vertices appearing consecutively is an edge. A digraph is strongly connected if there is a path between any pair of distinct vertices. A sequence of digraphs $\{\mathcal{G}_t := (\mathcal{I}, \mathcal{E}_t, \mathbf{A}_t)\}_{t \geq 1}$ is δ -nondegenerate, for $\delta \in \mathbb{R}_{>0}$, if the weights are uniformly bounded away from zero by δ whenever positive, i.e., for each $t \in \mathbb{Z}_{\geq 1}$, $a_{ij,t} := (\mathbf{A}_t)_{ij} > \delta$ whenever $a_{ij,t} > 0$. A sequence $\{\mathcal{G}_t\}_{t \geq 1}$ is B -jointly connected, for $B \in \mathbb{Z}_{\geq 1}$, if for each $k \in \mathbb{Z}_{\geq 1}$, the digraph $\mathcal{G}_{kB} \cup \dots \cup \mathcal{G}_{(k+1)B-1}$ is strongly connected. The Laplacian matrix $\mathbf{L} \in \mathbb{R}^{N \times N}$ of a digraph $\mathcal{G}$ is $\mathbf{L} := \operatorname{diag}(\mathbf{A} \mathbf{1}_N) - \mathbf{A}$. Note that $\mathbf{L} \mathbf{1}_N = 0$. The weighted out-degree and in-degree of $i \in \mathcal{I}$ are, respectively, $d_{\text{out}}(i) := \sum_{j=1}^N a_{ij}$ and $d_{\text{in}}(i) := \sum_{j=1}^N a_{ji}$. A digraph is weight-balanced if $d_{\text{out}}(i) = d_{\text{in}}(i)$ for all $i \in \mathcal{I}$, that is, $\mathbf{1}_N^\top \mathbf{L} = 0$. For convenience, we let $\mathbf{L}_\mathcal{K}$ denote the Laplacian of the complete graph with edge weights $1/N$, i.e., $\mathbf{L}_\mathcal{K} := \mathbf{I}_N - \mathbf{M}$, where $\mathbf{M} := \frac{1}{N} \mathbf{1}_N \mathbf{1}_N^\top$. Note that $\mathbf{L}_\mathcal{K}$ is idempotent, i.e., $\mathbf{L}_\mathcal{K}^2 = \mathbf{L}_\mathcal{K}$. For the reader's convenience, Table I collects the shorthand notation combining Laplacian matrices and Kronecker products used in the paper.

III. PROBLEM STATEMENT

This section introduces the problem of interest. We begin by describing the online convex optimization problem for one

$\mathbf{M} = \frac{1}{N} \mathbf{1}_N \mathbf{1}_N^\top$	$\mathbf{M} = \mathbf{M} \otimes \mathbf{I}_d$	$\hat{\mathbf{L}}_\mathcal{K} = \mathbf{I}_K \otimes \mathbf{L}_\mathcal{K}$
$\mathbf{L}_\mathcal{K} = \mathbf{I}_N - \mathbf{M}$	$\mathbf{L}_\mathcal{K} = \mathbf{L}_\mathcal{K} \otimes \mathbf{I}_d$	$\mathbf{L}_t = \mathbf{E} \otimes \mathbf{L}_t$
$\mathbf{L}_t = \operatorname{diag}(\mathbf{A}_t \mathbf{1}_N) - \mathbf{A}_t$	$\mathbf{L}_t = \mathbf{L}_t \otimes \mathbf{I}_d$	

TABLE I: Shorthand notation for graph matrices employed along the paper. Here, $\{\mathcal{G}_t\}_{t \geq 1}$, $K \in \mathbb{Z}_{\geq 1}$, and $\mathbf{E} \in \mathbb{R}^{K \times K}$.

player and then present the networked version, which is the focus of the paper. In online convex optimization, given a time horizon $T \in \mathbb{Z}_{\geq 1}$, in each round $t \in \{1, \dots, T\}$ a player chooses a point $x_t \in \mathbb{R}^d$. After committing to this choice, a convex cost function $f_t : \mathbb{R}^d \rightarrow \mathbb{R}$ is revealed. Consequently, the ‘cost’ incurred by the player is $f_t(x_t)$. Given the temporal sequence of objectives $\{f_t\}_{t=1}^T$, the regret of the player using $\{x_t\}_{t=1}^T$ with respect to a single choice $u \in \mathbb{R}^d$ in hindsight over a time horizon T is defined by

$$\mathcal{R}(u, \{f_t\}_{t=1}^T) := \sum_{t=1}^T f_t(x_t) - \sum_{t=1}^T f_t(u), \quad (1)$$

i.e., the difference between the total cost incurred by the online estimates $\{x_t\}_{t=1}^T$ and the cost of a single hindsight decision u . A logical choice, if it exists, is the best decision over a time horizon T had all the information been available a priori, i.e.,

$$u = x_T^* \in \arg \min_{x \in \mathbb{R}^d} \sum_{t=1}^T f_t(x).$$

In the case when no information is available about the evolution of the functions $\{f_t\}_{t=1}^T$, one is interested in designing algorithms whose worst-case regret is upper bounded sublinearly in the time horizon T with respect to any decision in hindsight. This ensures that, on average, the algorithm performs nearly as well as the best single decision in hindsight.

We now explain the distributed version of the online convex optimization problem where the online player is replaced by a network of N agents, each with access to partial information. In the round $t \in \{1, \dots, T\}$, agent $i \in \{1, \dots, N\}$ chooses a point x_t^i corresponding to what it thinks the network as a whole should have chosen. After committing to this choice, the agent has access to a convex cost function $f_t^i : \mathbb{R}^d \rightarrow \mathbb{R}$ and the network cost is then given by the evaluation of

$$f_t(x) := \sum_{i=1}^N f_t^i(x). \quad (2)$$

Note that this function is not known to any of the agents and is not available at any single location. In this scenario, the regret of agent $j \in \{1, \dots, N\}$ using $\{x_t^j\}_{t=1}^T$ with respect to a single choice u in hindsight over a time horizon T is

$$\mathcal{R}^j(u, \{f_t\}_{t=1}^T) := \sum_{t=1}^T \sum_{i=1}^N f_t^i(x_t^j) - \sum_{t=1}^T \sum_{i=1}^N f_t^i(u).$$

The goal then is to design coordination algorithms among the agents that guarantee that the worst-case agent regret is upper bounded sublinearly in the time horizon T with respect to any decision in hindsight. This would guarantee that each agent incurs an average cost over time with respect to the sum of local cost functions across the network that is nearly as low as the cost of the best single choice had all the information

been centrally available a priori. Since information is now distributed across the network, agents must collaborate with each other to determine their decisions for the next round. We assume that the network communication topology is time-dependent and described by a sequence of weight-balanced digraphs $\{\mathcal{G}_t\}_{t=1}^T = \{(\{1, \dots, N\}, \mathcal{E}_t, \mathbf{A}_t)\}_{t=1}^T$. At each round, agents can use historical observations of locally revealed cost functions and become aware through local communication of the choices made by their neighbors in the previous round.

IV. DYNAMICS FOR DISTRIBUTED ONLINE OPTIMIZATION

In this section we propose a distributed coordination algorithm to solve the networked online convex optimization problem described in Section III. In each round $t \in \{1, \dots, T\}$, agent $i \in \{1, \dots, N\}$ performs

$$\begin{aligned} x_{t+1}^i &= x_t^i + \sigma \left(a \sum_{j=1}^N \mathbf{a}_{ij,t} (x_t^j - x_t^i) + \sum_{j=1}^N \mathbf{a}_{ij,t} (z_t^j - z_t^i) \right) - \eta_t g_{x_t^i}, \\ z_{t+1}^i &= z_t^i - \sigma \sum_{j=1}^N \mathbf{a}_{ij,t} (x_t^j - x_t^i), \end{aligned} \quad (3)$$

where $g_{x_t^i} \in \partial f_t^i(x_t^i)$, the scalars $\sigma, a \in \mathbb{R}_{>0}$ are design parameters, and $\eta_t \in \mathbb{R}_{>0}$ is the learning rate at time t . Agent i is responsible for the variables x^i, z^i , and shares their values with its neighbors according to the time-dependent digraph $\mathcal{G}_t$. Note that (3) is both consistent with the notion of incremental access to information by individual agents and is distributed over $\mathcal{G}_t$: each agent updates its estimate by following a subgradient of the cost function revealed to it in the previous round while, at the same time, seeking to agree with its neighbors' estimates. The latter is implemented through a second-order process that employs proportional-integral feedback on the disagreement. Our design is inspired by and extends the distributed algorithms for distributed optimization of a sum of convex functions studied in [9], [7]. We use the term *online subgradient descent algorithm with proportional-integral disagreement feedback* to refer to (3).

We next rewrite the dynamics in compact form. To do so, we introduce the notation $\mathbf{x} := (x^1, \dots, x^N) \in (\mathbb{R}^d)^N$ and $\mathbf{z} := (z^1, \dots, z^N) \in (\mathbb{R}^d)^N$ to denote the aggregate of the agents' online decisions and the aggregate of the agents' auxiliary variables, respectively. For $t \in \{1, \dots, T\}$, we also define the convex function $\tilde{f}_t : (\mathbb{R}^d)^N \rightarrow \mathbb{R}$ by

$$\tilde{f}_t(\mathbf{x}) := \sum_{i=1}^N f_t^i(x^i). \quad (4)$$

When all agents agree on the same choice, one recovers the value of the network cost function (2), $\tilde{f}_t(\mathbb{1}_N \otimes x) = f_t(x)$. With this notation in place, the algorithm (3) takes the form

$$\begin{bmatrix} \mathbf{x}_{t+1} \\ \mathbf{z}_{t+1} \end{bmatrix} = \begin{bmatrix} \mathbf{x}_t \\ \mathbf{z}_t \end{bmatrix} - \sigma \begin{bmatrix} a\mathbf{L}_t & \mathbf{L}_t \\ -\mathbf{L}_t & 0 \end{bmatrix} \begin{bmatrix} \mathbf{x}_t \\ \mathbf{z}_t \end{bmatrix} - \eta_t \begin{bmatrix} \tilde{\mathbf{g}}_{\mathbf{x}_t} \\ 0 \end{bmatrix}, \quad (5)$$

where $\mathbf{L}_t := \mathbf{L}_t \otimes \mathbf{I}_d$ and $\tilde{\mathbf{g}}_{\mathbf{x}_t} = (g_{x_t^1}, \dots, g_{x_t^N}) \in \partial \tilde{f}_t(\mathbf{x}_t)$.

This compact-form representation suggests a more general class of distributed dynamics that includes (3) as a particular

case. For $K \in \mathbb{Z}_{\geq 1}$, let $E \in \mathbb{R}^{K \times K}$ be diagonalizable with real positive eigenvalues, and define $\mathbb{L}_t := E \otimes \mathbf{L}_t$. Consider the dynamics on $((\mathbb{R}^d)^N)^K$ defined by

$$\mathbf{v}_{t+1} = (\mathbf{I}_{KNd} - \sigma \mathbb{L}_t) \mathbf{v}_t - \eta_t \mathbf{g}_t, \quad (6)$$

where $\mathbf{g}_t \in ((\mathbb{R}^d)^N)^K$ takes the form

$$\mathbf{g}_t := (\tilde{\mathbf{g}}_{\mathbf{x}_t}, 0, \dots, 0), \quad (7)$$

and we use the decomposition $\mathbf{v} = (\mathbf{x}, \mathbf{v}^2, \dots, \mathbf{v}^K)$. Throughout the paper, our convergence results are formulated for this dynamics because of its generality, which we discuss in the following remark.

Remark IV.1. (Online subgradient descent algorithms with proportional and proportional-integral disagreement feedback): The online subgradient descent algorithm with proportional-integral disagreement feedback (3) corresponds to the dynamics (6) with the choices $K = 2$ and

$$E = \begin{bmatrix} a & 1 \\ -1 & 0 \end{bmatrix}.$$

For $a \in (2, \infty)$, E has positive eigenvalues $\lambda_{\min}(E) = \frac{a}{2} - \sqrt{(\frac{a}{2})^2 - 1}$ and $\lambda_{\max}(E) = \frac{a}{2} + \sqrt{(\frac{a}{2})^2 - 1}$. Interestingly, the online subgradient descent algorithm with proportional disagreement feedback proposed in [20] (without the projection component onto a bounded convex set) also corresponds to the dynamics (6) with the choices $K = 1$ and $E = [1]$. •

Our forthcoming exposition presents the technical approach to establish the properties of the distributed dynamics (6) with respect to the agent regret defined in Section III. An informal description of our main results is as follows. Under mild conditions on the connectivity of the communication network, a suitable choice of σ , and the assumption that the time-dependent local cost functions have bounded subgradient sets and uniformly bounded optimizers, the following bounds hold:

Logarithmic agent regret: if each local cost function is locally p -strongly convex and $\eta_t = \frac{1}{p^t}$, then any sequence generated by the dynamics (6) satisfies, for each $j \in \{1, \dots, N\}$,

$$\mathcal{R}^j(u, \{f_t\}_{t=1}^T) \in \mathcal{O}(\|u\|_2^2 + \log T).$$

Square-root agent regret: if each local cost function is convex (plus a mild geometric assumption) and, for $m = 0, 1, 2, \dots, \lceil \log_2 T \rceil$, we take $\eta_t = \frac{1}{\sqrt{2^m}}$ in each period of 2^m rounds $t = 2^m, \dots, 2^{m+1} - 1$, then any sequence generated by the dynamics (6) satisfies, for each $j \in \{1, \dots, N\}$,

$$\mathcal{R}^j(u, \{f_t\}_{t=1}^T) \in \mathcal{O}(\|u\|_2^2 \sqrt{T}).$$

In our technical approach to establish these sublinear agent regret bounds, we find it useful to consider the notion of network regret [18], [24] with respect to a single hindsight choice $u \in \mathbb{R}^d$ over the time horizon T ,

$$\mathcal{R}_{\mathcal{N}}(u, \{\tilde{f}_t\}_{t=1}^T) := \sum_{t=1}^T \tilde{f}_t(\mathbf{x}_t) - \sum_{t=1}^T \tilde{f}_t(\mathbb{1}_N \otimes u),$$

to capture the performance of the sequence of collective estimates $\{\mathbf{x}_t\}_{t=1}^T \subseteq (\mathbb{R}^d)^N$. Our proof strategy builds on this concept and relies on bounding the following terms:

- (i) both the network regret and the difference between the agent and network regrets;
- (ii) the cumulative disagreement of the collective estimates;
- (iii) the sequence of collective estimates uniformly in the time horizon.

Section V presents the formal discussion for these results. The combination of these steps allows us in Section VI to formally establish the sublinear agent regret bounds outlined above.

V. REGRET ANALYSIS

This section presents the results outlined above on bounding the agent and network regrets, the cumulative disagreement of the collective estimates, and the sequence of collective estimates for executions of the distributed dynamics (6). These results are instrumental later in the derivation of the sublinear agent regret bounds, but are also of independent interest.

A. Bounds on network and agent regret

Our first result relates the agent and network regrets for any sequence of collective estimates (regardless of the algorithm that generates them) in terms of their cumulative disagreement.

Lemma V.1. (Bound on the difference between agent and network regret): For $T \in \mathbb{Z}_{\geq 1}$, let $\{f_t^1, \dots, f_t^N\}_{t=1}^T$ be convex functions on $\mathbb{R}^d$ with H -bounded subgradient sets. Then, any sequence $\{\mathbf{x}_t\}_{t=1}^T \subset (\mathbb{R}^d)^N$ satisfies, for any $j \in \{1, \dots, N\}$ and $\mathbf{u} \in \mathbb{R}^d$,

$$\mathcal{R}^j(u, \{f_t\}_{t=1}^T) \leq \mathcal{R}_N(u, \{\tilde{f}_t\}_{t=1}^T) + NH \sum_{t=1}^T \|\mathbf{L}_K \mathbf{x}_t\|_2,$$

where $\mathbf{L}_K := \mathbf{L}_K \otimes \mathbf{I}_d$.

Proof: Since $\tilde{f}_t(\mathbf{1}_N \otimes x) = f_t(x)$ for all $x \in \mathbb{R}^d$, we have

$$\begin{aligned} \mathcal{R}^j(u, \{f_t\}_{t=1}^T) - \mathcal{R}_N(u, \{\tilde{f}_t\}_{t=1}^T) \\ = \sum_{t=1}^T (\tilde{f}_t(\mathbf{1}_N \otimes x_t^j) - \tilde{f}_t(\mathbf{x}_t)). \end{aligned} \quad (8)$$

The convexity of $\tilde{f}_t$ implies that, for any $\xi \in \partial \tilde{f}_t(\mathbf{1}_N \otimes x_t^j)$,

$$\begin{aligned} \tilde{f}_t(\mathbf{1}_N \otimes x_t^j) - \tilde{f}_t(\mathbf{x}_t) &\leq \xi^\top (\mathbf{1}_N \otimes x_t^j - \mathbf{x}_t) \\ &\leq \|\xi\|_2 \|\mathbf{1}_N \otimes x_t^j - \mathbf{x}_t\|_2 \leq \sqrt{N}H \|\mathbf{1}_N \otimes x_t^j - \mathbf{x}_t\|_2, \end{aligned} \quad (9)$$

where we have used the Cauchy-Schwarz inequality and the fact that the subgradient sets are H -bounded. In addition,

$$\begin{aligned} \|\mathbf{1}_N \otimes x_t^j - \mathbf{x}_t\|_2^2 &= \sum_{i=1}^N \|x_t^j - x_t^i\|_2^2 \\ &\leq \frac{1}{2} \sum_{j=1}^N \sum_{i=1}^N \|x_t^j - x_t^i\|_2^2 = N \mathbf{x}_t^\top \mathbf{L}_K \mathbf{x}_t. \end{aligned} \quad (10)$$

The fact that $\mathbf{L}_K^2 = \mathbf{L}_K = \mathbf{L}_K^\top$ allows us to write $\mathbf{x}_t^\top \mathbf{L}_K \mathbf{x}_t = \|\mathbf{L}_K \mathbf{x}_t\|_2^2$. The result now follows using (9) and (10) in conjunction with (8). ■

Next, we bound the network regret for executions of the coordination algorithm (6) in terms of the learning rates and the cumulative disagreement. The bound holds regardless of the connectivity of the communication network as long as the digraph remains weight-balanced.

Lemma V.2. (Bound on network regret): For $T \in \mathbb{Z}_{\geq 1}$, let $\{f_t^1, \dots, f_t^N\}_{t=1}^T$ be convex functions on $\mathbb{R}^d$ with H -bounded subgradient sets. Let the sequence $\{\mathbf{x}_t\}_{t=1}^T$ be generated by the coordination algorithm (6) over a sequence of arbitrary weight-balanced digraphs $\{\mathcal{G}_t\}_{t=1}^T$. Then, for any $\mathbf{u} \in \mathbb{R}^d$, and any sequence of learning rates $\{\eta_t\}_{t=1}^T \subset \mathbb{R}_{>0}$,

$$\begin{aligned} 2\mathcal{R}_N(u, \{\tilde{f}_t\}_{t=1}^T) &\leq \sum_{t=2}^T \|\mathbf{M} \mathbf{x}_t - \mathbf{u}\|_2^2 \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} - p_t(\mathbf{u}, \mathbf{x}_t) \right) \\ &\quad + 2\sqrt{N}H \sum_{t=1}^T \|\mathbf{L}_K \mathbf{x}_t\|_2 + NH^2 \sum_{t=1}^T \eta_t + \frac{1}{\eta_1} \|\mathbf{M} \mathbf{x}_1 - \mathbf{u}\|_2^2, \end{aligned}$$

where $\mathbf{M} := \mathbf{M} \otimes \mathbf{I}_d$, $\mathbf{u} := \mathbf{1}_N \otimes \mathbf{u}$ and $p_t : (\mathbb{R}^d)^N \times (\mathbb{R}^d)^N \rightarrow \mathbb{R}_{\geq 0}$ is the modulus of strong convexity of $\tilde{f}_t$.

Proof: Left-multiplying the dynamics (6) by the block-diagonal matrix $\text{diag}(1, 0, \dots, 0) \otimes \mathbf{M} \in \mathbb{R}^{(Nd)K \times (Nd)K}$, and using $\mathbf{M} \mathbf{L}_t = 0$, we obtain the following projected dynamics

$$\mathbf{M} \mathbf{x}_{t+1} = \mathbf{M} \mathbf{x}_t - \eta_t \mathbf{M} \tilde{g}_{\mathbf{x}_t}. \quad (11)$$

Note that this dynamics is decoupled from the dynamics of $\mathbf{v}_t^2, \dots, \mathbf{v}_t^K$. Subtracting $\mathbf{u}$ and taking the norm on both sides, we get $\|\mathbf{M} \mathbf{x}_{t+1} - \mathbf{u}\|_2^2 = \|\mathbf{M} \mathbf{x}_t - \mathbf{u} - \eta_t \mathbf{M} \tilde{g}_{\mathbf{x}_t}\|_2^2$, so that

$$\begin{aligned} \|\mathbf{M} \mathbf{x}_{t+1} - \mathbf{u}\|_2^2 &= \|\mathbf{M} \mathbf{x}_t - \mathbf{u}\|_2^2 + \eta_t^2 \|\mathbf{M} \tilde{g}_{\mathbf{x}_t}\|_2^2 - 2\eta_t (\mathbf{M} \tilde{g}_{\mathbf{x}_t})^\top (\mathbf{M} \mathbf{x}_t - \mathbf{u}) \\ &= \|\mathbf{M} \mathbf{x}_t - \mathbf{u}\|_2^2 + \eta_t^2 \|\mathbf{M} \tilde{g}_{\mathbf{x}_t}\|_2^2 - 2\eta_t \tilde{g}_{\mathbf{x}_t}^\top (\mathbf{M} \mathbf{x}_t - \mathbf{u}), \end{aligned} \quad (12)$$

where we have used $\mathbf{M}^2 = \mathbf{M}$ and $\mathbf{M} \mathbf{u} = \mathbf{u}$. Regarding the last term, note that

$$\begin{aligned} -\tilde{g}_{\mathbf{x}_t}^\top (\mathbf{M} \mathbf{x}_t - \mathbf{u}) &= -\tilde{g}_{\mathbf{x}_t}^\top (\mathbf{M} \mathbf{x}_t - \mathbf{x}_t) - \tilde{g}_{\mathbf{x}_t}^\top (\mathbf{x}_t - \mathbf{u}) \\ &\leq \tilde{g}_{\mathbf{x}_t}^\top \mathbf{L}_K \mathbf{x}_t + \tilde{f}_t(\mathbf{u}) - \tilde{f}_t(\mathbf{x}_t) - \frac{p_t(\mathbf{u}, \mathbf{x}_t)}{2} \|\mathbf{u} - \mathbf{x}_t\|_2^2, \end{aligned}$$

where we have used $\mathbf{L}_K = \mathbf{I}_{Nd} - \mathbf{M}$. Substituting into (12), we obtain

$$\begin{aligned} \|\mathbf{M} \mathbf{x}_{t+1} - \mathbf{u}\|_2^2 &\leq \|\mathbf{M} \mathbf{x}_t - \mathbf{u}\|_2^2 + \eta_t^2 \|\mathbf{M} \tilde{g}_{\mathbf{x}_t}\|_2^2 \\ &\quad + 2\eta_t \left(\tilde{g}_{\mathbf{x}_t}^\top \mathbf{L}_K \mathbf{x}_t + \tilde{f}_t(\mathbf{u}) - \tilde{f}_t(\mathbf{x}_t) - \frac{p_t(\mathbf{u}, \mathbf{x}_t)}{2} \|\mathbf{u} - \mathbf{x}_t\|_2^2 \right), \end{aligned}$$

so that, reordering terms,

$$\begin{aligned} 2(\tilde{f}_t(\mathbf{x}_t) - \tilde{f}_t(\mathbf{u})) &\leq \frac{1}{\eta_t} (\|\mathbf{M} \mathbf{x}_t - \mathbf{u}\|_2^2 - \|\mathbf{M} \mathbf{x}_{t+1} - \mathbf{u}\|_2^2) \\ &\quad - p_t(\mathbf{u}, \mathbf{x}_t) \|\mathbf{x}_t - \mathbf{u}\|_2^2 + 2\tilde{g}_{\mathbf{x}_t}^\top \mathbf{L}_K \mathbf{x}_t + \eta_t \|\mathbf{M} \tilde{g}_{\mathbf{x}_t}\|_2^2. \end{aligned} \quad (13)$$

Next, we bound each of the terms appearing in the last line of (13). For the term $p_t(\mathbf{u}, \mathbf{x}_t) \|\mathbf{x}_t - \mathbf{u}\|_2^2$, note that

$$\begin{aligned} \|\mathbf{x}_t - \mathbf{u}\|_2^2 &= \|(\mathbf{M} + \mathbf{L}_K)(\mathbf{x}_t - \mathbf{u})\|_2^2 = \|\mathbf{M}(\mathbf{x}_t - \mathbf{u})\|_2^2 \\ &\quad + \|\mathbf{L}_K(\mathbf{x}_t - \mathbf{u})\|_2^2 + 2(\mathbf{x}_t - \mathbf{u})^\top \mathbf{M} \mathbf{L}_K (\mathbf{x}_t - \mathbf{u}) \\ &= \|\mathbf{M} \mathbf{x}_t - \mathbf{u}\|_2^2 + \|\mathbf{L}_K \mathbf{x}_t\|_2^2, \end{aligned} \quad (14)$$

where we have used $\mathbf{M} \mathbf{L}_K = 0$ and $\mathbf{M} \mathbf{u} = \mathbf{u}$. Regarding the term $2\tilde{g}_{\mathbf{x}_t}^\top \mathbf{L}_K \mathbf{x}_t$, note that $\|\tilde{g}_{\mathbf{x}_t}\|_2^2 = \sum_{i=1}^N \|g_{x_t^i}\|_2^2 \leq NH^2$

because the subgradient sets are bounded by H . Hence, using the Cauchy-Schwarz inequality,

$$\tilde{g}_{x_t}^\top \mathbf{L}_K \mathbf{x}_t \leq \|\tilde{g}_{x_t}\|_2 \|\mathbf{L}_K \mathbf{x}_t\|_2 \leq \sqrt{N}H \|\mathbf{L}_K \mathbf{x}_t\|_2. \quad (15)$$

Finally, regarding the term $\eta_t \|\mathbf{M} \tilde{g}_{x_t}\|_2^2$ in (13), note that

$$\begin{aligned} \|\mathbf{M} \tilde{g}_{x_t}\|_2^2 &= \|\mathbf{1}_N \otimes \frac{1}{N} \sum_{i=1}^N g_{x_t^i}\|_2^2 = N \left\| \frac{1}{N} \sum_{i=1}^N g_{x_t^i} \right\|_2^2 \\ &= \frac{1}{N} \sum_{l=1}^d \left(\sum_{i=1}^N g_{x_t^i} \right)_l^2 \leq \frac{1}{N} \sum_{l=1}^d \left(N \sum_{i=1}^N (g_{x_t^i})_l^2 \right) \\ &= \sum_{i=1}^N \sum_{l=1}^d (g_{x_t^i})_l^2 = \sum_{i=1}^N \|g_{x_t^i}\|_2^2 \leq NH^2, \end{aligned} \quad (16)$$

where in the first inequality we have used the inequality of quadratic and arithmetic means [25]. The result now follows from summing the expression in (13) over the time horizon T , discarding the negative terms, and using the upper bounds in (14)-(16). ■

The combination of Lemmas V.1 and V.2 provides a bound on the agent regret in terms of the learning rates and the cumulative disagreement of the collective estimates. This motivates our next section.

B. Bound on cumulative disagreement

In this section we study the evolution of the disagreement among the agents' estimates under (6). Our analysis builds on the input-to-state stability (ISS) properties of the linear part of the dynamics with respect to the agreement subspace, where we treat the subgradient term as a perturbation. Consequently, here we study the dynamics

$$\mathbf{v}_{t+1} = (\mathbf{I}_{KNd} - \sigma \mathbb{L}_t) \mathbf{v}_t + \mathbf{d}_t, \quad (17)$$

where $\{\mathbf{d}_t\}_{t \geq 1} \subset ((\mathbb{R}^d)^N)^K$ is an arbitrary sequence of disturbances. Our first result shows that, for the purpose of studying the ISS properties of (17), the dynamics can be decoupled into K first-order linear consensus dynamics.

Lemma V.3. (Decoupling into a collection of first-order consensus dynamics): Given a diagonalizable matrix $E \in \mathbb{R}^{K \times K}$ with real eigenvalues, let S_E be the matrix of eigenvectors in the decomposition $E = S_E D_E S_E^{-1}$, with $D_E = \text{diag}(\lambda_1(E), \dots, \lambda_K(E))$. Then, under the change of variables

$$\mathbf{w}_t := (S_E^{-1} \otimes \mathbf{I}_{Nd}) \mathbf{v}_t, \quad (18)$$

the dynamics (17) is equivalently represented by the collection of first-order dynamics on $(\mathbb{R}^d)^N$ defined by

$$\mathbf{w}_{t+1}^l = (\mathbf{I}_{Nd} - \sigma \lambda_l(E) \mathbf{L}_t) \mathbf{w}_t^l + \mathbf{e}_t^l, \quad (19)$$

where $l \in \{1, \dots, K\}$, $\mathbf{w}_t = (\mathbf{w}_t^1, \dots, \mathbf{w}_t^K) \in ((\mathbb{R}^d)^N)^K$ and

$$\mathbf{e}_t^l := ((S_E^{-1} \otimes \mathbf{I}_{Nd}) \mathbf{d}_t)^l \in (\mathbb{R}^d)^N. \quad (20)$$

Moreover, for each $t \in \mathbb{Z}_{\geq 1}$,

$$\|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2 \leq \|S_E\|_2 \sqrt{K} \max_{1 \leq l \leq K} \|\mathbf{L}_K \mathbf{w}_t^l\|_2, \quad (21)$$

where $\hat{\mathbf{L}}_K := \mathbf{I}_K \otimes \mathbf{L}_K$.

Proof: We start by noting that

$$\begin{aligned} \mathbb{L}_t &= S_E D_E S_E^{-1} \otimes \mathbf{I}_{Nd} \mathbf{L}_t \mathbf{I}_{Nd} \\ &= (S_E \otimes \mathbf{I}_{Nd}) (D_E \otimes \mathbf{L}_t) (S_E \otimes \mathbf{I}_{Nd})^{-1}, \end{aligned}$$

and therefore we obtain the factorization

$$\mathbf{I}_{KNd} - \sigma \mathbb{L}_t = (S_E \otimes \mathbf{I}_{Nd}) (\mathbf{I}_{KNd} - \sigma D_E \otimes \mathbf{L}_t) (S_E \otimes \mathbf{I}_{Nd})^{-1}.$$

Now, under the change of variables (18), the dynamics (17) takes the form

$$\mathbf{w}_{t+1} = (\mathbf{I}_{KNd} - \sigma D_E \otimes \mathbf{L}_t) \mathbf{w}_t + (S_E^{-1} \otimes \mathbf{I}_{Nd}) \mathbf{d}_t, \quad (22)$$

which corresponds to the set of dynamics (19). Moreover,

$$\begin{aligned} \hat{\mathbf{L}}_K \mathbf{v}_t &= (\mathbf{I}_K \otimes \mathbf{L}_K) (S_E \otimes \mathbf{I}_{Nd}) \mathbf{w}_t \\ &= (S_E \otimes \mathbf{I}_{Nd}) (\mathbf{I}_K \otimes \mathbf{L}_K) \mathbf{w}_t = (S_E \otimes \mathbf{I}_{Nd}) \hat{\mathbf{L}}_K \mathbf{w}_t. \end{aligned}$$

Hence, the sub-multiplicativity of the norm together with [26, Fact 9.12.22] for the norms of Kronecker products, yields

$$\begin{aligned} \|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2 &\leq \|S_E \otimes \mathbf{I}_{Nd}\|_2 \|\hat{\mathbf{L}}_K \mathbf{w}_t\|_2 = \|S_E\|_2 \|\hat{\mathbf{L}}_K \mathbf{w}_t\|_2 \\ &= \|S_E\|_2 \left(\sum_{l=1}^K \|\mathbf{L}_K \mathbf{w}_t^l\|_2^2 \right)^{1/2} \leq \|S_E\|_2 \sqrt{K} \max_{1 \leq l \leq K} \|\mathbf{L}_K \mathbf{w}_t^l\|_2, \end{aligned}$$

as claimed. ■

In the next result, we use Lemma V.3 to bound the cumulative disagreement of the collective estimates over time.

Proposition V.4. (Input-to-state stability and cumulative disagreement of (17) over jointly connected weight-balanced digraphs): Let $E \in \mathbb{R}^{K \times K}$ be a diagonalizable matrix with real positive eigenvalues and $\{\mathcal{G}_s\}_{s \geq 1}$ a sequence of B -jointly connected, δ -nondegenerate, weight-balanced digraphs. For $\tilde{\delta}' \in (0, 1)$, let

$$\tilde{\delta} := \min \left\{ \tilde{\delta}', (1 - \tilde{\delta}') \frac{\lambda_{\min}(E) \delta}{\lambda_{\max}(E) d_{\max}} \right\}, \quad (23)$$

where

$$d_{\max} := \max \{ d_{\text{out},t}(k) : k \in \mathcal{I}, 1 \leq t \leq T \}. \quad (24)$$

Then, for any choice

$$\sigma \in \left[\frac{\tilde{\delta}}{\lambda_{\min}(E) \delta}, \frac{1 - \tilde{\delta}}{\lambda_{\max}(E) d_{\max}} \right], \quad (25)$$

the dynamics (17) over $\{\mathcal{G}_s\}_{s \geq 1}$ is input-to-state stable with respect to the nullspace of the matrix $\hat{\mathbf{L}}_K$. Specifically, for any $t \in \mathbb{Z}_{\geq 1}$ and any $\{\mathbf{d}_s\}_{s=1}^{t-1} \subset ((\mathbb{R}^d)^N)^K$,

$$\begin{aligned} \|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2 &\leq C_{\mathcal{I}} \|\mathbf{v}_1\|_2 \left(1 - \frac{\tilde{\delta}}{4N^2} \right)^{\lceil \frac{t-1}{B} \rceil} \\ &\quad + C_{\mathcal{U}} \max_{1 \leq s \leq t-1} \|\mathbf{d}_s\|_2, \end{aligned} \quad (26)$$

where

$$C_{\mathcal{I}} := \kappa(S_E) \sqrt{K} \left(\frac{4}{3} \right)^2, \quad C_{\mathcal{U}} := \frac{C_{\mathcal{I}}}{1 - \left(1 - \frac{\tilde{\delta}}{4N^2} \right)^{1/B}}. \quad (27)$$

And the cumulative disagreement satisfies, for $T \in \mathbb{Z}_{\geq 1}$,

$$\sum_{t=1}^T \|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2 \leq C_{\mathcal{U}} \left(\|\mathbf{v}_1\|_2 + \sum_{t=1}^{T-1} \|\mathbf{d}_t\|_2 \right). \quad (28)$$

Proof: The strategy of the proof is the following. We use Lemma V.3 to decouple (17) into K copies (for each eigenvalue of E) of the same first-order linear consensus dynamics. We then analyze the convergence properties of the latter using [27, Th. 1.2]. Finally, we bound the disagreement in the original network variables using again Lemma V.3.

We start by noting that the selection of $\tilde{\delta}$ makes the set in (25) nonempty and consequently the selection of σ feasible. We write the dynamics (19), omitting the dependence on $l \in \{1, \dots, K\}$ for the sake of clarity, as

$$\mathbf{y}_{t+1} = (\mathbf{I}_{Nd} - \hat{\sigma} \mathbf{L}_t) \mathbf{y}_t + \mathbf{e}_t, \quad (29)$$

where $\hat{\sigma} := \sigma \lambda_l(E) > 0$ and $\mathbf{e}_t := \mathbf{e}_t^l$. From (20), we have

$$\|\mathbf{e}_t\|_2 \leq \|(\mathbf{S}_E^{-1} \otimes \mathbf{I}_{Nd}) \mathbf{d}_t\|_2 \leq \|\mathbf{S}_E^{-1}\|_2 \|\mathbf{d}_t\|_2, \quad (30)$$

for each $t \in \mathbb{Z}_{\geq 1}$. Next, let

$$\mathbf{P}_t := \mathbf{I}_N - \hat{\sigma} \mathbf{L}_t = \hat{\sigma} \mathbf{A}_t + \mathbf{I}_N - \hat{\sigma} \mathbf{D}_{\text{out},t}, \quad (31)$$

and define $\Phi(k, s) := (\mathbf{P}_k \mathbf{P}_{k-1} \dots \mathbf{P}_{s+1} \mathbf{P}_s) \otimes \mathbf{I}_d$, for each $k \geq s \geq 1$. The trajectory of (29) can then be expressed as

$$\mathbf{y}_{t+1} = \Phi(t, 1) \mathbf{y}_1 + \sum_{s=1}^{t-1} \Phi(t, s+1) \mathbf{e}_s + \mathbf{e}_t,$$

for $t \geq 2$. If we now multiply this equation by $\mathbf{L}_K$, take norms on each side, and use the triangular inequality, we obtain

$$\begin{aligned} \|\mathbf{L}_K \mathbf{y}_{t+1}\|_2 &\leq \|\mathbf{L}_K \Phi(t, 1) \mathbf{y}_1\|_2 \\ &\quad + \sum_{s=1}^{t-1} \|\mathbf{L}_K \Phi(t, s+1) \mathbf{e}_s\|_2 + \|\mathbf{L}_K \mathbf{e}_t\|_2 \\ &= \sqrt{V(\Phi(t, 1) \mathbf{y}_1)} + \sum_{s=1}^{t-1} \sqrt{V(\Phi(t, s+1) \mathbf{e}_s)} + \sqrt{V(\mathbf{e}_t)}, \end{aligned} \quad (32)$$

where $V : (\mathbb{R}^d)^N \rightarrow \mathbb{R}$ is defined by

$$V(\mathbf{y}) := \|\mathbf{L}_K \mathbf{y}\|_2^2 = \sum_{i=1}^N \|\mathbf{y}^i - (\mathbf{M} \mathbf{y})^i\|_2^2.$$

Our next step is to verify the hypotheses of [27, Th. 1.2] to conclude from [27, (1.23)] that, for every $\mathbf{y} \in (\mathbb{R}^d)^N$ and every $k \geq s \geq 1$, the following holds,

$$V(\Phi(k, s) \mathbf{y}) \leq \left(1 - \frac{\tilde{\delta}}{2N^2}\right)^{\lceil \frac{k-s+1}{B} \rceil - 2} V(\mathbf{y}). \quad (33)$$

Consider the matrices $\{\mathbf{P}_t\}_{t \geq 1}$ defined in (31). Since the digraphs are weight-balanced, i.e., $\mathbf{1}_N^\top \mathbf{L}_t = 0$, we have $\mathbf{1}_N^\top \mathbf{P}_t = \mathbf{1}_N^\top$, and since $\mathbf{L}_t \mathbf{1}_N = 0$, it follows that $\mathbf{P}_t \mathbf{1}_N = \mathbf{1}_N$. Moreover, according to (24) and (25), for each $t \in \mathbb{Z}_{\geq 1}$,

$$\begin{aligned} (\mathbf{P}_t)_{ii} &\geq 1 - \hat{\sigma} d_{\text{out},t}(i) = 1 - \sigma \lambda_l(E) d_{\text{out},t}(i) \\ &\geq 1 - \sigma \lambda_{\max}(E) d_{\max} \geq \tilde{\delta}, \end{aligned}$$

for every $i \in \mathcal{I}$. On the other hand, for $i \neq j$, $(\mathbf{P}_t)_{ij} = \hat{\sigma} a_{ij,t} \geq 0$ and therefore, if $a_{ij,t} > 0$, then the nondegeneracy of the adjacency matrices together with (25) implies that

$$(\mathbf{P}_t)_{ij} = \sigma \lambda_l(E) a_{ij,t} \geq \sigma \lambda_{\min}(E) \delta \geq \tilde{\delta}.$$

Summarizing, the matrices in the sequence $\{\mathbf{P}_t\}_{t \geq 1}$ are doubly stochastic with entries uniformly bounded away from 0 by $\tilde{\delta}$

whenever positive. These are the sufficient conditions in [27, Th. 1.2], along with B -joint connectivity, to guarantee that (33) holds. Plugging (33) into (32), and noting that

$$\rho_{\tilde{\delta}} := 1 - \frac{\tilde{\delta}}{4N^2} \geq \sqrt{1 - \frac{\tilde{\delta}}{2N^2}},$$

because $(1 - x/2)^2 \geq 1 - x$ for any $x \in [0, 1]$, we get

$$\|\mathbf{L}_K \mathbf{y}_{t+1}\|_2 \leq \rho_{\tilde{\delta}}^{\lceil \frac{t}{B} \rceil - 2} \|\mathbf{y}_1\|_2 + \sum_{s=1}^t \rho_{\tilde{\delta}}^{\lceil \frac{t-s}{B} \rceil - 2} \|\mathbf{e}_s\|_2. \quad (34)$$

Here we have used that $\sqrt{V(\mathbf{y})} \leq \|\mathbf{L}_K\|_2 \|\mathbf{y}\|_2 \leq \|\mathbf{y}\|_2$ because $\|\mathbf{L}_K\|_2 = 1$ (as $\hat{\mathbf{L}}_K$ is symmetric and all its nonzero eigenvalues are equal to 1). We now proceed to bound $\|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2$ in terms of $\mathbf{v}_1$ and the inputs $\{\mathbf{d}_t\}_{t \geq 1}$ of the original dynamics (17). To do this, we rely on Lemma V.3. In fact, from (21), and using (34) for each of the K first-order consensus algorithms, we obtain

$$\begin{aligned} \|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2 &\leq \|\mathbf{S}_E\|_2 \sqrt{K}. \\ \max_{1 \leq l \leq K} \left\{ \rho_{\tilde{\delta}}^{\lceil \frac{t-1}{B} \rceil - 2} \|\mathbf{w}_1^l\|_2 + \sum_{s=1}^{t-1} \rho_{\tilde{\delta}}^{\lceil \frac{t-1-s}{B} \rceil - 2} \|\mathbf{e}_s\|_2 \right\}. \end{aligned}$$

Recalling now (18), so that $\|\mathbf{w}_1^l\|_2 \leq \|\mathbf{w}_1\|_2 \leq \|\mathbf{S}_E^{-1}\|_2 \|\mathbf{v}_1\|_2$ for each $l \in \{1, \dots, K\}$, and using also (30), we obtain

$$\begin{aligned} \|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2 &\leq \|\mathbf{S}_E\|_2 \sqrt{K} \rho_{\tilde{\delta}}^{-2} \left(\rho_{\tilde{\delta}}^{\lceil \frac{t-1}{B} \rceil} \|\mathbf{S}_E^{-1}\|_2 \|\mathbf{v}_1\|_2 \right. \\ &\quad \left. + \sum_{s=1}^{t-1} \rho_{\tilde{\delta}}^{\lceil \frac{t-1-s}{B} \rceil} \|\mathbf{S}_E^{-1}\|_2 \|\mathbf{d}_s\|_2 \right), \end{aligned} \quad (35)$$

for all $t \geq 2$ (and for $t = 1$ the inequality holds trivially). Equation (26) follows from (35) noting two facts. First, $\sum_{k=0}^{\infty} r^k = \frac{1}{1-r}$ for any $r \in (0, 1)$ and in particular for $r = \rho_{\tilde{\delta}}^{1/B}$. Second, since $\tilde{\delta} \in (0, 1)$, we have

$$\rho_{\tilde{\delta}}^{-1} = \frac{1}{1 - \tilde{\delta}/(4N^2)} \leq \frac{1}{1 - 1/(4N^2)} = \frac{4N^2}{4N^2 - 1} \leq \frac{4}{3}.$$

To obtain (28), we sum (35) over the time horizon T to get

$$\begin{aligned} \sum_{t=1}^T \|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2 &\leq \kappa(\mathbf{S}_E) \sqrt{K} \rho_{\tilde{\delta}}^{-2} \left(\frac{1}{1 - \rho_{\tilde{\delta}}^{1/B}} \|\mathbf{v}_1\|_2 \right. \\ &\quad \left. + \sum_{t=2}^T \sum_{s=1}^{t-1} \rho_{\tilde{\delta}}^{\lceil \frac{t-1-s}{B} \rceil} \|\mathbf{d}_s\|_2 \right), \end{aligned}$$

and using $r = \rho_{\tilde{\delta}}^{1/B}$ for brevity, the last sum is bounded as

$$\begin{aligned} \sum_{t=2}^T \sum_{s=1}^{t-1} r^{t-1-s} \|\mathbf{d}_s\|_2 &= \sum_{s=1}^{T-1} \sum_{t=s+1}^T r^{t-1-s} \|\mathbf{d}_s\|_2 \\ &= \sum_{s=1}^{T-1} \|\mathbf{d}_s\|_2 \sum_{t=s+1}^T r^{t-1-s} \leq \frac{1}{1-r} \sum_{s=1}^{T-1} \|\mathbf{d}_s\|_2. \end{aligned}$$

This yields (28) and the proof is complete. $\blacksquare$

The combination of the bound on the cumulative disagreement stated in Proposition V.4 with the bound on agent regret that follows from Lemmas V.1 and V.2 leads us to the next result.

Corollary V.5. (Bound on agent regret under the dynamics (6) for arbitrary learning rates): For $T \in \mathbb{Z}_{\geq 1}$, let

$\{f_t^1, \dots, f_t^N\}_{t=1}^T$ be convex functions on $\mathbb{R}^d$ with H -bounded subgradient sets. Let $E \in \mathbb{R}^{K \times K}$ be a diagonalizable matrix with real positive eigenvalues and $\{\mathcal{G}_t\}_{t \geq 1}$ a sequence of B -jointly connected, δ -nondegenerate, weight-balanced digraphs. If σ is chosen according to (25), then the agent regret associated to a sequence $\{\mathbf{x}_t = (x_t^1, \dots, x_t^N)\}_{t=1}^T$ generated by the coordination algorithm (6) satisfies, for any $j \in \{1, \dots, N\}$, $u \in \mathbb{R}^d$, and $\{\eta_t\}_{t=1}^T \subset \mathbb{R}_{>0}$, the bound

$$\begin{aligned} & 2\mathcal{R}^j(u, \{f_t\}_{t=1}^T) \\ & \leq N \sum_{t=2}^T \left\| \frac{1}{N} \sum_{i=1}^N x_t^i - u \right\|_2^2 \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} - p_t(\mathbf{u}, \mathbf{x}_t) \right) \\ & \quad + 4NH C_{\mathcal{U}} \|\mathbf{v}_1\|_2 + NH^2 (4\sqrt{N} C_{\mathcal{U}} + 1) \sum_{t=1}^T \eta_t \\ & \quad + \frac{N}{\eta_1} \left\| \frac{1}{N} \sum_{i=1}^N x_1^i - u \right\|_2^2, \end{aligned} \quad (36)$$

where $C_{\mathcal{U}}$ is given in (27) and $p_t : (\mathbb{R}^d)^N \times (\mathbb{R}^d)^N \rightarrow \mathbb{R}_{\geq 0}$ is the modulus of strong convexity of $\tilde{f}_t$.

Proof: From Lemmas V.1 and V.2, we can write

$$\begin{aligned} & 2\mathcal{R}^j(u, \{f_t\}_{t=1}^T) \\ & \leq \sum_{t=2}^T \|\mathbf{M}\mathbf{x}_t - \mathbf{u}\|_2^2 \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} - p_t(\mathbf{u}, \mathbf{x}_t) \right) \\ & \quad + (2\sqrt{N}H + 2NH) \sum_{t=1}^T \|\mathbf{L}_{\mathcal{K}}\mathbf{x}_t\|_2 + NH^2 \sum_{t=1}^T \eta_t \\ & \quad + \frac{1}{\eta_1} \|\mathbf{M}\mathbf{x}_1 - \mathbf{u}\|_2^2. \end{aligned} \quad (37)$$

On the one hand, note that

$$\|\mathbf{M}\mathbf{x}_t - \mathbf{u}\|_2^2 = N \left\| \frac{1}{N} \sum_{i=1}^N x_t^i - u \right\|_2^2. \quad (38)$$

On the other hand, taking $\mathbf{d}_s = -\eta_s(\tilde{g}_{\mathbf{x}_s}, 0, \dots, 0) \in ((\mathbb{R}^d)^N)^K$ in (28) of Proposition V.4 and noting that $\|\mathbf{d}_s\|_2 = \eta_s \|\tilde{g}_{\mathbf{x}_s}\|_2 \leq \eta_s \sqrt{N}H$, we get

$$\sum_{t=1}^T \|\mathbf{L}_{\mathcal{K}}\mathbf{x}_t\|_2 \leq \sum_{t=1}^T \|\hat{\mathbf{L}}_{\mathcal{K}}\mathbf{v}_t\|_2 \leq C_{\mathcal{U}} \left(\|\mathbf{v}_1\|_2 + \sum_{t=1}^{T-1} \eta_t \sqrt{N}H \right).$$

We obtain the result by substituting this and (38) into (37), and using the bound $2\sqrt{N}H + 2NH \leq 4NH$. ■

To establish the desired logarithmic and square-root regret bounds we need a suitable selection of learning rates in the bound obtained in Corollary V.5. This step is enabled by the final ingredient in our analysis: bounding the evolution of the online estimates and all the auxiliary states uniformly in the time horizon T . We tackle this next.

C. Bound on the trajectories uniform in the time horizon

Here we show that the trajectories of (6) are bounded uniformly in the time horizon. For this, we first bound the mean of the online estimates and then use the ISS property of the disagreement evolution studied in the previous section.

Our first result establishes a useful bound on how far from the origin one should be so that a certain important inclusion among convex cones is satisfied. This plays a key role in the technical developments of this section.

Lemma V.6. (Convex cone inclusion): Given $\beta \in (0, 1]$, $\epsilon \in (0, \beta)$, and any scalars $C_{\mathcal{X}}, C_{\mathcal{U}} \in \mathbb{R}_{>0}$, let

$$\hat{r}_{\beta} := \frac{C_{\mathcal{X}} + C_{\mathcal{U}}}{\beta \sqrt{1 - \epsilon^2} - \epsilon \sqrt{1 - \beta^2}}. \quad (39)$$

Then, $\hat{r}_{\beta} \in (C_{\mathcal{X}} + C_{\mathcal{U}}, \infty)$ and, for any $x \in \mathbb{R}^d \setminus \mathcal{B}(0, \hat{r}_{\beta})$,

$$\bigcup_{w \in \bar{\mathcal{B}}(-x, C_{\mathcal{X}} + C_{\mathcal{U}})} \mathcal{F}_{\beta}(w) \subseteq \mathcal{F}_{\epsilon}(-x), \quad (40)$$

where the set on the left is convex.

Proof: Throughout the proof, we consider the functions arccos and arcsin in the domain $[0, 1]$. Since $\epsilon \in (0, \beta)$ and $\beta \in (0, 1]$, it follows that $\arccos(\epsilon) - \arccos(\beta) \in (0, \pi/2)$. Now, using the angle-difference formula and noting that $\sin(\arccos(\alpha)) = \sqrt{1 - \alpha^2}$ for any $\alpha \in [0, 1]$, we have

$$\sin(\arccos(\epsilon) - \arccos(\beta)) = (\sqrt{1 - \epsilon^2})\beta - \epsilon\sqrt{1 - \beta^2},$$

which belongs to the set $(0, 1)$ by the observation above. Therefore, $\hat{r}_{\beta} \in (C_{\mathcal{X}} + C_{\mathcal{U}}, \infty)$. Let $x \in \mathbb{R}^d \setminus \mathcal{B}(0, \hat{r}_{\beta})$. Since $\|x\|_2 \geq \hat{r}_{\beta} > C_{\mathcal{X}} + C_{\mathcal{U}}$, then $\bar{\mathcal{B}}(-x, C_{\mathcal{X}} + C_{\mathcal{U}}) \subseteq \mathbb{R}^d \setminus \{0\}$, and the intersection of $\bigcup_{w \in \bar{\mathcal{B}}(-x, C_{\mathcal{X}} + C_{\mathcal{U}})} \mathcal{F}_{\beta}(w)$ with any plane passing through the origin and $-x$ forms a two-dimensional cone (cf. Figure 1) with angle

$$2 \arcsin\left(\frac{C_{\mathcal{X}} + C_{\mathcal{U}}}{\|x\|_2}\right) + 2 \arccos(\beta). \quad (41)$$

In the case of the intersection of $\mathcal{F}_{\epsilon}(-x)$ with any plane passing through the origin and $-x$, the angle is $2 \arccos(\epsilon)$ (which is less than π because $\epsilon < \beta \leq 1$). Now, given the axial symmetry of both cones with respect to the line passing through the origin and $-x$, (40) is satisfied if and only if

$$\arcsin\left(\frac{C_{\mathcal{X}} + C_{\mathcal{U}}}{\|x\|_2}\right) + \arccos(\beta) \leq \arccos(\epsilon), \quad (42)$$

as implied by $\|x\|_2 \geq \hat{r}_{\beta}$ because $\sin$ is increasing in $(0, \pi/2)$. On the other hand, the inclusion (40) also guarantees that $\bigcup_{w \in \bar{\mathcal{B}}(-x, C_{\mathcal{X}} + C_{\mathcal{U}})} \mathcal{F}_{\beta}(w)$ is a convex cone because each $\mathcal{F}_{\beta}(w)$ is convex, the union is taken over elements in a convex set, and (42) implies that the angle in (41) is less than π . ■

The following result bounds the mean of the online estimates for arbitrary learning rates uniformly in the time horizon.

Lemma V.7. (Uniform bound on the mean of the online estimates): For $T \in \mathbb{Z}_{\geq 1}$, let $\{f_t^1, \dots, f_t^N\}_{t=1}^T$ be convex functions on $\mathbb{R}^d$ with H -bounded subgradient sets and nonempty sets of minimizers. Let $\bigcup_{t=1}^T \bigcup_{i=1}^N \argmin(f_t^i) \subseteq \bar{\mathcal{B}}(0, C_{\mathcal{X}})$ for some $C_{\mathcal{X}} \in \mathbb{R}_{>0}$ independent of T , and assume $\{f_t^1, \dots, f_t^N\}_{t=1}^T$ are β -central on $\mathbb{R}^d \setminus \bar{\mathcal{B}}(0, C_{\mathcal{X}})$ for some $\beta \in (0, 1]$. Let $E \in \mathbb{R}^{K \times K}$ be a diagonalizable matrix with real positive eigenvalues and $\{\mathcal{G}_s\}_{s \geq 1}$ a sequence of B -jointly connected, δ -nondegenerate, weight-balanced digraphs. Let σ be chosen according to (25) and denote by $\{\mathbf{x}_t = (x_t^1, \dots, x_t^N)\}_{t=1}^T$ the sequence generated by the coordination algorithm (6). For $t \in \{1, \dots, T\}$, let $\bar{x}_t := \frac{1}{N} \sum_{i=1}^N x_t^i$ denote

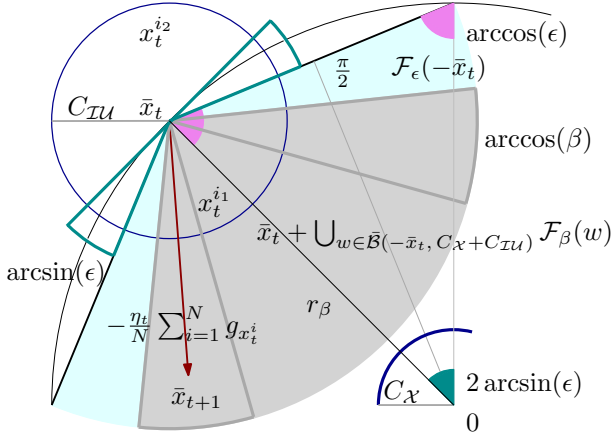


Fig. 1: Visual aid in two dimensions for the proof of Lemmas V.6 and V.7 (where the shaded cones are actually infinite).

the mean of the online estimates. Then, for any sequence of learning rates $\{\eta_t\}_{t=1}^T \subset \mathbb{R}_{>0}$,

$$\|\bar{x}_t\|_2 \leq r_\beta + H \max_{s \geq 1} \eta_s, \quad (43)$$

where, for some $\epsilon \in (0, \beta)$,

$$r_\beta := \max \left\{ \frac{C_X + C_{IU}}{\beta \sqrt{1 - \epsilon^2} - \epsilon \sqrt{1 - \beta^2}}, \frac{H}{2\epsilon} \max_{s \geq 1} \eta_s \right\} \quad (44)$$

(which is well defined as shown in Lemma V.6), and

$$C_{IU} := C_I \|v_1\|_2 + C_U \sqrt{N} H \max_{s \geq 1} \eta_s, \quad (45)$$

where C_U and C_I are given in (27).

Proof: To guide the reasoning, Figure 1 depicts some of the elements of the proof and intends to be a visual aid. The dynamics of the mean of the agents' estimates is described by (11), which in fact corresponds to N copies of

$$\bar{x}_{t+1} = \bar{x}_t - \eta_t \frac{1}{N} \sum_{i=1}^N g_{x_t^i}, \quad (46)$$

where $g_{x_t^i} \in \partial f_t^i(x_t^i)$. Our proof strategy is based on showing that, for any $t \in \{1, \dots, T\}$, if $\bar{x}_t$ belongs to the set

$$\mathbb{R}^d \setminus \mathcal{B}(0, r_\beta), \quad (47)$$

then $\|\bar{x}_{t+1}\|_2 \leq \|\bar{x}_t\|_2$. To establish this fact, we study both the direction and the magnitude of the increment $-\frac{\eta_t}{N} \sum_{i=1}^N g_{x_t^i}$ in (46). Since the subgradients in the latter expression are not evaluated at the mean, but at the agents' estimates, we first show that the agents' estimates are sufficiently close to the mean. According to the input-to-state stability property (26) from Proposition V.4 with the choice $\mathbf{d}_s = -\eta_s(\hat{g}_{x_s}, 0, \dots, 0) \in ((\mathbb{R}^d)^N)^K$, so that $\|\mathbf{d}_s\|_2 = \eta_s \|\hat{g}_{x_s}\|_2 \leq \eta_s \sqrt{N} H$, we get

$$\begin{aligned} \|\mathbf{L}_K \mathbf{x}_t\|_2 &\leq \|\hat{\mathbf{L}}_K \mathbf{v}_t\|_2 \leq C_I \|\mathbf{v}_1\|_2 \left(1 - \frac{\tilde{\delta}}{4N^2}\right)^{\lceil \frac{t-1}{B} \rceil} \\ &+ C_U \sqrt{N} H \max_{1 \leq s \leq t-1} \eta_s \leq C_{IU}, \end{aligned} \quad (48)$$

where C_{IU} is defined in (45). Hence,

$$\begin{aligned} \max_i \|x_t^i - \bar{x}_t\|_2 &\leq \left(\sum_{i=1}^N \|x_t^i - \bar{x}_t\|_2^2 \right)^{1/2} \\ &= \|\mathbf{L}_K \mathbf{x}_t\|_2 \leq C_{IU}. \end{aligned} \quad (49)$$

This allows to exploit the starting assumption that $\bar{x}_t$ belongs to the set (47) when we study the increment $-\frac{\eta_t}{N} \sum_{i=1}^N g_{x_t^i}$.

Regarding the direction of the increment $-\frac{\eta_t}{N} \sum_{i=1}^N g_{x_t^i}$, the β -centrality of the function f_t^i for each $i \in \{1, \dots, N\}$ and $t \in \{1, \dots, T\}$ on $\mathbb{R}^d \setminus \bar{\mathcal{B}}(0, C_X)$ implies that, for any $z \in \mathbb{R}^d \setminus \bar{\mathcal{B}}(0, C_X)$, we have

$$-\partial f_t^i(z) \subseteq \bigcup_{y \in \argmin(f_t^i)} \mathcal{F}_\beta(y - z) \subseteq \bigcup_{y \in \bar{\mathcal{B}}(0, C_X)} \mathcal{F}_\beta(y - z), \quad (50)$$

where the last inclusion follows from the hypothesis that $\bigcup_{t=1}^T \bigcup_{i=1}^N \argmin(f_t^i) \subseteq \bar{\mathcal{B}}(0, C_X)$. Now, using the change of variables $w := y - z$, we have

$$\bigcup_{\substack{y \in \bar{\mathcal{B}}(0, C_X) \\ z \in \bar{\mathcal{B}}(x, C_{IU})}} \mathcal{F}_\beta(y - z) = \bigcup_{w \in \bar{\mathcal{B}}(-x, C_X + C_{IU})} \mathcal{F}_\beta(w). \quad (51)$$

The representation on the right shows that the set is convex whenever x belongs to the set in (47) thanks to Lemma V.6 (essentially because $\mathcal{F}_\beta(w)$ is convex, the union is taken over elements in a convex set, and the intersection with any plane passing through $-x$ and the origin is a two-dimensional cone with angle less than π). Hence, taking the union when $z \in \bar{\mathcal{B}}(x, C_{IU})$ on both sides of (50) and using (51), we obtain

$$\begin{aligned} \text{conv} \left(\bigcup_{z \in \bar{\mathcal{B}}(x, C_{IU})} -\partial f_t^i(z) \right) &\subseteq \bigcup_{w \in \bar{\mathcal{B}}(-x, C_X + C_{IU})} \mathcal{F}_\beta(w) \\ &\subseteq \mathcal{F}_\epsilon(-x), \end{aligned}$$

where the last inclusion holds for any x in the set (47) by Lemma V.6 (noting from (39) that $r_\beta \geq \hat{r}_\beta$). Taking now $x = \bar{x}_t$ and noting that $x_t^i \in \bar{\mathcal{B}}(\bar{x}_t, C_{IU})$ by (49), we deduce

$$-\frac{1}{N} \sum_{i=1}^N g_{x_t^i} \in \text{conv} \left(\bigcup_{z \in \bar{\mathcal{B}}(\bar{x}_t, C_{IU})} -\partial f_t^i(z) \right) \subseteq \mathcal{F}_\epsilon(-\bar{x}_t).$$

This guarantees that $\bar{x}_{t+1} = \bar{x}_t - \frac{\eta_t}{N} \sum_{i=1}^N g_{x_t^i}$ is contained in a convex cone with vertex at $\bar{x}_t$ and strictly contained in the semi-space tangent to the ball $\bar{\mathcal{B}}(0, \|\bar{x}_t\|_2)$ at $\bar{x}_t$ (with a tolerance-angle between them of $\arcsin(\epsilon)$).

Regarding the magnitude $\|-\frac{\eta_t}{N} \sum_{i=1}^N g_{x_t^i}\|_2 \leq H \eta_t$, we need to show, based on the starting assumption that $\bar{x}_t$ belongs to the set (47), that $H \max_{s \geq 1} \eta_s$ is no larger than the chords of angle $\arccos(\epsilon)$ with respect to the radii of $\bar{\mathcal{B}}(0, r_\beta)$. Now, any such chord defines an isosceles triangle in the plane containing the chord and the segment joining the origin and $\bar{x}_t$. Since the angle subtended by the chord at the origin is $2 \arcsin(\epsilon)$, then the length of the chord is $2 r_\beta \epsilon$. Therefore, since $r_\beta \geq \frac{H}{2\epsilon} \max_{s \geq 1} \eta_s$ by the hypothesis (44), we conclude that the length of the chord is larger or equal than $H \max_{s \geq 1} \eta_s$. This guarantees that $\bar{x}_{t+1} = \bar{x}_t - \frac{\eta_t}{N} \sum_{i=1}^N g_{x_t^i}$ is in the ball $\bar{\mathcal{B}}(0, \|\bar{x}_t\|_2)$.

The above argument guarantees that, if the starting assumption that $\bar{x}_t$ belongs to the set (47) holds, then $\|\bar{x}_{t+1}\|_2 \leq \|\bar{x}_t\|_2$. However, if the starting assumption is not true, then the previous inequality might not hold. Since the magnitude of the increment in an arbitrary direction is upper bounded by $H \max_{s \geq 1} \eta_s$, adding this value to the threshold r_β in the definition of (47) yields the desired bound (43) for $\{\bar{x}_t\}_{t=1}^T$, uniformly in T . ■

The next result bounds the online estimates for arbitrary learning rates in terms of the initial conditions and the uniform bound on the sets of local minimizers. The fact that the bound includes the auxiliary states follows from the ISS property and the invariance of the mean of the auxiliary states.

Proposition V.8. (Boundedness of the online estimates and the auxiliary states): Under the hypotheses of Lemma V.7, the trajectories of the coordination algorithm (6) are uniformly bounded in the time horizon T , for any $\{\eta_t\}_{t=1}^T \subset \mathbb{R}_{>0}$, as

$$\|v_t\|_2 \leq C(\beta),$$

for $t \in \{1, \dots, T\}$, where

$$C(\beta) := \sqrt{N}(r_\beta + H \max_{s \geq 1} \eta_s) + \sqrt{K}\|v_1\|_2 + C_{\mathcal{U}}, \quad (52)$$

and where r_β is given in (44) and $C_{\mathcal{U}}$ in (45).

Proof: We start by noting the useful decomposition $v_t = (I_K \otimes \mathbf{M})v_t + \hat{\mathbf{L}}_K v_t$. Using the triangular inequality, we obtain

$$\begin{aligned} \|v_t\|_2 &\leq \|(I_K \otimes \mathbf{M})v_t\|_2 + \|\hat{\mathbf{L}}_K v_t\|_2 \\ &\leq \|\mathbf{M}x_t\|_2 + \sum_{l=2}^K \|\mathbf{M}v_t^l\|_2 + \|\hat{\mathbf{L}}_K v_t\|_2. \end{aligned}$$

The first term can be upper bounded by noting that $\|\mathbf{M}x_t\|_2 = \sqrt{N}\|\frac{1}{N} \sum_{i=1}^n x_t^i\|_2$ and invoking (43) in Lemma V.7. The second term does not actually depend on t . To see this, we use the fact that $(I_K \otimes \mathbf{M})(I_{KNd} - \sigma \mathbb{L}_t) = I_K \otimes \mathbf{M}$ in the dynamics (6) with the choice (7) to obtain the following invariance property of the mean of the auxiliary states,

$$\mathbf{M}v_{t+1}^l = \mathbf{M}v_t^l = \mathbf{M}v_1^l$$

for $l \in \{2, \dots, K\}$. Then, using the sub-multiplicativity of the norm and [26, Fact 9.12.22] for the norms of Kronecker products in $\|\mathbf{M}\|_2 = \|\mathbf{M} \otimes I_d\|_2 = \|\mathbf{M}\|_2 \|I_d\|_2 = 1$, we get

$$\sum_{l=2}^K \|\mathbf{M}v_1^l\|_2 \leq \|\mathbf{M}\|_2 \sum_{l=2}^K \|v_1^l\|_2 \leq \sqrt{K}\|v_1\|_2,$$

where the last inequality follows from the inequality of arithmetic and quadratic means [25]. Finally, the third term is upper bounded in (48), and the result follows. ■

The previous statements about uniform boundedness of the trajectories can also be concluded when the objectives are strongly convex. The following result says that local strong convexity and bounded subgradient sets imply β -centrality.

Lemma V.9. (Local strong convexity and bounded subgradients implies centrality away from the minimizer): Let $h : \mathbb{R}^d \rightarrow \mathbb{R}$ be a convex function on $\mathbb{R}^d$ that is also γ -strongly convex on $\bar{\mathcal{B}}(y, \zeta)$, for some $\gamma, \zeta \in \mathbb{R}_{>0}$ and $y \in \mathbb{R}^d$.

Then, for any $x \in \mathbb{R}^d \setminus \bar{\mathcal{B}}(y, \zeta)$ and $g_x \in \partial h(x)$, $g_y \in \partial h(y)$,

$$(g_x - g_y)^\top (x - y) \geq \gamma \zeta \|x - y\|_2. \quad (53)$$

If in addition h has H -bounded subgradient sets and $0 \in \partial h(y)$, then h is $\frac{\gamma \zeta}{H}$ -central in $\mathbb{R}^d \setminus \bar{\mathcal{B}}(y, \zeta)$. (Note that if $0 \in \partial h(y)$, then $\arg \min_{x \in \mathbb{R}^d} h(x) = \{y\}$ is a singleton by strong convexity in the ball $\bar{\mathcal{B}}(y, \zeta)$.)

Proof: Given any $y \in \mathbb{R}^d$ and $x \in \mathbb{R}^d \setminus \bar{\mathcal{B}}(y, \zeta)$, let $\tilde{x} \in \bar{\mathcal{B}}(y, \zeta)$ be any point in the line segment between x and y . Consequently, for some $\nu \in (0, 1)$, we can write

$$\tilde{x} - y = \nu(x - y) = \frac{\nu}{1 - \nu}(x - \tilde{x}). \quad (54)$$

Then, for any $g_x \in \partial h(x)$, $g_y \in \partial h(y)$, and $g_{\tilde{x}} \in \partial h(\tilde{x})$,

$$\begin{aligned} (g_x - g_y)^\top (x - y) &= (g_x - g_{\tilde{x}} + g_{\tilde{x}} - g_y)^\top (x - y) \\ &= \frac{1}{1 - \nu}(g_x - g_{\tilde{x}})^\top (x - \tilde{x}) + \frac{1}{\nu}(g_{\tilde{x}} - g_y)^\top (\tilde{x} - y) \\ &\geq 0 + \frac{\gamma}{\nu}\|\tilde{x} - y\|_2^2 = \gamma\|\tilde{x} - y\|_2\|x - y\|_2, \end{aligned}$$

where in the inequality we have used convexity for the first term and strong convexity for the second term. To derive (53) we choose $\tilde{x}$ satisfying $\|\tilde{x} - y\|_2 = \zeta$, while the second part follows taking $g_y = 0$ in (53) and multiplying the right-hand side by $\frac{\|g_x\|_2}{H}$ because the latter quantity is less than 1. ■

VI. LOGARITHMIC AND SQUARE-ROOT AGENT REGRET

In this section, we build on our technical results of the previous section: the general agent regret bound for arbitrary learning rates (cf. Corollary V.5), and the uniform boundedness of the trajectories of the general dynamics (6) (cf. Proposition V.8). Equipped with these results, we are ready to select the learning rates to deduce the agent regret bounds outlined in Section IV. Our first main result establishes the logarithmic agent regret for the general dynamics (6) under harmonic learning rates.

Theorem VI.1. (Logarithmic agent regret for the dynamics (6)): For $T \in \mathbb{Z}_{\geq 1}$, let $\{f_t^1, \dots, f_t^N\}_{t=1}^T$ be convex functions on $\mathbb{R}^d$ with H -bounded subgradient sets and nonempty sets of minimizers. Let $\cup_{t=1}^T \cup_{i=1}^N \arg \min(f_t^i) \subseteq \bar{\mathcal{B}}(0, C_{\mathcal{X}}/2)$ for some $C_{\mathcal{X}} \in \mathbb{R}_{>0}$ independent of T , and assume $\{f_t^1, \dots, f_t^N\}_{t=1}^T$ are p -strongly convex on $\bar{\mathcal{B}}(0, C(\frac{pC_{\mathcal{X}}}{2H}))$, for some $p \in \mathbb{R}_{>0}$, where $C(\cdot)$ is defined in (52). Let $E \in \mathbb{R}^{K \times K}$ be a diagonalizable matrix with real positive eigenvalues and $\{\mathcal{G}_t\}_{t \geq 1}$ a sequence of B -jointly connected, δ -nondegenerate, weight-balanced digraphs. Let σ be chosen according to (25) and denote by $\{x_t = (x_t^1, \dots, x_t^N)\}_{t=1}^T$ the sequence generated by the coordination algorithm (6). Then, taking $\eta_t = \frac{1}{\tilde{p}t}$, for any $\tilde{p} \in (0, p]$, the following regret bound holds for any $j \in \{1, \dots, N\}$ and $u \in \mathbb{R}^d$:

$$\begin{aligned} 2\mathcal{R}^j(u, \{f_t\}_{t=1}^T) &\leq \frac{NH^2(4\sqrt{N}C_{\mathcal{U}} + 1)}{\tilde{p}}(1 + \log T) \\ &\quad + 4NHC_{\mathcal{U}}\|v_1\|_2 + N\tilde{p}\|\frac{1}{N} \sum_{i=1}^N x_1^i - u\|_2^2, \end{aligned} \quad (55)$$

where $C_{\mathcal{U}}$ is given by (27).

Proof: First we note that $C_{\mathcal{X}} < C(\frac{pC_{\mathcal{X}}}{2H})$ because r_{β} in (44) is a lower bound for the function $C(\cdot)$ in (52) and $C_{\mathcal{X}} < r_{\beta}$ as a consequence of Lemma V.6. Thus, the fact that each f_t^i is p -strongly convex on $\bar{B}(0, C(\frac{pC_{\mathcal{X}}}{2H}))$ implies that it is also p -strongly convex on $\bar{B}(0, C_{\mathcal{X}})$. Let x_t^{*i} denote the unique minimizer of f_t^i . Then, $\arg\min(f_t^i) \subseteq \bar{B}(0, C_{\mathcal{X}}/2)$ implies that $\bar{B}(x_t^{*i}, C_{\mathcal{X}}/2) \subseteq \bar{B}(0, C_{\mathcal{X}})$. The application of Lemma V.9 with $\gamma = p$, $\zeta = C_{\mathcal{X}}/2$ and $y = x_t^{*i}$ implies then that each f_t^i is β' -central on $\mathbb{R}^d \setminus \bar{B}(0, C_{\mathcal{X}})$ for any $\beta' \leq \frac{pC_{\mathcal{X}}/2}{H}$. Hence, the hypotheses of Proposition V.8 are satisfied with $\beta = \frac{pC_{\mathcal{X}}}{2H}$ and therefore the estimates satisfy the bound $\|x_t\|_2 \leq \|v_t\|_2 \leq C(\frac{pC_{\mathcal{X}}}{2H})$ for $t \geq 1$, independent of T , which means they are confined to the region where the modulus of strong convexity of each f_t^i is p . Now, the modulus of strong convexity of $\tilde{f}_t$ is the same as for the functions $\{f_t^i\}_{i=1}^N$. That is, for each $\tilde{\xi}_y = (\xi_{y^1}, \dots, \xi_{y^N}) \in \partial \tilde{f}_t(y)$ and $\tilde{\xi}_x = (\xi_{x^1}, \dots, \xi_{x^N}) \in \partial \tilde{f}_t(x)$, for all $y, x \in \bar{B}(0, C(\frac{pC_{\mathcal{X}}}{2H})) \subset (\mathbb{R}^d)^N$, one has

$$\begin{aligned} (\tilde{\xi}_y - \tilde{\xi}_x)^\top (y - x) &= \sum_{i=1}^N (\xi_{y^i} - \xi_{x^i})^\top (y^i - x^i) \\ &\geq p \sum_{i=1}^N \|y^i - x^i\|_2^2 = p \|y - x\|_2^2. \end{aligned}$$

Thus, for all $y, x \in \bar{B}(0, C(\frac{pC_{\mathcal{X}}}{2H}))$, we can take $p_t(y, x) = p$ in (36) and hence Corollary V.5 implies the result by noting

$$\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} - p_t(u, x_t) = \tilde{p}t - \tilde{p}(t-1) - p = \tilde{p} - p \leq 0,$$

so the first sum in (36) can be bounded by 0. Finally, $\sum_{t=1}^T \eta_t = \frac{1}{\tilde{p}} \sum_{t=1}^T \frac{1}{t} < \frac{1}{\tilde{p}} (1 + \log T)$. ■

Our second main result establishes the square-root agent regret for the general dynamics (6). Its proof follows from Corollary V.5, this time by using a bounding technique called the Doubling Trick [13, Sec. 2.3.1] in the learning rates selection.

Theorem VI.2. (Square-root agent regret): For $T \in \mathbb{Z}_{\geq 1}$, let $\{f_t^1, \dots, f_t^N\}_{t=1}^T$ be convex functions on $\mathbb{R}^d$ with H -bounded subgradient sets and nonempty sets of minimizers. Let $\cup_{t=1}^T \cup_{i=1}^N \arg\min(f_t^i) \subseteq \bar{B}(0, C_{\mathcal{X}})$ for some $C_{\mathcal{X}} \in \mathbb{R}_{>0}$ independent of T , and assume $\{f_t^1, \dots, f_t^N\}_{t=1}^T$ are also β -central on $\mathbb{R}^d \setminus \bar{B}(0, C_{\mathcal{X}})$ for some $\beta \in (0, 1]$. Let $E \in \mathbb{R}^{K \times K}$ be a diagonalizable matrix with real positive eigenvalues and $\{\mathcal{G}_t\}_{t \geq 1}$ a sequence of B -jointly connected, δ -nondegenerate, weight-balanced digraphs. Let σ be chosen according to (25) and denote by $\{x_t = (x_t^1, \dots, x_t^N)\}_{t=1}^T$ the sequence generated by the coordination algorithm (6). Consider the following choice of learning rates called Doubling Trick scheme: for $m = 0, 1, 2, \dots, \lceil \log_2 T \rceil$, we take $\eta_t = \frac{1}{\sqrt{2^m}}$ in each period of 2^m rounds $t = 2^m, \dots, 2^{m+1} - 1$. Then, the following regret bound holds for any $j \in \{1, \dots, N\}$ and $u \in \mathbb{R}^d$:

$$2\mathcal{R}^j(u, \{f_t\}_{t=1}^T) \leq \frac{\sqrt{2}}{\sqrt{2}-1} \alpha \sqrt{T}, \quad (56)$$

where

$$\begin{aligned} \alpha := N^{3/2} H^2 C_{\mathcal{U}} C(\beta) &\left(\frac{4}{\sqrt{N}H} + \frac{4}{C(\beta)} + \frac{1}{\sqrt{N}C_{\mathcal{U}}C(\beta)} \right) \\ &+ N(r_{\beta} + H + \|u\|_2)^2, \end{aligned}$$

where $C_{\mathcal{U}}$ is given in (27) and $C(\cdot)$ is defined in (52).

Proof: We divide the proof in two steps. In step (i), we use the general agent regret bound of Corollary V.5 making a choice of constant learning rates over a fixed known time horizon T' . In step (ii), we use multiple times this bound together with the Doubling Trick [13, Sec. 2.3.1] to produce an implementation procedure in which no knowledge of the time horizon is required. Regarding (i), the choice $\eta_t = \eta'$ for all $t \in \{1, \dots, T'\}$ in (36) yields

$$\begin{aligned} 2\mathcal{R}^j(u, \{f_t\}_{t=1}^{T'}) &\leq 4NH C_{\mathcal{U}} \|v_1\|_2 \\ &+ NH^2 (4\sqrt{N}C_{\mathcal{U}} + 1) T' \eta' + \frac{N}{\eta'} \left\| \frac{1}{N} \sum_{i=1}^N x_1^i - u \right\|_2^2, \end{aligned} \quad (57)$$

where the first sum in (36) is upper-bounded by 0 because $\frac{1}{\eta'} - \frac{1}{\eta'} - p_t(u, x_t) \leq 0$. Taking now $\eta' = 1/\sqrt{T'}$ in (57), factoring out $\sqrt{T'}$ and using $1 \leq \sqrt{T'}$, we obtain

$$\begin{aligned} 2\mathcal{R}^j(u, \{f_t\}_{t=1}^{T'}) &\leq \left(4NH C_{\mathcal{U}} \|v_1\|_2 \right. \\ &\left. + NH^2 (4\sqrt{N}C_{\mathcal{U}} + 1) + N \left\| \frac{1}{N} \sum_{i=1}^N x_1^i - u \right\|_2^2 \right) \sqrt{T'}. \end{aligned} \quad (58)$$

This bound is of the type $2\mathcal{R}^j(u, \{f_t\}_{t=1}^{T'}) \leq \alpha' \sqrt{T'}$, where α' depends on the initial conditions. This leads to step (ii). According to the Doubling Trick [13, Sec. 2.3.1], for $m = 0, 1, \dots, \lceil \log_2 T \rceil$, the dynamics is executed in each period of $T' = 2^m$ rounds $t = 2^m, \dots, 2^{m+1} - 1$, where at the beginning of each period the initial conditions are the final values in the previous period. The regret bound for each period is $\alpha' \sqrt{T'} = \alpha_m \sqrt{2^m}$, where α_m is the multiplicative constant in (58) that depends on the initial conditions in the corresponding period. To eliminate the dependence on the latter, by Proposition V.8, we have that $\|v_t\|_2 \leq C(\beta)$, for $C(\cdot)$ in (52) with $\max_{s \geq 1} \eta_s = 1$. Also, using (43), we have

$$\left\| \frac{1}{N} \sum_{i=1}^N x_t^i - u \right\|_2 \leq \|\bar{x}_t\|_2 + \|u\|_2 \leq r_{\beta} + H + \|u\|_2.$$

Since $C(\beta)$ only depends on the initial conditions at the beginning of the implementation procedure, the regret on each period is now bounded as $\alpha_m \sqrt{2^m} \leq \alpha \sqrt{2^m}$, for α in the statement. Consequently, the total regret can be bounded by

$$\sum_{m=0}^{\lceil \log_2 T \rceil} \alpha \sqrt{2^m} = \alpha \frac{1 - \sqrt{2}^{\lceil \log_2 T \rceil + 1}}{1 - \sqrt{2}} \leq \alpha \frac{1 - \sqrt{2}^T}{1 - \sqrt{2}} \leq \frac{\sqrt{2}}{\sqrt{2}-1} \alpha \sqrt{T},$$

which yields the desired bound. ■

Remark VI.3. (Asymptotic dependence of logarithmic agent regret bound on network properties): Here we analyze the asymptotic dependence of the logarithmic regret bound in Theorem VI.1 on the number of agents. It is not difficult to see that, when $N \rightarrow \infty$, then

$$\frac{C_{\mathcal{U}}}{C_{\mathcal{I}}} = \frac{1}{1 - (1 - \frac{\delta}{4N^2})^{1/B}} \sim \frac{4N^2 B}{\delta}.$$

Hence, for any B that guarantees B -joint connectivity, the asymptotic behavior as $N \rightarrow \infty$ of the bound (55) scales as

$$\frac{N^{3+1/2}B}{\tilde{\delta}}o(T), \quad (59)$$

where $\lim_{T \rightarrow \infty} \frac{o(T)}{T} = 0$. In contrast to (59), the asymptotic dependence on the number of agents in [20], [21], which assume strong connectivity every time step and a doubly stochastic adjacency matrix A , is

$$\frac{N^{1+1/2}}{1 - \sigma_2(A)}o(T), \quad (60)$$

where $\sigma_2(A)$ is the second smallest singular value of A . (Here we are taking into account the fact that [21] divides the regret by the number of agents.) The bounds (59) and (60) are comparable in the case of sparse connected graphs that fail to be good expanders (i.e., for sparse graphs with low algebraic connectivity given by the second smallest eigenvalue of the Laplacian). This is the most reasonable comparison given our joint-connectivity assumption. For simplicity, we examine the case of undirected graphs because then $1 - \sigma_2(A) = 1 - \lambda_2(A) = \lambda_2(L)$, where $L = \text{diag}(A\mathbf{1}_N) - A = I - A$ is the Laplacian corresponding to A . The sparsity of the graph implies that $\delta, d_{\max} \approx 1$, so that one can compute the maximum feasible $\tilde{\delta}$ from (23) to be $\tilde{\delta}^* := (1 + \frac{\lambda_{\max}(E)d_{\max}}{\lambda_{\min}(E)\delta})^{-1} \approx (1 + \lambda_{\max}(E)/\lambda_{\min}(E))^{-1}$. The algebraic connectivity $\lambda_2(L)$ can vary even for sparse graphs. Paths and cycles are two examples of graphs that fail to be expander graphs and their algebraic connectivity [28] is $2(1 - \cos(\pi/N))$ and $2(1 - \cos(2\pi/N))$, respectively (for edge-weights equal to 1), and thus proportional to $1 - \cos(1/N) \sim \frac{1}{N^2}$ when $N \rightarrow \infty$. With these values of $\tilde{\delta}^*$ and $\lambda_2(L)$ (up to a constant independent of N), (59) and (60) become

$$N^{3+1/2}B o(T) \quad \text{and} \quad N^{3+1/2}o(T),$$

respectively. Expression (59) highlights the trade-offs between the degree of parallelization and the regret behavior for a given time horizon. Such trade-offs must be considered in the light of factors like the serial processor speed and the rate of data-collection as well as the cost and bandwidth limitations of transmitting spatially distributed data. •

VII. SIMULATION: APPLICATION TO MEDICAL DIAGNOSIS

In this section we illustrate the performance of the coordination algorithm (6) in a binary classification problem from medical diagnosis. We specifically consider the online gradient descent with proportional and with proportional-integral disagreement feedback. Inspired by [29], we consider a clinical decision problem involving the use of Computerized Tomography (CT) for patients with minor head injury. We consider a network of hospitals that works cooperatively to develop a set of rules to determine whether a patient requires immediately a CT for possible neurological intervention, or if an alternative follow-up protocol should be applied to further inform the decision. The hospitals estimate local prediction models using the data collected from their patients while

coordinating their efforts according to (6) to benefit from the model parameters updated by other hospitals.

We start by describing the data collected by the hospitals. Suppose that at round t , hospital i collects a vector $w_t^i \in \mathbb{R}^c$ encoding a set of features corresponding to patient data. In our case, $c = 10$ and the components of w_t^i correspond to factors or symptoms like “age”, “amnesia before impact”, “open skull fracture”, “loss of consciousness”, “vomiting”, etc. The ultimate goal of each hospital is to decide if any acute brain finding would be revealed by the CT, and the true answer is denoted by $y_t^i \in \{-1, 1\}$, where $-1 = \text{“no”}$ and $1 = \text{“yes”}$ are the two possible classes. The true assessment is only found once the CT or the follow-up protocol have been used.

To cast this scenario in the networked online optimization framework described in Section III, it is enough to specify the cost function $f_t^i : \mathbb{R}^d \rightarrow \mathbb{R}$ for each hospital $i \in \{1, \dots, N\}$ and each round $t \in \{1, \dots, T\}$. In this scenario, the cost function measures the fitness of the model parameters estimated by the hospital with respect to the data collected from its patients, as we explain next. Each hospital i seeks to estimate a vector of model parameters $x_t^i \in \mathbb{R}^d$, $d = c + 1$, that weigh the correspondence between the symptoms and the actual brain damage (up to an additional affine term). More precisely, hospital i employs a model h to assign the quantity $h(x_t^i, w_t^i)$, called decision or prediction, to the data point w_t^i using the estimated model parameters x_t^i . For instance, a linear predictor is based on the model $h(x_t^i, w) = x_t^{i\top}(w_t^i, 1)$, with the corresponding class predictor being $\text{sign}(h(x_t^i, w_t^i))$. The loss incurred by hospital i is then $f_t^i(x_t^i) = l(x_t^i, w_t^i, y_t^i)$, where the loss function l is decreasing in the so-called margin $y_t^i h(x_t^i, w_t^i)$. This is because correct predictions (when the margin is positive) should be penalized less than incorrect predictions (when the margin is negative). Common loss functions are the logistic (smooth) function, $l(x, w, y) = \log(1 + e^{-2y h(x, w)})$ or the hinge (nonsmooth) function, $l(x, w, y) = \max\{0, 1 - y h(x, w)\}$.

In the scenario just described, each hospital $i \in \{1, \dots, N\}$ updates to x_{t+1}^i its estimated model parameters x_t^i according to the dynamics (6) as the data (w_t^i, y_t^i) becomes available. We simulate here two cases, the online gradient descent with proportional disagreement feedback, corresponding to $K = 1$ and $E = [1]$, and the online gradient descent with proportional-integral disagreement feedback, corresponding to

$$K = 2 \quad \text{and} \quad E = \begin{bmatrix} a & 1 \\ -1 & 0 \end{bmatrix};$$

cf. Remark IV.1. Both the online and distributed aspects of our approach are relevant for this kind of large-scale supervised learning. On one hand, data streams can be analyzed rapidly and with low storage to produce a real-time service using first-order information of the corresponding cost functions (for single data points or for mini-batches). On the other hand, hospitals can benefit from the prediction models updated by other hospitals. Under (6), each hospital i only shares the provisional vector of model parameters x_t^i with neighboring hospitals and maintains its patient data (w_t^i, y_t^i) private. In addition, the joint connectivity assumption is a flexible con-

dition on how frequently hospitals communicate with each other. With regards to communication latency, note that the potential delays in communication among hospitals are small compared to the rate at which data is collected from patients. Also, the fitness of provisional local models can always be computed with respect to mini-batches of variable size when one hospital collects a different amount of data than others in the given time scales of coordination.

In our simulation, a network of 5 hospitals uses the time-varying communication topology shown in Figure 2. This results in the executions displayed in Figure 3, where provisional local models are shown to asymptotically agree and achieve sublinear regret with respect to the best model obtained in hindsight with all the data centrally available. For completeness, the plots also compare their performance against a centralized online gradient descent algorithm [11], [16].

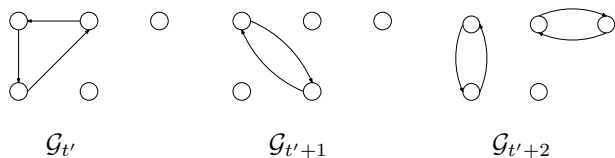


Fig. 2: The communication topology corresponds to the periodic repetition of the displayed sequence of weight-balanced digraphs (where all nonzero edge weights are 1). The resulting sequence is 3-jointly connected, 1-nondegenerate, and the maximum out-degree is 1, i.e., $B = 3$, $\delta = 1$, and $d_{\max} = 1$.

VIII. CONCLUSIONS

We have studied a networked online convex optimization scenario where each agent has access to partial information that is increasingly revealed over time in the form of a local cost function. The goal of the agents is to generate a sequence of decisions that achieves sublinear regret with respect to the best single decision in hindsight had all the information been centrally available. We have proposed a class of distributed coordination algorithms that allow agents to fuse their local decision parameters and incorporate the information of the local objectives as it becomes available. Our algorithm design uses first-order local information about the cost functions revealed in the previous round, in the form of subgradients, and only requires local communication of decision parameters among neighboring agents over a sequence of weight-balanced, jointly connected digraphs. We have shown that our distributed strategies achieve the same logarithmic and square-root agent regret bounds that centralized implementations enjoy. We have also characterized the dependence of the agent regret bounds on the network parameters. Our technical approach has built on an innovative combination of network and agent regret bounds, the cumulative disagreement of the collective estimates, and the boundedness of the sequence of collective estimates uniformly in the time horizon. Future work will include the refinement of the regret bounds when partial knowledge about the evolution of the cost functions is available, the study of the impact of practical implementation considerations such as disturbances, noise, communication

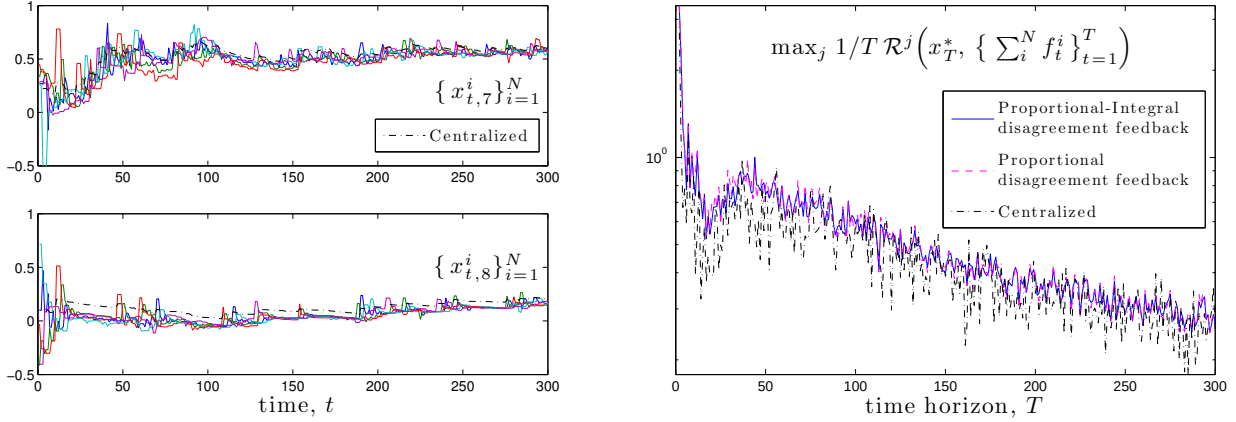
delays, and asynchronism in the algorithm performance, and the application to large-scale learning scenarios involving the distributed interaction of many users and devices.

ACKNOWLEDGMENTS

The authors thank the anonymous reviewers for their useful feedback that helped us improved the presentation of the paper.

REFERENCES

- [1] D. P. Bertsekas and J. N. Tsitsiklis, *Parallel and Distributed Computation: Numerical Methods*. Athena Scientific, 1997.
- [2] J. N. Tsitsiklis, D. P. Bertsekas, and M. Athans, "Distributed asynchronous deterministic and stochastic gradient optimization algorithms," *IEEE Transactions on Automatic Control*, vol. 31, no. 9, pp. 803–812, 1986.
- [3] A. Nedic and A. Ozdaglar, "Distributed subgradient methods for multi-agent optimization," *IEEE Transactions on Automatic Control*, vol. 54, no. 1, pp. 48–61, 2009.
- [4] B. Johansson, M. Rabi, and M. Johansson, "A randomized incremental subgradient method for distributed optimization in networked systems," *SIAM Journal on Control and Optimization*, vol. 20, no. 3, pp. 1157–1170, 2009.
- [5] J. C. Duchi, A. Agarwal, and M. J. Wainwright, "Dual averaging for distributed optimization: convergence analysis and network scaling," *IEEE Transactions on Automatic Control*, vol. 57, no. 3, pp. 592–606, 2012.
- [6] M. Zhu and S. Martínez, "On distributed convex optimization under inequality and equality constraints," *IEEE Transactions on Automatic Control*, vol. 57, no. 1, pp. 151–164, 2012.
- [7] B. Ghahsifard and J. Cortés, "Distributed continuous-time convex optimization on weight-balanced digraphs," *IEEE Transactions on Automatic Control*, vol. 59, no. 3, pp. 781–786, 2014.
- [8] D. P. Bertsekas, "Incremental gradient, subgradient, and proximal methods for convex optimization: a survey," in *Optimization for Machine Learning* (S. Sra, S. Nowozin, and S. J. Wright, eds.), pp. 85–120, Cambridge, MA: MIT Press, 2011.
- [9] J. Wang and N. Elia, "A control perspective for centralized and distributed convex optimization," in *IEEE Conf. on Decision and Control*, (Orlando, Florida), pp. 3800–3805, 2011.
- [10] D. Mateos-Núñez and J. Cortés, "Noise-to-state stable distributed convex optimization on weight-balanced digraphs," in *IEEE Conf. on Decision and Control*, (Florence, Italy), pp. 2781–2786, 2013.
- [11] M. Zinkevich, "Online convex programming and generalized infinitesimal gradient ascent," in *Proceedings of the Twentieth International Conference on Machine Learning*, (Washington, D.C.), pp. 928–936, 2003.
- [12] N. Cesa-Bianchi and G. Lugosi, *Prediction, Learning, and Games*. Cambridge University Press, 2006.
- [13] S. Shalev-Shwartz, *Online Learning and Online Convex Optimization*, vol. 12 of *Foundations and Trends in Machine Learning*. Now Publishers Inc, 2012.
- [14] E. Hazan, "The convex optimization approach to regret minimization," in *Optimization for Machine Learning* (S. Sra, S. Nowozin, and S. J. Wright, eds.), pp. 287–304, Cambridge, MA: MIT Press, 2011.
- [15] S. Shalev-Shwartz and Y. Singer, "Convex repeated games and fenchel duality," in *Advances in Neural Information Processing Systems* (B. Schölkopf, J. Platt, and T. Hoffman, eds.), vol. 19, (Cambridge, MA), MIT Press, 2007.
- [16] E. Hazan, A. Agarwal, and S. Kale, "Logarithmic regret algorithms for online convex optimization," *Machine Learning*, vol. 69, no. 2-3, pp. 169–192, 2007.
- [17] H. Wang and A. Banerjee, "Online alternating direction method," in *International Conference on Machine Learning*, (Edinburgh, Scotland), pp. 1119–1126, July 2012.
- [18] O. Dekel, R. Gilad-Bachrach, O. Shamir, and L. Xiao, "Optimal distributed online prediction using mini-batches," *Journal of Machine Learning Research*, vol. 13, no. Jan, pp. 165–202, 2012.
- [19] M. Raginsky, N. Kiarashi, and R. Willett, "Decentralized online convex programming with local information," in *American Control Conference*, (San Francisco, CA), pp. 5363–5369, 2011.



(a) Agents' estimates using proportional-integral disagreement feedback

(b) Average regret

Fig. 3: Performance of the online gradient descent algorithms with proportional and proportional-integral disagreement feedback (with $a = 4$ in the latter). The dynamics involves $N = 5$ agents communicating over the periodic sequence of digraphs displayed in Figure 2. Each local objective $f_t^i : \mathbb{R}^d \rightarrow \mathbb{R}$, with $d = 11$, is given by $f_t^i(x) = l(x, w_t^i, y_t^i)$ with loss function $l(x, w, y) = \log(1 + e^{-2yx^\top(w,1)})$, for the data set from <http://www.stats4stem.org/r-headinjury-data.html>. Since $w \in \{0, 1\}^{d-1}$ and $y \in \{-1, 1\}$, we have $\|\partial_x l(x, w, y)\| \leq \|-2y(w, 1)\| \leq 2d$, so the local cost functions are globally Lipschitz with $H = 2d$. In both cases the learning rates are $\eta_t = 1/\sqrt{t}$ (with the same asymptotic behavior as for the Doubling Trick scheme employed in Theorem VI.2). With $\tilde{\delta}' = 0.01$ in (23), so that $\tilde{\delta} = \tilde{\delta}'$, equation (25) yields $\sigma \in (0.01/\lambda_{\min}(E), 0.99/\lambda_{\max}(E))$, so we take $\sigma = 0.1$ for both dynamics. The initial condition x_1 is randomly generated and, in the second-order case, we take $z_1 = \mathbb{1}_5 \otimes \mathbb{1}_{11}$. Plot (a), top, shows the evolution of the 7th coordinate of each agent's estimate, which is the gain associated to the feature “assessed by clinician as high risk for neurological intervention,” versus the evolution of the centralized estimate. The centralized estimate is computed by the centralized online gradient descent $c_{t+1} = c_t - \eta_t \frac{1}{N} \sum_{i=1}^N \nabla f_t^i(c_t)$ with the same learning rates. Plot (a), bottom, shows a similar comparison for the gain associated to the feature “loss of consciousness.” Plot (b) depicts the temporal average regret of the two distributed dynamics versus the centralized online gradient descent algorithm. The scale is logarithmic and the evolutions are bounded by a line of negative slope, as should correspond to a regret bound proportional to $\log(\sqrt{T}/T) = -\frac{1}{2} \log T$. The global optimal solution in hindsight, x_T^* , for each time horizon T , is computed offline using centralized gradient descent. (As a side note, the agent regret of an algorithm can be sublinear regardless of the design of the cost functions, which ultimately determines the pertinence of the centralized model in hindsight and hence the pertinence of the online distributed performance.)

- [20] F. Yan, S. Sundaram, S. V. N. Vishwanathan, and Y. Qi, “Distributed autonomous online learning: regrets and intrinsic privacy-preserving properties,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 25, no. 11, pp. 2483–2493, 2013.
- [21] S. Hosseini, A. Chapman, and M. Mesbahi, “Online distributed optimization via dual averaging,” in *IEEE Conf. on Decision and Control*, (Florence, Italy), 2013.
- [22] Y. Nesterov, *Introductory Lectures on Convex Optimization: A Basic Course*, vol. 87 of *Applied Optimization*. Springer, 2004.
- [23] F. Bullo, J. Cortés, and S. Martínez, *Distributed Control of Robotic Networks*. Applied Mathematics Series, Princeton University Press, 2009. Electronically available at <http://coordinationbook.info>.
- [24] K. I. Tsianos and M. G. Rabbat, “Distributed strongly convex optimization,” 2012. arXiv:1207.3031v2.
- [25] P. S. Bullen, *Handbook of Means and Their Inequalities*, vol. 560 of *Mathematics and Its Applications*. Dordrecht, The Netherlands: Kluwer Academic Publishers, 2003.
- [26] D. S. Bernstein, *Matrix Mathematics*. Princeton, NJ: Princeton University Press, 2005.
- [27] A. Nedić and A. Ozdalgar, “Cooperative distributed multi-agent optimization,” in *Convex Optimization in Signal Processing and Communications* (Y. Eldar and D. Palomar, eds.), Cambridge University Press, 2010.
- [28] M. Fiedler, “Algebraic connectivity of graphs,” *Czechoslovak Mathematical Journal*, vol. 23, no. 2, pp. 298–305, 1973.
- [29] I. G. Stiell, G. A. Wells, K. Vandemheen, C. Clement, H. Lesiuk, A. Laupacis, R. D. McKnight, R. Verbeek, R. Brison, D. Cass, M. A. Eisenhauer, G. H. Greenberg, and J. Worthington, “The Canadian CT

Head Rule for patients with minor head injury,” *The Lancet*, vol. 357, no. 9266, pp. 1391–1396, 2001.



David Mateos-Núñez received the Licenciatura degree in mathematics with distinction from the Universidad Autónoma de Madrid, Spain, in 2011. He then entered the Ph.D. program in Mechanical and Aerospace Engineering at the University of California, San Diego, under supervision from Professor Jorge Cortés. He received his M.S. in Mechanical Engineering in 2012 and is currently working towards finishing his Ph.D. His research focuses on distributed optimization and networked systems.



Jorge Cortés received the Licenciatura degree in mathematics from the Universidad de Zaragoza, Spain, in 1997, and the Ph.D. degree in engineering mathematics from the Universidad Carlos III de Madrid, Spain, in 2001. He held postdoctoral positions at the University of Twente, The Netherlands, and the University of Illinois at Urbana-Champaign, USA. He was an Assistant Professor with the Department of Applied Mathematics and Statistics at the University of California, Santa Cruz from 2004 to 2007. He is currently a Professor in the Department of Mechanical and Aerospace Engineering, University of California, San Diego. His research interests include cooperative control, network science, distributed decision making, game theory, and geometric mechanics.