

This is the peer reviewed version of the following article:

Dynamic Scheduling for Over-the-Air Federated Edge Learning with Energy Constraints / Sun, Y., Zhou, S., Niu, Z., Gunduz, D.. - In: IEEE JOURNAL ON SELECTED AREAS IN COMMUNICATIONS. - ISSN 0733-8716. - 40:1(2022), pp. 227-242. [10.1109/JSAC.2021.3126078]

Terms of use:

The terms and conditions for the reuse of this version of the manuscript are specified in the publishing policy. For all terms of use and more information see the publisher's website.

15/09/2026 13:27

Dynamic Scheduling for Over-the-Air Federated Edge Learning with Energy Constraints

Yuxuan Sun, *Member, IEEE*, Sheng Zhou, *Member, IEEE*,
Zhisheng Niu, *Fellow, IEEE*, Deniz Gündüz, *Senior Member, IEEE*

Abstract

Machine learning and wireless communication technologies are jointly facilitating an intelligent edge, where federated edge learning (FEEL) is a promising training framework. As wireless devices involved in FEEL are resource limited in terms of communication bandwidth, computing power and battery capacity, it is important to carefully schedule them to optimize the training performance. In this work, we consider an over-the-air FEEL system with analog gradient aggregation, and propose an energy-aware dynamic device scheduling algorithm to optimize the training performance under energy constraints of devices, where both communication energy for gradient aggregation and computation energy for local training are included. The consideration of computation energy makes dynamic scheduling challenging, as devices are scheduled before local training, but the communication energy for over-the-air aggregation depends on the l_2 -norm of local gradient, which is known after local training. We thus incorporate estimation methods into scheduling to predict the gradient norm. Taking the estimation error into account, we characterize the performance gap between the proposed algorithm and its offline counterpart. Experimental results show that, under a highly unbalanced local data distribution, the proposed algorithm can increase the accuracy by 4.9% on CIFAR-10 dataset compared with the myopic benchmark, while satisfying the energy constraints.

Index Terms

Y. Sun, S. Zhou and Z. Niu are with the Beijing National Research Center for Information Science and Technology, Department of Electronic Engineering, Tsinghua University, Beijing 100084, China (e-mail: sunyuxuan@tsinghua.edu.cn, sheng.zhou@tsinghua.edu.cn, niuzhs@tsinghua.edu.cn).

D. Gündüz is with the Department of Electrical and Electronic Engineering, Imperial College London, London SW7 2BT, UK (e-mail: d.gunduz@imperial.ac.uk).

Part of this work has been presented in IEEE ICC 2020 [1].

Federated edge learning, over-the-air computation, energy constraints, dynamic scheduling, Lyapunov optimization.

I. INTRODUCTION

Many emerging applications at the wireless edge, such as autonomous driving, virtual reality and Internet of things (IoT), are powered by modern machine learning (ML) techniques. Data-driven approaches also penetrate into the wireless network itself for channel estimation, encoding and decoding, resource allocation, etc. [2], [3]. The complex ML models for these applications need to be trained over massive datasets, while data samples are usually generated by edge devices. Traditional centralized training methods can hardly be competent, as collecting data at one location would create network congestion, lead to extremely high transmission cost and may cause privacy concerns. On the other hand, computing capabilities of base stations (BSs) and edge devices, such as mobile phones, smart vehicles and IoT sensors, are becoming increasingly powerful, enabling intensive computations at the edge [4]. In this context, federated learning (FL) is considered as a promising training framework that can exploit distributed data and computational resources with limited communication and privacy leakage [5]–[7]. In FL, multiple devices train a shared model collaboratively with local data, and a central *parameter server (PS)* coordinates training and aggregates global model periodically.

The limited communication resource and non-independent and identically distributed (i.i.d.) data, i.e., the distribution of local data at one device is not identical with that of other devices or the global data, are the two major challenges in FL [8], [9]. Current methods to improve the communication efficiency of FL mainly include model compression [12]–[15], device scheduling [16], [17], and enabling multiple local iterations [15], [18]. Under non-i.i.d. data, the training performance can be improved by sharing global i.i.d. data with devices [9] or the PS [19], introducing data redundancy [1], or scheduling devices based on their importance [17].

In a wireless network, FL can be carried out among wireless edge devices coordinated by a BS, called federated edge learning (FEEL). In FEEL, participating devices are often resource limited in terms of wireless bandwidth, computing capability and battery capacity. A key challenge is to design device scheduling and resource allocation algorithms that optimize the training performance under device energy constraints and training delay budget. Considering the communication energy constraints, an energy-efficient bandwidth allocation policy is proposed

to maximize the fraction of scheduled devices in [20], while an online algorithm is designed to maximize the sum utility of scheduling in [21]. Due to the timeliness requirements of FEEL tasks at the wireless edge [22], training delay is also becoming a key performance metric. The total communication delay for training is minimized by joint device selection and wireless resource allocation in [23], while the total training delay taking into account both local computations and model transmission is minimized in [24] by balancing the trade-off between the average delay per round and the total number of rounds required for convergence. Communication delay is combined with the importance of each update for probabilistic scheduling in [25]. A hierarchical FEEL framework is proposed in [26], where the end-to-end training delay is minimized by the joint optimization of update interval and model compression. The trade-off between total energy consumption for communication and computation and the training delay is further considered in [27]–[29], yielding a joint design of local computation speed and wireless resource allocation.

The literature above mainly focuses on the implementation of FEEL via digital wireless communications. However, the unique communication requirement of FEEL, i.e., the PS only needs the *average* of local model updates rather than each individual vector, makes the separate design of learning and communication protocol highly suboptimal [30]. A new solution called *over-the-air computation* is facilitated to further improve the communication efficiency [31]–[35], which is achieved by synchronizing the devices to transmit their local gradients or models in an analog fashion, and exploiting the superposition property of a wireless multiple access channel (MAC) to do the summation over-the-air. Power limits of devices can highly degrade the training performance, which yields the design of power allocation schemes over noisy channels [32], fading channels [33] and broadband fading channels [34]. Power control algorithms that take into account the importance of updates [36], uplink and downlink noise [37], [38], gradient statistics [39] and non-i.i.d. data [40] are further proposed. While over-the-air FEEL requires accurate channel state information (CSI), it is shown in [41] that multiple antennas can help to relax the CSI requirement. The impact of imperfect CSI or synchronization across devices is considered in [42], and a digital realization of over-the-air FEEL is further proposed in [43], based on one-bit gradient quantization and majority voting.

Existing papers on over-the-air FEEL mainly consider average power constraints for communication, but have not considered the computation energy for local model training, which is in fact non-negligible for edge devices. In this work, we aim to optimize the training perfor-

mance under total energy constraints of devices by designing an energy-aware dynamic device scheduling algorithm, where energy is consumed for both communication and computation. The introduction of computation energy makes the scheduling decisions challenging due to the *causality of decision making and energy consumption*. This is because, in over-the-air FEEL, the communication energy of each device for gradient aggregation depends on the l_2 -norm of its local gradient estimate, which can only be obtained *after computation*. However, online scheduling decision should be made at the start of each training round *before computation*. Note that this issue does not arise in the case of digital communication, as the transmission power can be chosen independently of the local update.

The main contributions of this work are summarized as follows:

- 1) We characterize the convergence bound of the considered over-the-air FEEL system, based on which we formulate a device scheduling problem to optimize the training performance under the total energy budget of each device. Both the communication energy for gradient aggregation and the computation energy for local gradient calculation are included.
- 2) Due to the unavailability of future system states, we design an energy-aware dynamic device scheduling algorithm based on Lyapunov optimization, where a virtual queue is constructed to indicate the up-to-date energy deficit and enable online decision making.
- 3) To further address the challenge that communication energy is unknown at the device scheduling point, we propose estimation methods to predict the l_2 -norm of the local gradients upon scheduling, and characterize the theoretical performance guarantee of the proposed scheduling algorithm by taking the error of energy estimation into consideration.
- 4) Experiments on MNIST and CIFAR-10 datasets validate that the proposed dynamic device scheduling algorithm can achieve higher model accuracies compared with the myopic benchmark, while satisfying the energy limits. Under a highly-non-i.i.d. scenario, the accuracy can be increased by 4.9%. The impact of design parameters on the training performance and energy consumption are also evaluated to provide guidelines for practical implementations.

The rest of this paper is organized as follows. In Section II, we introduce the system model and problem formulation. In Section III, we carry out convergence analysis. The energy-aware dynamic device scheduling algorithm is developed in Section IV with its performance guarantee. Experimental results are shown in Section V, and conclusions are given in Section VI.

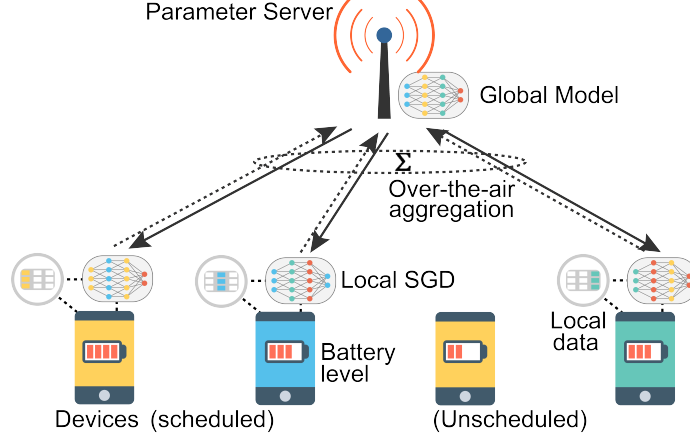


Fig. 1. Illustration of the considered over-the-air FEEL system.

II. SYSTEM MODEL AND PROBLEM FORMULATION

A. System Overview

As shown in Fig. 1, we consider a FEEL system with one PS and N devices, denoted by $\mathcal{N} = \{1, 2, \dots, N\}$. Each device $n \in \mathcal{N}$ has a local dataset $\mathcal{D}_n$ with D data samples, and the global dataset is denoted by $\mathcal{D} = \bigcup_{n=1, \dots, N} \mathcal{D}_n$ with ND data samples.

Given a single data sample $\mathbf{x} \in \mathcal{D}$, a loss function $f(\mathbf{w}, \mathbf{x})$ is used to measure the fitting performance of an s -dimensional model vector $\mathbf{w} \in \mathbb{R}^s$. At device $n \in \mathcal{N}$, the local loss function $F_n(\mathbf{w})$ is defined as the average loss over local data samples, i.e.,

$$F_n(\mathbf{w}) \triangleq \frac{1}{D} \sum_{\mathbf{x} \in \mathcal{D}_n} f(\mathbf{w}, \mathbf{x}). \quad (1)$$

The goal of the FEEL task is to train a shared global model $\mathbf{w}$ that minimizes the global loss function $F(\mathbf{w})$, which is defined as

$$F(\mathbf{w}) \triangleq \frac{1}{N} \sum_{n=1}^N F_n(\mathbf{w}) = \frac{1}{ND} \sum_{n=1}^N \sum_{\mathbf{x} \in \mathcal{D}_n} f(\mathbf{w}, \mathbf{x}). \quad (2)$$

Under the coordination of the PS, the FEEL system iterates the following three steps until the termination condition is satisfied: 1) the PS broadcasts the up-to-date global model to a subset of devices, which are scheduled to participate in the current training process; 2) the scheduled devices compute their local gradients with local datasets; and 3) the PS aggregates the local gradients over a wireless MAC and updates the global model. Each iteration consisting of these three steps is called a training round, which is indexed by t in the following. The termination conditions that are commonly used for FEEL include the convergence of the global model, or

reaching a preset maximum number of training rounds. Since we consider an energy-limited wireless scenario, we set the total number of training rounds to T .

B. Local Gradient Computation

At the start of the t -th training round, the PS schedules a subset of devices $\mathcal{B}_t \subseteq \mathcal{N}$, and broadcasts the global model vector $\mathbf{w}_{t-1}$ obtained in the last round to these scheduled devices. Let $\beta_{n,t} \in \{0, 1\}$ be an indicator, with $\beta_{n,t} = 1$ if device n is scheduled to participate in the t -th training round, and $\beta_{n,t} = 0$ otherwise. Thus $\mathcal{B}_t = \{n | \beta_{n,t} = 1, n \in \mathcal{N}\}$. We also assume that the broadcast of $\mathbf{w}_{t-1}$ is error-free since the PS is a more capable node with sufficient power.

Each scheduled device $n \in \mathcal{B}_t$ computes the local gradient estimate $\tilde{\mathbf{g}}_{n,t}$ by running the stochastic gradient descent (SGD) algorithm on a local mini-batch $\mathcal{L}_{n,t} \subseteq \mathcal{D}_n$, according to

$$\tilde{\mathbf{g}}_{n,t} = \frac{1}{L_b} \sum_{\mathbf{x} \in \mathcal{L}_{n,t}} \nabla f(\mathbf{w}_{t-1}, \mathbf{x}), \quad (3)$$

where $L_b = |\mathcal{L}_{n,t}| \leq D$ is the batch size, and $\mathcal{L}_{n,t}$ is uniformly selected at random from the local dataset $\mathcal{D}_n$. We remark here that, a single-iteration gradient update is considered in this work, but the proposed algorithm can be extended to a more general case where multiple local iterations are carried out in each training round. Also note that, the batch size is considered as a hyper-parameter rather than an optimization variable, and set to an identical value across devices in this work, since local data might be non-i.i.d. across devices and local training should guarantee the fairness by exploiting the same amount of data.

We assume that for device n , the computation energy for calculating the local gradient on a single data sample is e_n , which can be estimated according to the number of floating point operations (FLOPs) of the ML model and the computation frequency of the device [29]. Therefore, the computation energy consumption $E_{n,t}^{[\text{cp}]}$ at device n in round t is given by

$$E_{n,t}^{[\text{cp}]} = e_n L_b. \quad (4)$$

C. Gradient Aggregation Over-the-Air

We assume that the devices transmit their local gradients over a noisy wireless MAC in an analog fashion for global gradient aggregation. To enable the summation of local gradients over-the-air, transmissions are synchronized across all the scheduled devices, and the transmit power

of each device is aligned with the others. To be specific, let $h_{n,t}$ be the wireless channel gain between device n and the PS, which is assumed to remain constant during one transmission period. Note that, the device scheduling policy designed in this work is applicable to arbitrary channel models. Moreover, as local gradient computation takes time and the wireless channel is time variant, the channel gain observed at the start of each training round may not be precise. The observation error, i.e., the difference between the observed channel gain that determines the device scheduling and its true value during transmission, will also be considered in the following. Let σ_t be the power scalar that determines the received SNR at the PS. Then the transmit power $p_{n,t}$ of each scheduled device $n \in \mathcal{B}_t$ is set to

$$p_{n,t} = \frac{\sigma_t}{h_{n,t}}, \quad (5)$$

and $p_{n,t}\tilde{\mathbf{g}}_{n,t}$ is transmitted from device n to the PS. The communication energy consumption $E_{n,t}^{[\text{tr}]}$ at device n in round t is then given by

$$E_{n,t}^{[\text{tr}]} = \|p_{n,t}\tilde{\mathbf{g}}_{n,t}\|_2^2 = \frac{\sigma_t^2}{h_{n,t}^2} \|\tilde{\mathbf{g}}_{n,t}\|_2^2, \quad (6)$$

where $\|\mathbf{x}\|_2$ represents the l_2 -norm of vector $\mathbf{x}$. Therefore, if device n is scheduled in the t -th round, the total energy consumption $E_{n,t}$ for computation and communication is

$$E_{n,t} = E_{n,t}^{[\text{tr}]} + E_{n,t}^{[\text{cp}]} = \frac{\sigma_t^2}{h_{n,t}^2} \|\tilde{\mathbf{g}}_{n,t}\|_2^2 + e_n L_b. \quad (7)$$

At the PS, the received signal $\mathbf{y}_t$ is given by

$$\mathbf{y}_t = \sum_{n \in \mathcal{B}_t} h_{n,t} p_{n,t} \tilde{\mathbf{g}}_{n,t} + \mathbf{z}_t = \sigma_t \sum_{n \in \mathcal{B}_t} \tilde{\mathbf{g}}_{n,t} + \mathbf{z}_t, \quad (8)$$

where $\mathbf{z}_t \in \mathbb{R}^s$ is an additive white Gaussian noise vector, in which each entry is i.i.d. and follows Gaussian distribution with zero mean and variance σ_0^2 .

In FEEL, we aim to update the global model vector $\mathbf{w}_t$ according to

$$\mathbf{w}_t = \mathbf{w}_{t-1} - \eta_t \frac{\sum_{n \in \mathcal{B}_t} \tilde{\mathbf{g}}_{n,t}}{|\mathcal{B}_t|}, \quad (9)$$

where η_t is the learning rate in the t -th training round, and $|\cdot|$ denotes the cardinality of a set. Due to channel noise, the actual global model is updated according to

$$\mathbf{w}_t = \mathbf{w}_{t-1} - \eta_t \frac{\mathbf{y}_t}{\sigma_t |\mathcal{B}_t|} = \mathbf{w}_{t-1} - \eta_t \left(\frac{\sum_{n \in \mathcal{B}_t} \tilde{\mathbf{g}}_{n,t}}{|\mathcal{B}_t|} + \frac{\mathbf{z}_t}{\sigma_t |\mathcal{B}_t|} \right). \quad (10)$$

D. Problem Formulation

Given the total number of training rounds T and the initial global model vector $\mathbf{w}_0$, we aim to minimize the expected global loss $\mathbb{E}[F(\mathbf{w}_T)]$ under the energy constraints of devices, by optimizing the device scheduling $\{\beta_{n,t}\}$ and power scalar $\{\sigma_t\}$. The expectation $\mathbb{E}[F(\mathbf{w}_T)]$ is taken over the randomness of channel noise and data sampling for local SGD. The problem is formulated as

$$\mathcal{P}1 : \min_{\{\sigma_t, \beta_{n,t}\}_{t=1}^T} \mathbb{E}[F(\mathbf{w}_T)] \quad (11a)$$

$$\text{s.t.} \quad \sum_{t=1}^T \beta_{n,t} E_{n,t} \leq \bar{E}_n, \quad \forall n, \quad (11b)$$

$$\sigma_t > 0, \quad \beta_{n,t} \in \{0, 1\}, \quad \forall n, t. \quad (11c)$$

In the first constraint (11b), $\bar{E}_n$ represents the total energy budget of device n , and the inequality indicates that for each device, the total energy consumption for both local gradient computation and wireless communication over T training rounds cannot exceed its given budget. The second constraint limits the ranges of optimization variables.

Based on the law of telescoping sums, problem $\mathcal{P}1$ can be re-written as

$$\mathcal{P}2 : \min_{\{\sigma_t, \beta_{n,t}\}_{t=1}^T} \sum_{t=1}^T \mathbb{E}[F(\mathbf{w}_t)] - \mathbb{E}[F(\mathbf{w}_{t-1})] \quad (12)$$

$$\text{s.t.} \quad \text{constraints (11b), (11c).}$$

There are three major challenges to solve problem $\mathcal{P}2$:

1) *The inexplicit form of the objective function:* Since the neural network architectures for ML might be deep and diverse, and the evolution of the model vector is very complex during the training process, it is hard to express $\mathbb{E}[F(\mathbf{w}_T)]$ or $\mathbb{E}[F(\mathbf{w}_t)] - \mathbb{E}[F(\mathbf{w}_{t-1})]$ in closed form.

2) *The unavailability of future information:* Optimally solving problem $\mathcal{P}2$ requires the system states of future training rounds due to the existence of total energy constraints, which are not available in practice. Therefore, we aim to design an online scheduling algorithm in this work, which only relies on the system states in the current training round.

3) *The causality of decision making and energy consumption:* A unique characteristic of over-the-air FEEL is that the communication energy (6) depends on the l_2 -norm of local gradient through $\|\tilde{\mathbf{g}}_{n,t}\|_2^2$, which can only be acquired after computing the gradient in each round.

However, online device scheduling decision should be made before gradient computation, in order not to consume computation energy at unscheduled devices, or even not to transmit global updates to these devices. This means that the exact energy consumption in the current training round is unknown upon decision making. Moreover, the channel gain $h_{n,t}$ observed at the start of each round may not be precise, making the estimation of communication energy more inaccurate.

To address these challenges, we first substitute the objective function with its upper bound based on the convergence analysis in Section III. Then in Section IV, we design an online device scheduling algorithm based on Lyapunov optimization, where the unknown instantaneous states for decision making, including the l_2 -norm of local gradients and the wireless channel gains, are substituted with their estimates, and in particular, the impact of the estimation error on the performance of the proposed algorithm is analyzed.

III. CONVERGENCE ANALYSIS AND PROBLEM TRANSFORMATION

In this section, we provide an upper bound for the objective function in problem $\mathcal{P}2$ based on the convergence analysis, and transform the original optimization problem to an alternative form with explicit expressions.

For the simplicity of notation, we define the local full gradient on device n in the t -th round as $\mathbf{g}_{n,t} \triangleq \nabla F_n(\mathbf{w}_{t-1}) = \frac{1}{D} \sum_{\mathbf{x} \in \mathcal{D}_n} \nabla f(\mathbf{w}_{t-1}, \mathbf{x})$, the global full gradient in round t as $\mathbf{g}_t \triangleq \nabla F(\mathbf{w}_{t-1}) = \frac{1}{N} \sum_{n=1}^N \mathbf{g}_{n,t}$, and the optimum loss as $F^* \triangleq \min_{\mathbf{w} \in \mathbb{R}^s} F(\mathbf{w})$.

To facilitate the convergence analysis, we make the following assumptions according to the state-of-the-art literature, including [12]–[15], [32]–[34], [43], etc.

Assumption 1. *Stochastic gradient is unbiased and variance-bounded, i.e., for any device n and training round t , taking the expectation over stochastic data sampling, we have*

$$\mathbb{E}_{\mathbf{x}_n} [\nabla f(\mathbf{w}_{t-1}, \mathbf{x}_n)] = \mathbb{E}_{\mathcal{L}_{n,t}} [\tilde{\mathbf{g}}_{n,t}] = \mathbf{g}_t, \quad \mathbb{E}_{\mathbf{x}_n} [\|\nabla f(\mathbf{w}_{t-1}, \mathbf{x}_n) - \mathbf{g}_t\|_2^2] \leq G^2, \forall n, t, \quad (13)$$

where $\mathbf{x}_n \in \mathcal{D}_n$ is a data sample, $\mathcal{L}_{n,t} \subseteq \mathcal{D}_n$ is a stochastic mini-batch, and G is a constant.

Assumption 2. *Loss functions $F_1(\mathbf{w}), \dots, F_N(\mathbf{w})$ are l -smooth, i.e., for $\forall \mathbf{v}, \mathbf{w} \in \mathbb{R}^s$ and $n \in \mathcal{N}$,*

$$F_n(\mathbf{v}) - F_n(\mathbf{w}) \leq \nabla F_n^T(\mathbf{w})(\mathbf{v} - \mathbf{w}) + \frac{l}{2} \|\mathbf{v} - \mathbf{w}\|_2^2. \quad (14)$$

Assumption 3. Loss functions $F_1(\mathbf{w}), \dots, F_N(\mathbf{w})$ are μ -strongly convex, i.e., for $\forall \mathbf{v}, \mathbf{w} \in \mathbb{R}^s$ and $n \in \mathcal{N}$,

$$F_n(\mathbf{v}) - F_n(\mathbf{w}) \geq \nabla F_n^T(\mathbf{w})(\mathbf{v} - \mathbf{w}) + \frac{\mu}{2} \|\mathbf{v} - \mathbf{w}\|_2^2. \quad (15)$$

A. Convergence Analysis

Based on the assumptions above, we provide a single-round convergence guarantee in the following lemma by characterizing the upper bound of $\mathbb{E}[F(\mathbf{w}_t)] - \mathbb{E}[F(\mathbf{w}_{t-1})]$.

Lemma 1. Given the global model vector $\mathbf{w}_{t-1}$ and the set of scheduled devices $\mathcal{B}_t$ at the beginning of round t , the single-round convergence is upper-bounded by

$$\mathbb{E}[F(\mathbf{w}_t)] - \mathbb{E}[F(\mathbf{w}_{t-1})] \leq -\eta_t \left(1 - \frac{l\eta_t}{2}\right) \|\mathbf{g}_t\|_2^2 + \frac{l\eta_t^2}{2} \left(\frac{G^2}{L_b|\mathcal{B}_t|} + \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2} \right), \quad (16)$$

where the expectation is taken over the randomness of channel noise and SGD.

Proof. See Appendix A. □

Based on Lemma 1, the convergence performance of over-the-air FEEL after T training rounds is given in the following theorem.

Theorem 1. Given the global model vector $\mathbf{w}_0$ and any device scheduling sequence $\{\mathcal{B}_t, t = 1, \dots, T\}$, after T rounds of training,

$$\mathbb{E}[F(\mathbf{w}_T)] - F^* \leq (\mathbb{E}[F(\mathbf{w}_0)] - F^*) \prod_{i=1}^T (1 - \mu\eta_i) + \sum_{i=1}^{T-1} A_i \prod_{j=i+1}^T (1 - \mu\eta_j) + A_T, \quad (17)$$

where $A_t \triangleq \frac{\eta_t}{2} \left(\frac{G^2}{L_b|\mathcal{B}_t|} + \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2} \right)$ and the learning rate satisfies $\eta_t \leq \min\{\frac{1}{l}, 1\}$, $\forall t$.

Proof. See Appendix B. □

According to Lemma 1 and Theorem 1, we can see that the number of devices $|\mathcal{B}_t|$ scheduled in each round makes a key contribution to the convergence rate of training. While existing papers [1], [20], [21] maximize the weighted fraction of devices scheduled over time, we provide a more reasonable objective function based on the theoretical characterization. We also remark that, although Assumption 1 indicates i.i.d. local data, our proposed algorithm can also work well under non-i.i.d. data as being validated in the experiments in Section V.

B. Problem Transformation

As discussed in Section II-D, the objective function $\sum_{t=1}^T \mathbb{E}[F(\mathbf{w}_t)] - \mathbb{E}[F(\mathbf{w}_{t-1})]$ in problem $\mathcal{P}2$ cannot be expressed explicitly. To make the optimization problem tractable, we substitute the objective function with its convergence bound according to Lemma 1, and formulate an alternative optimization problem:

$$\begin{aligned} \mathcal{P}3 : \quad & \min_{\{\sigma_t, \beta_{n,t}\}_{t=1}^T} \sum_{t=1}^T -\eta_t \left(1 - \frac{l\eta_t}{2}\right) \|\mathbf{g}_t\|_2^2 + \frac{l\eta_t^2}{2} \left(\frac{G^2}{L_b \sum_{n=1}^N \beta_{n,t}} + \frac{\sigma_0^2 s}{\sigma_t^2 \left(\sum_{n=1}^N \beta_{n,t}\right)^2} \right) \\ \text{s.t.} \quad & \text{constraints (11b), (11c),} \end{aligned} \quad (18)$$

where we recall that σ_t and $\beta_{n,t}$ are the power scalar and worker scheduling indicator, respectively.

Moreover, due to the unavailability of future system states, we aim to design an online algorithm to solve problem $\mathcal{P}3$, and ignore the impact of current decision on the future system states. As the global full gradient $\mathbf{g}_t$ defined on the whole dataset is fixed given the global model vector $\mathbf{w}_{t-1}$ at the start of training round t , and the learning rate η_t and smoothness parameter l are hyper-parameters, the first term in (18) is a constant. Therefore, we ignore this term and transform the optimization problem to

$$\begin{aligned} \mathcal{P}4 : \quad & \min_{\{\sigma_t, \beta_{n,t}\}_{t=1}^T} \sum_{t=1}^T \frac{l\eta_t^2}{2} \left(\frac{G^2}{L_b \sum_{n=1}^N \beta_{n,t}} + \frac{\sigma_0^2 s}{\sigma_t^2 \left(\sum_{n=1}^N \beta_{n,t}\right)^2} \right) \\ \text{s.t.} \quad & \text{constraints (11b), (11c).} \end{aligned} \quad (19)$$

IV. ENERGY-AWARE DYNAMIC DEVICE SCHEDULING ALGORITHM

In this section, we propose an energy-aware dynamic device scheduling algorithm that solves problem $\mathcal{P}4$ in an online fashion. To address the challenge brought by the causality of decision making and communication energy consumption, we first propose two heuristics to estimate the l_2 -norm of local gradient estimates. Then, we design an online scheduling algorithm based on Lyapunov optimization, and characterize the worst-case performance of the proposed algorithm, which takes the error of energy estimation into consideration. Finally, we provide some practical considerations for real implementations.

A. Estimating the l_2 -Norm of Local Gradients

We propose two heuristics in the following to estimate the l_2 -norm of local gradients $\|\tilde{\mathbf{g}}_{n,t}\|_2^2$ at the start of each training round t .

1) Compute the l_2 -norm of local gradients with a smaller mini-batch (EST-C):

An additional step is introduced at the start of each training round. To be specific, the PS broadcasts the up-to-date global model $\mathbf{w}_{t-1}$ to all the devices. Then, each device randomly selects a mini-batch $\mathcal{L}'_{n,t} \subseteq \mathcal{D}_n$ with batch size L_e to calculate a local gradient estimate $\tilde{\mathbf{g}}_{n,t}^{[\text{est}]}$:

$$\tilde{\mathbf{g}}_{n,t}^{[\text{est}]} = \frac{1}{L_e} \sum_{\mathbf{x} \in \mathcal{L}'_{n,t}} \nabla f(\mathbf{w}_{t-1}, \mathbf{x}). \quad (20)$$

The computation energy should be modified as $E_{n,t}^{[\text{cp}]} = \beta_{n,t} e_n (L_b - L_e) + e_n L_e$.

We further assume that each device can upload the value of $\|\tilde{\mathbf{g}}_{n,t}^{[\text{est}]}\|_2^2$ to the PS with negligible cost, which is used as the estimation of the l_2 -norm of local gradient for device n in round t .

Let $G_{n,t}^2$ be the variance of the stochastic gradient on a single data sample for device n in round t . Following the same proof of Lemma 1 as (38), the exact value of gradient norm and its estimation have the following expectations:

$$\mathbb{E} [\|\tilde{\mathbf{g}}_{n,t}\|_2^2] = \|\mathbf{g}_{n,t}\|_2^2 + \frac{G_{n,t}^2}{L_b}, \quad \mathbb{E} [\|\tilde{\mathbf{g}}_{n,t}^{[\text{est}]}\|_2^2] = \|\mathbf{g}_{n,t}\|_2^2 + \frac{G_{n,t}^2}{L_e}. \quad (21)$$

As L_b is typically much larger than L_e , the expressions above indicate that the estimation may suffer from a large deviation due to the gradient variance.

2) Estimate with past information (EST-P):

A simpler and more straightforward way is to use the most recent l_2 -norm of local gradient to estimate the current one at each device. Let $t_n = \arg \max_t \{\beta_{n,t} = 1\}$ be the most recent round in which device n is scheduled. The estimated l_2 -norm of the current local gradient estimate is:

$$\|\tilde{\mathbf{g}}_{n,t}^{[\text{est}]}\|_2^2 = \|\tilde{\mathbf{g}}_{n,t_n}\|_2^2. \quad (22)$$

We will show in Fig. 2 and Fig. 3 in Section V that, under our considered datasets, estimation by EST-P is more accurate due to the strong temporal correlation of gradients. Moreover, compared with EST-C that requires additional computation and communication, EST-P method is computation-free and only needs each device to report the l_2 -norm of local gradients when it is scheduled. Therefore, we use EST-P to estimate the l_2 -norm of local gradient in the following device scheduling algorithm.

B. Energy-Aware Dynamic Device Scheduling Algorithm

To enable online scheduling without any future information while satisfying the total energy constraints of devices, we construct a virtual queue $q_{n,t}$ for each device n to indicate the gap between the cumulative energy consumption till round t and the budget, evolved as

$$q_{n,t+1} = \max \left\{ q_{n,t} + \beta_{n,t} E_{n,t} - \frac{\bar{E}_n}{T}, 0 \right\}, \quad (23)$$

with initial value $q_{n,1} = 0, \forall n \in \mathcal{N}$.

Recall that the causality of device scheduling and energy consumption leads to the unawareness of $E_{n,t}$ at the start of round t . Based on the estimated l_2 -norm of local gradient $\|\tilde{\mathbf{g}}_{n,t}^{[\text{est}]}\|_2^2$ by EST-P and the wireless channel gain $\tilde{h}_{n,t}$ observed at the beginning of round t , the estimated energy consumption of device n at round t , denoted by $\tilde{E}_{n,t}$, is given by

$$\tilde{E}_{n,t} = \frac{\sigma_t^2}{\tilde{h}_{n,t}^2} \|\tilde{\mathbf{g}}_{n,t}^{[\text{est}]}\|_2^2 + e_n L_b. \quad (24)$$

Let $U_t \triangleq \frac{\ln \eta_t^2}{2} \left(\frac{G^2}{L_b \sum_{n=1}^N \beta_{n,t}} + \frac{\sigma_0^2 s}{\sigma_t^2 (\sum_{n=1}^N \beta_{n,t})^2} \right)$. Inspired by the drift-plus-penalty algorithm of Lyapunov optimization [44], the online scheduling aims to solve the following problem:

$$\mathcal{P5} : \min_{\{\sigma_t, \beta_{n,t}\}} V U_t + \sum_{n=1}^N \beta_{n,t} q_{n,t} \tilde{E}_{n,t} \quad (25a)$$

$$\text{s.t. } \sigma_t > 0, \beta_{n,t} \in \{0, 1\}, \forall n, \quad (25b)$$

where V is an adjustable weight parameter to balance the loss U_t and energy consumption. Compared to the classical drift-plus-penalty algorithm where all the states in the current round are known, the drift term $q_{n,t} \tilde{E}_{n,t}$ in $\mathcal{P5}$ is an approximation, and thus we call it *estimated-drift-plus-penalty* algorithm.

Notice that problem $\mathcal{P5}$ is a mixed integer non-linear programming problem, which is still very difficult to solve. Meanwhile, existing work has shown that the convergence performance of FEEL with over-the-air gradient aggregation is not very sensitive to the power scalar σ_t , as long as the received SNR or the power limit of each device is larger than a threshold [32]. Therefore, we further decouple the optimization variables in $\mathcal{P5}$ by considering the power scalar σ_t as a hyper-parameter, and then develop the optimal solution to the online device scheduling problem.

1) Received SNR and Power Scalar σ_t

The power scalar σ_t is chosen as follows. In the t -th round, the expected received SNR at the PS side is denoted by γ_t , given by

$$\gamma_t = \mathbb{E} \left[\frac{\left\| \sigma_t \sum_{n \in \mathcal{B}_t} \tilde{\mathbf{g}}_{n,t} \right\|_2^2}{\left\| \mathbf{z}_t \right\|_2^2} \right] = \frac{\sigma_t^2}{\sigma_0^2 s} \left\| \sum_{n \in \mathcal{B}_t} \tilde{\mathbf{g}}_{n,t} \right\|_2^2. \quad (26)$$

Let γ_0 be a pre-defined SNR threshold. The power scalar is set according to

$$\sigma_t = \frac{\gamma_0 \sigma_0^2 \sqrt{s}}{\min_{n \in \mathcal{N}} \left\| \tilde{\mathbf{g}}_{n,t} \right\|_2} \approx \frac{\gamma_0 \sigma_0^2 \sqrt{s}}{\min_{n \in \mathcal{N}} \left\| \tilde{\mathbf{g}}_{n,t}^{[\text{est}]} \right\|_2}, \quad (27)$$

such that the expectation of the received SNR can meet the threshold γ_0 even in the worst case when a single device is scheduled. Recall that $\left\| \tilde{\mathbf{g}}_{n,t} \right\|_2$ is unknown and thus approximated by $\left\| \tilde{\mathbf{g}}_{n,t}^{[\text{est}]} \right\|_2$ according to the EST-P method.

2) Optimal Online Device Scheduling

Given the power scalar σ_t , the device scheduling $\{\beta_{n,t}\}$ in the t -th round aims to solve

$$\mathcal{P6}: \min_{\{\beta_{n,t}\}} VU_t + \sum_{n=1}^N \beta_{n,t} q_{n,t} \tilde{E}_{n,t} \quad (28a)$$

$$\text{s.t. } \beta_{n,t} \in \{0, 1\}, \forall n. \quad (28b)$$

An optimal solution to problem $\mathcal{P6}$ is shown in Algorithm 1. In Line 1, we sort $\mathcal{C}_t = \{q_{n,t} \tilde{E}_{n,t}, \forall n\}$ in the ascending order, and let $C_t^{[m]}$ be the m -th smallest value of $\mathcal{C}_t$. Many sorting algorithms such as Heapsort or Mergesort can be used, with worst-case complexity $O(N \log N)$. In Lines 2-4, we iterate over the possible number of scheduled devices $k \in \{1, \dots, N\}$, and calculate the corresponding minimum estimated-drift-plus-penalty $v_t(k)$ according to

$$v_t(k) \triangleq \frac{l\eta_t^2}{2} \left(\frac{G^2}{L_b k} + \frac{\sigma_0^2 s}{\sigma_t^2 k^2} \right) + \sum_{n=1}^k C_t^{[k]}. \quad (29)$$

The optimal number of devices k^* to be scheduled is obtained by finding the minimum $v_t(k)$ according to Line 5, and k^* devices with smallest estimated drift $q_{n,t} \tilde{E}_{n,t}$ are scheduled, as shown Lines 6-8. Besides Line 1, all the other steps are with complexity $O(N)$, and thus the complexity of Algorithm 1 is $O(N \log N)$.

3) The Complete Algorithm

The proposed energy-aware dynamic device scheduling algorithm is summarized in Algorithm 2. In the t -th training round, the PS makes device scheduling decision by solving $\mathcal{P6}$ based on the estimated energy consumption and the virtual queue, which is run in an online fashion without

Algorithm 1 Optimal Online Device Scheduling to $\mathcal{P}6$

- 1: Sort $\mathcal{C}_t = \{q_{n,t}\tilde{E}_{n,t}, \forall n\}$ and let $C_t^{[m]}$ be the m -th smallest value of $\mathcal{C}_t$.
 - 2: **for** $k = 1, \dots, N$ **do**
 - 3: Calculate $v_t(k)$ according to (29).
 - 4: **end for**
 - 5: Get $k^* = \arg \min_k \{v_t(k) \mid k = 1, \dots, N\}$.
 - 6: **for** $n = 1, \dots, N$ **do**
 - 7: Let $\beta_{n,t} = 1$ if $q_{n,t}\tilde{E}_{n,t} \leq C_t^{[k^*]}$, and $\beta_{n,t} = 0$ otherwise.
 - 8: **end for**
-

Algorithm 2 Energy-Aware Dynamic Device Scheduling Algorithm

- 1: **Initialization:** initialize global model $\mathbf{w}_0$. Each device n runs local SGD according to (3) to report $\|\tilde{\mathbf{g}}_{n,0}\|_2^2$ to the PS, and let $q_{n,1} = 0$.
 - 2: **for** $t = 1, \dots, T$ **do**
 - 3: The PS set σ_t according to (27), acquires channel gains $\tilde{h}_{n,t}$ and calculates the estimated energy consumption $\tilde{E}_{n,t}$ according to (24) for all devices.
 - 4: The PS schedules a subset of devices $\mathcal{B}_t$ by solving $\mathcal{P}6$ according to Algorithm 1.
 - 5: The PS broadcasts $\mathbf{w}_{t-1}$ and σ_t to the scheduled devices $n \in \mathcal{B}_t$.
 - 6: Each scheduled device $n \in \mathcal{B}_t$ updates local gradient $\tilde{\mathbf{g}}_{n,t}$ according to (3), and transmits $\frac{\sigma_t}{h_{n,t}}\tilde{\mathbf{g}}_{n,t}$ simultaneously with all the other scheduled devices.
 - 7: The PS receives $\mathbf{y}_t$ and updates the global model $\mathbf{w}_t$ according to (10).
 - 8: Each scheduled device $n \in \mathcal{B}_t$ reports $E_{n,t}$ and the PS updates the virtual queue $q_{n,t}$ for all devices according to (23).
 - 9: **end for**
-

any future information. The weight parameter V and the virtual queue states $\{q_{n,t}, \forall n\}$ jointly balance the training gain of the FEEL task and the energy consumption of devices. In particular, a larger V puts more emphasis on scheduling more devices so as to accelerate the convergence rate. Meanwhile, a larger $q_{n,t}$ indicates that the cumulative energy consumption of device n till the current round far exceeds the budget, so that the device tends to save energy. As shown in Algorithm 1, the optimal solution to $\mathcal{P}6$ also indicates that devices with smaller values of

$q_{n,t}\tilde{E}_{n,t}$ are always scheduled first, as their energy is relatively sufficient.

Then, the up-to-date global model vector $\mathbf{w}_{t-1}$ is broadcast to the scheduled devices, who run local SGD to compute local gradients $\tilde{\mathbf{g}}_{n,t}$ in parallel. After computation, local gradients are aggregated over-the-air and the global model is updated by the PS. Finally, the PS collects the actual energy consumption of each scheduled device, which also contains the information of local gradient norm $\|\tilde{\mathbf{g}}_{n,t}\|_2$, and updates the virtual queue states for all the devices to guide the scheduling decision in the next training round.

C. Performance Analysis

The performance of the proposed dynamic device scheduling algorithm is characterized by comparing with its optimal offline counterpart, which is achieved by solving the optimal device scheduling $\{\beta_{n,t}^*\}$ to problem $\mathcal{P}4$ while taking $\{\sigma_t\}$ as a pre-defined hyper-parameter sequence. Let $\sum_{t=1}^T U_t^*$ be the offline optimal cumulative loss of $\mathcal{P}4$ by scheduling device sequence $\{\beta_{n,t}^*\}$, and define $\sum_{t=1}^T U_t^\dagger$ as the cumulative loss of the proposed algorithm, which is achieved by solving the online device scheduling problem $\mathcal{P}6$ in each round. To enable the theoretical analysis, we neglect the impact of current scheduling decision on the future gradient norm for the offline counterpart. Meanwhile, we do not limit the distributions of wireless channel or local gradient norms, which can be non-stationary over time.

The performance guarantee of the proposed algorithm is shown in the following theorem.

Theorem 2. *Compared to the offline optimal solution, the cumulative loss of Algorithm 2 can be bounded by*

$$\sum_{t=1}^T U_t^\dagger \leq \sum_{t=1}^T U_t^* + \frac{\theta_0 T^2 + T(T-1)\delta_0 \sum_{n=1}^N \theta_n}{V}, \quad (30)$$

and the total energy consumption of Algorithm 2 can be bounded by

$$\sum_{t=1}^T \beta_{n,t} E_{n,t} \leq \bar{E}_n + \sqrt{2V \sum_{t=1}^T U_t^* + 2\theta_0 T^2 + 2T(T-1)\delta_0 \sum_{n=1}^N \theta_n}, \quad (31)$$

where $\delta_0 \triangleq \max_{\{n,t\}} \left\{ \left| \tilde{E}_{n,t} - E_{n,t} \right| \right\}$, $\theta_0 \triangleq \sum_{n=1}^N \frac{1}{2} \theta_n^2$ and $\theta_n \triangleq \max_t \left\{ \left| E_{n,t} - \frac{\bar{E}_n}{T} \right| \right\}$.

Proof. See Appendix C. □

Theorem 2 shows that, the training performance of the proposed energy-aware dynamic device scheduling algorithm can be bounded with respect to its optimal offline counterpart, while the deviation between the cumulative energy consumption of each device and its budget is also bounded. The worst-case performance can be improved by reducing the upper bound of the energy overuse θ_n and the maximum energy estimation error δ_0 . Moreover, the trade-off between the training performance of the FEEL task and maximum energy consumption of each device can be balanced by the weight parameter V .

We also remark here that, compared to the state-of-the-art that also applies Lyapunov optimization to solve scheduling problem under energy constraints [1], [21], [45], our analysis further shows the impact of estimation error on the performance bound.

D. Implementation Issues

To enable the efficient implementation of the proposed algorithm in a real system, we provide some practical considerations as follows.

1) Communication Rescheduling

The key motivation of rescheduling is to avoid using significantly more energy than expected when the estimation error of $\tilde{E}_{n,t}$ is large. To be specific, after local gradient computation, each scheduled device can learn its exact energy consumption $E_{n,t}$ by calculating the local gradient norm $\|\tilde{\mathbf{g}}_{n,t}\|_2^2$ and acquiring the accurate channel gain $h_{n,t}$. If $E_{n,t} - \tilde{E}_{n,t} \leq \delta_h$, where $\delta_h > 0$ is a given threshold, then the device is scheduled for gradient aggregation. Otherwise, the device backs off from the communication step.

2) Minimum Value of Virtual Queue

The typical evolution of virtual queue is given in (23), in which the minimum queue value is set to 0. In problem $\mathcal{P}6$, $q_{n,t} = 0$ indicates that the energy consumption is not considered in the current scheduling round, and thus the device is scheduled. However, the energy consumption $E_{n,t}$ might be large, leading to a large deviation θ_n and thus a poor worst-case performance. To avoid such cases, we instead set $q_{\min} > 0$ as the minimum value of the virtual queue in practice.

3) Estimations of Smoothness Parameter l and Variance Bound G^2

Our algorithm is designed based on the convergence analysis under Assumptions in Section III. These hyper-parameters should be estimated in practice. According to the definition of

smoothness, l is estimated by the maximum value of $\frac{\|\tilde{g}_{n,t} - \tilde{g}_{n,t-1}\|}{\|\mathbf{w}_{t-1} - \mathbf{w}_{t-2}\|}$ during training, while each device can count the variance of local gradients to set a reasonable variance bound G^2 .

V. EXPERIMENTS

In this section, we evaluate the proposed energy-aware dynamic device scheduling algorithm for an image classification task using both MNIST¹ and CIFAR-10² datasets. We consider $N = 10$ devices and both i.i.d. and non-i.i.d. datasets on devices. For the i.i.d. case, the training dataset of MNIST with 60000 samples (or CIFAR-10 with 50000 samples) is randomly partitioned into N disjoint subsets, and each device holds one subset. For the non-i.i.d. case, we sort the data samples by their labels, and each device holds a disjoint subset of data with m labels (represented by ‘non-i.i.d. (m)’ in the following). Note that the data distributions are more skewed for smaller m , and they become i.i.d. when m is equal to the total number of classes in the dataset.

For MNIST, we train a multilayer perceptron (MLP) which has a 784-unit input layer with ReLU activation, a 64-unit hidden layer, and a 10-unit softmax output layer, with 50890 parameters in total. The total number of rounds is set to $T = 200$, and 10 local iterations are carried out per round with batch size $L_b = 64$. In each round, the total computation energy is 1J for each device. For CIFAR-10, we train a convolutional neural network (CNN) with the following structure: two 3×3 convolution layers each with 32 channels and followed by a 2×2 max pooling layer, two 3×3 convolution layers each with 64 channels and followed by a 2×2 max pooling layer, a fully connected layer with 120 units, and finally a 10-unit softmax output layer. Each convolution or fully connected layer is activated by ReLU, and the total number of model parameters is 258898. We train the model for $T = 10000$ rounds, and one mini-batch is run per round with batch size $L_b = 64$. Local computation energy per round per device is set to 10J.

For both MNIST and CIFAR-10, the learning rate η_t is set to 0.05, $\forall t$, a momentum of 0.9 is adopted, and cross entropy is adopted as the loss function. The wireless channel follows Rayleigh fading with scale parameter 1, and by default we assume that the accurate channel gain can be observed, i.e., $\tilde{h}_{n,t} = h_{n,t}$. The variance of channel noise is $\sigma_0^2 = 10^{-6}$. The power scalar is selected according to (27), where the default SNR threshold is $\gamma_0 = 5$. For the dynamic

¹<http://yann.lecun.com/exdb/mnist/>

²<https://www.cs.toronto.edu/~kriz/cifar.html>

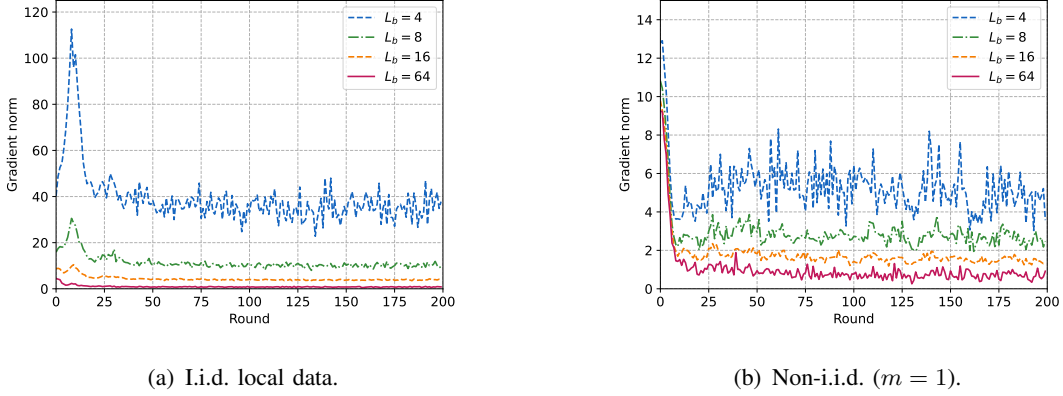


Fig. 2. The l_2 -norm of local gradients and their estimated values on the MNIST dataset.

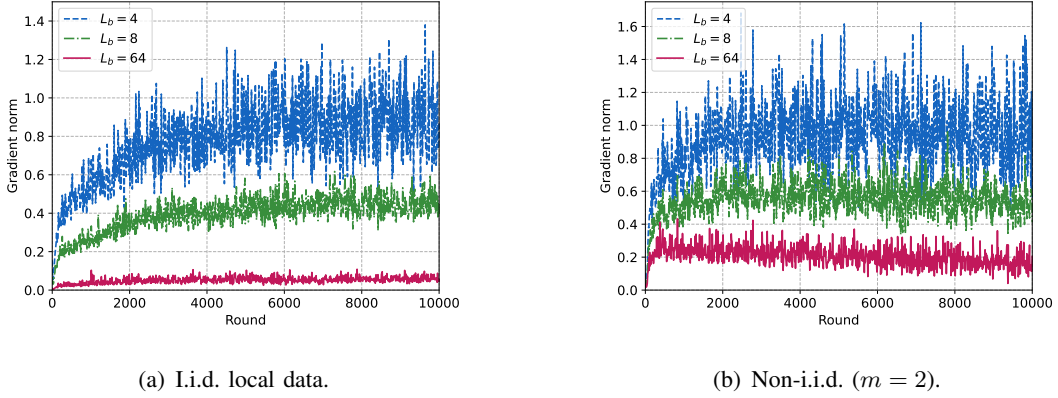


Fig. 3. The l_2 -norm of local gradients and their estimated values on the CIFAR-10 dataset.

scheduling algorithm, the minimum value of virtual queue is $q_{\min} = 0.1$, and the maximum estimation error $\delta_h = 0.5\tilde{E}_{n,t}$ is allowed for communication reschedule.

A. l_2 -Norm of Local Gradients

In Fig. 2 and Fig. 3, we first evaluate the EST-C and EST-P methods proposed in Section IV-A that estimate the l_2 -norm of local gradient, by observing the temporal variations of the gradients. To eliminate the impact of device scheduling, we do not limit the energy consumption and all devices are scheduled. The batch size L_b used for the model training is 64. In each round, each device further computes its local gradient with smaller batch sizes $L_b = 4, 8$ and 16 and records the corresponding estimated gradient norm, which is adopted by the EST-C method as $\|\tilde{\mathbf{g}}_{n,t}^{[\text{est}]} \|_2^2$. For the EST-P method, $\|\tilde{\mathbf{g}}_{n,t}^{[\text{est}]} \|_2^2$ will be given the value of the l_2 -norm of gradients with $L_b = 64$ at a certain round before t . Each curve is averaged over 50 and 20 runs for MNIST

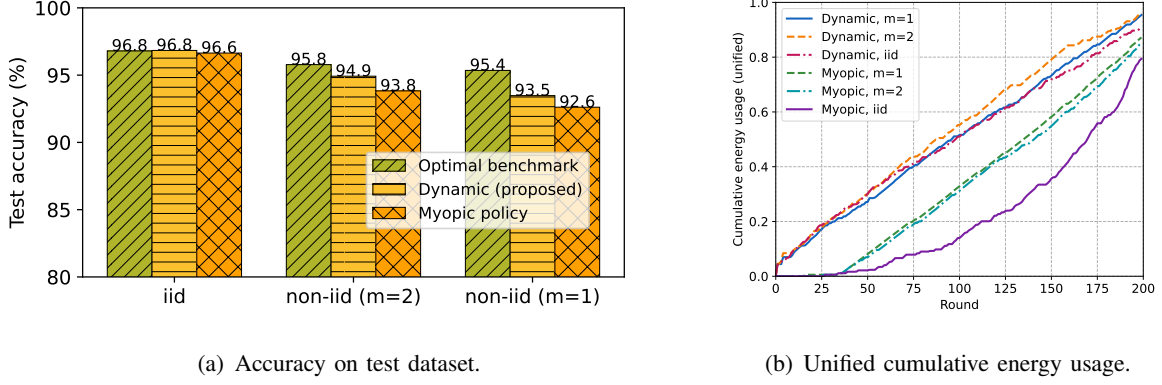


Fig. 4. Performance of the proposed dynamic scheduling algorithm and benchmarks on MNIST.

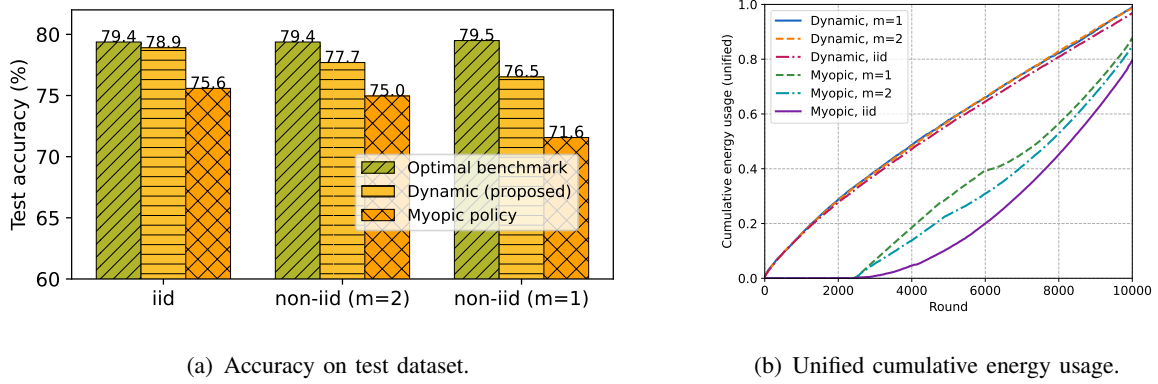


Fig. 5. Performance of the proposed dynamic scheduling algorithm and benchmarks on CIFAR-10.

and CIFAR-10, respectively.

As shown in Fig. 2 and Fig. 3, the gradient norms achieved by different batch sizes are highly varying, and a smaller batch size yields a higher l_2 -norm of gradient due to the non-negligible gradient variance, which is consistent with the analysis in (21). This result indicates that the EST-C method cannot provide an accurate estimation of gradient norm. Meanwhile, with a fixed batch size, such as $L_b = 64$, the gradient norm has a strong temporal correlation. Therefore, the EST-P method can provide a much better estimate of the gradient norm, which is embedded in the proposed dynamic device scheduling algorithm.

B. Performance of the Proposed Device Scheduling Algorithm

We compare the performance of the proposed scheduling algorithm with two benchmarks:

1) *Optimal benchmark*: Devices do not have energy limitations, so that all of them participate in each training round.

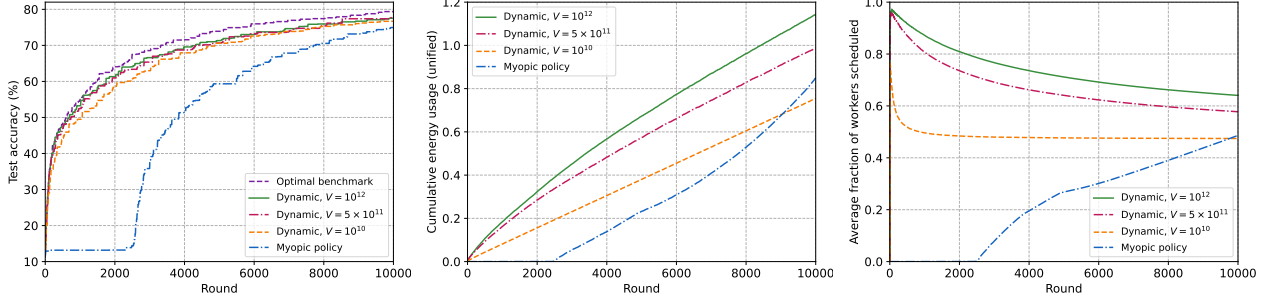
2) *Myopic policy*: For each device n , the maximum energy that can be used in round t is given by the remaining energy divided by the remaining number of rounds, i.e., $\frac{\bar{E}_n - \sum_{\tau=1}^{t-1} \beta_{n,\tau} E_{n,\tau}}{T-t+1}$.

In Fig. 4, we compare the training performance and energy consumption of the proposed dynamic scheduling algorithm with the optimal and myopic benchmarks on MNIST. Let $\bar{E} = 1\text{J}$ be the energy budget per round, and the total energy budget of each device is $\bar{E}_n = T\bar{E}, \forall n$. For non-i.i.d. data with 1 label per device, the weight parameter is $V = 5 \times 10^7$, while for the other two cases, $V = 10^8$. The training performance is characterized by the accuracy of the MLP model on the test dataset, as shown in Fig. 4(a). Results show that our proposed dynamic scheduling algorithm achieves the optimal accuracy under i.i.d. data, and always outperforms the myopic policy. The maximum value of the unified cumulative energy usage across devices till the t -th round, given by $\max_{n \in \mathcal{N}} \frac{\sum_{\tau=1}^t \beta_{n,\tau} E_{n,\tau}}{t\bar{E}}$, is plotted in Fig. 4(b). For the myopic policy, the energy required for computation and communication exceeds the budget at the beginning of training, thus no devices can be scheduled. However, our proposed algorithm enables devices to use energy in a more flexible way, thus improving the training performance.

Similar comparison is made on CIFAR-10 dataset in Fig. 5, where $\bar{E} = 8\text{J}$ and $V = 5 \times 10^{11}$. Note that compared to the local computation energy required per round (10J), the energy budget is relatively limited, and the advantage of the proposed dynamic scheduling over the myopic policy is more prominent in such a scenario. In particular, under the highly non-i.i.d. case with $m = 1$, dynamic scheduling improves the accuracy by 4.9% compared to the myopic policy, by utilizing 10% more energy in a more balanced manner. We can also see that our proposed algorithm can satisfy the energy constraints of devices under both datasets (at the end of training, the unified energy usage is smaller than 1).

In the following, we further explore the impact of key parameters on the training performance and energy consumption with CIFAR-10, as it is more challenging than MNIST. We focus on the non-i.i.d. case, where each device has a local subset with $m = 2$ labels.

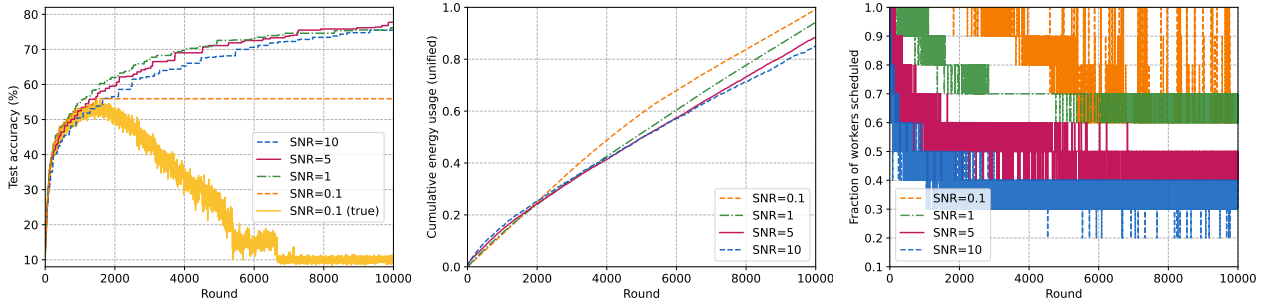
Fig. 6 validates that the weight parameter V can balance the trade-off between the training performance and energy consumption, where $\bar{E} = 8\text{J}$. As V increases, devices use energy in a more aggressive manner, leading to a higher energy usage and more scheduled devices, so as to accelerate the convergence. However, if V is too large, such as $V = 10^{12}$, energy is not given enough attention and finally the limit is violated. In practical systems, V should be judiciously selected to optimize the training performance while satisfying the energy constraints.



(a) Accuracy on test dataset.

(b) Unified cumulative energy usage.

(c) Scheduled devices.

Fig. 6. Performance of the proposed algorithm under different weight parameter V on CIFAR-10.

(a) Accuracy on test dataset.

(b) Unified cumulative energy usage.

(c) Scheduled devices.

Fig. 7. Performance of the proposed algorithm under different received SNR thresholds on CIFAR-10.

The impact of the received SNR threshold γ_0 on the training performance and energy consumption with the proposed dynamic scheduling algorithm is shown in Fig. 7, where $\bar{E} = 8J$, and $V = 2.5 \times 10^{11}$. In Fig. 7(a), the curve marked with ‘true’ plots the model accuracy on the test dataset in the current round, with SNR threshold 0.1, while the other curves present the best test accuracy up-to-date. The maximum cumulative energy usage and instantaneous fraction of devices that are scheduled in each round is shown in Fig. 7(b) and Fig. 7(c), respectively. Clearly, a smaller SNR threshold helps to save communication energy, and thus more devices can be scheduled in each round. However, the cumulative noise might degrade the accuracy or even diverge the training if the SNR is too low, for instance when $SNR = 0.1$. On the other hand, a larger SNR, such as $SNR = 10$, makes communication more energy-consuming, which also degrades the training performance due to fewer participants. A proper value of the received SNR threshold should be given to balance the negative impact of noise and the energy consumption. As shown in Fig. 7, $SNR = 5$ is the best choice under our simulation setting.

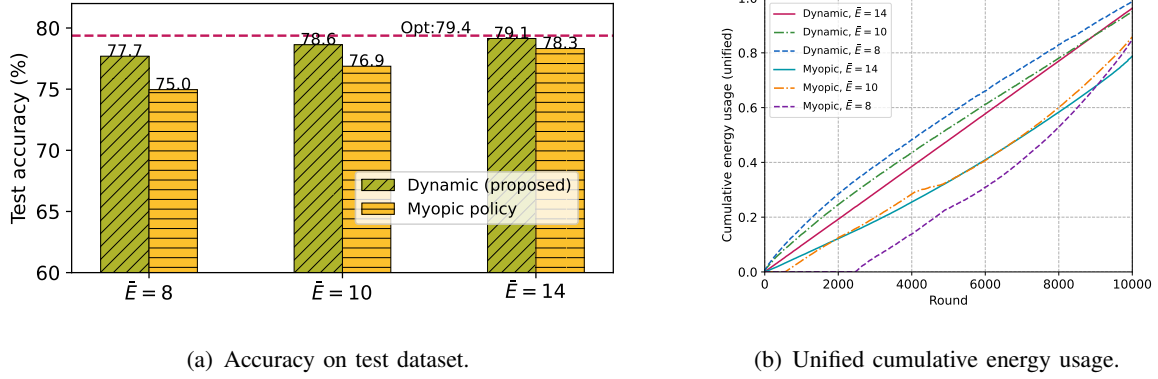


Fig. 8. Performance of the proposed algorithm and benchmarks under different energy budgets on CIFAR-10.

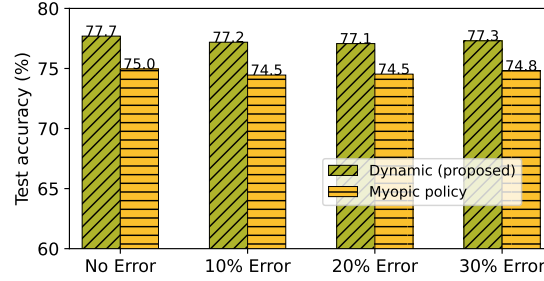


Fig. 9. Performance of the proposed algorithm and myopic benchmark under different channel estimation errors on CIFAR-10.

We compare our proposed algorithm with optimal and myopic benchmarks under different energy budgets in Fig. 8. For $\bar{E} = 14\text{J}$, we set $V = 10^{10}$, and for $\bar{E} = 8\text{J}$ or 10J , we let $V = 5 \times 10^{11}$. Our proposed dynamic scheduling algorithm always outperforms the myopic benchmark by achieving higher accuracy and utilizing energy more efficiently, and approaches the optimal accuracy as $\bar{E}$ increases. Moreover, the accuracy gap between the proposed algorithm and myopic policy is 2.7%, 1.7% and 0.8% for $\bar{E} = 8, 10$ and 14 , respectively, indicating that the dynamic scheduling algorithm is particularly promising under the energy-limited regime.

Finally, we evaluate the robustness of the proposed dynamic scheduling algorithm by introducing channel observation errors, where $\bar{E} = 8\text{J}$ and $V = 5 \times 10^{11}$. In Fig. 9, the first group of bars are obtained without channel observation error, i.e., $\tilde{h}_{n,t} = h_{n,t}$. The second to the forth group of bars suffer inaccurate channel observations. For example, if the error is 20%, then $\tilde{h}_{n,t}$ is uniformly distributed within $[0.8h_{n,t}, 1.2h_{n,t}]$. A larger observation error further leads to less accurate energy estimations $\tilde{E}_{n,t}$. However, the training performance of the proposed

dynamic scheduling algorithm only suffers a tiny degradation, which validates its robustness in practical scenarios. We also mention that the myopic policy also performs well under different observation errors compared to the error-free case. Nevertheless, the proposed algorithm still beats the myopic policy by a significant margin in all the scenarios.

VI. CONCLUSIONS

We have investigated the device scheduling problem for FEEL with over-the-air gradient aggregation, aiming to optimize the training performance under joint communication and computation energy limits of devices. Convergence analysis has been carried out showing the importance of device participation to the training performance, and an energy-aware dynamic device scheduling algorithm has been developed. In particular, we have noticed the existence of unobservable states, mainly the l_2 -norm of local gradients, for online decision making in over-the-air FEEL, and proposed an estimated-drift-plus-penalty solution based on the Lyapunov optimization framework accordingly. We have characterized a theoretical guarantee for the proposed dynamic scheduling algorithm by taking the deviation of estimated states into consideration. Experiments on MNIST and CIFAR-10 datasets have been carried out to validate the theoretical findings. Compared to the myopic benchmark, we have shown a significant 4.9% accuracy improvement on CIFAR-10 for a highly non-i.i.d. data distribution and stringent energy constraints.

As future directions, heterogeneous data distributions across devices can be considered, where local datasets represent different number of classes. We would like to observe if data diversity of a device should be taken into account for the scheduling decision. The trade-off between training delay and energy consumption in over-the-air FEEL is also worth further investigation.

APPENDIX A

PROOF OF LEMMA 1

For the simplicity of notation, let $\tilde{\mathbf{g}}_t \triangleq \frac{\sum_{n \in \mathcal{B}_t} \tilde{\mathbf{g}}_{n,t}}{|\mathcal{B}_t|}$, $\tilde{\mathbf{z}}_t \triangleq \frac{\mathbf{z}_t}{\sigma_t |\mathcal{B}_t|}$. Thus the global model is updated according to

$$\mathbf{w}_t = \mathbf{w}_{t-1} - \eta_t (\tilde{\mathbf{g}}_t + \tilde{\mathbf{z}}_t). \quad (32)$$

According to Assumption 2, the gap of loss between two adjacent rounds can be bounded by

$$F(\mathbf{w}_t) - F(\mathbf{w}_{t-1}) \leq \nabla F(\mathbf{w}_{t-1})^\top (\mathbf{w}_t - \mathbf{w}_{t-1}) + \frac{l}{2} \|\mathbf{w}_t - \mathbf{w}_{t-1}\|_2^2 \quad (33)$$

$$= -\eta_t \mathbf{g}_t^\top (\tilde{\mathbf{g}}_t + \tilde{\mathbf{z}}_t) + \frac{l\eta_t^2}{2} \|\tilde{\mathbf{g}}_t + \tilde{\mathbf{z}}_t\|_2^2. \quad (34)$$

Recall that each entry in the noise vector $\mathbf{z}_t$ follows Gaussian distribution with zero mean and variance σ_0^2 . Taking the expectation over noise, and considering the fact that the channel noise and local gradient are independent, we can obtain

$$\mathbb{E}_{\mathbf{z}_t} [\|\tilde{\mathbf{g}}_t + \tilde{\mathbf{z}}_t\|_2^2] = \|\tilde{\mathbf{g}}_t\|_2^2 + \mathbb{E}_{\mathbf{z}_t} [\|\tilde{\mathbf{z}}_t\|_2^2] = \|\tilde{\mathbf{g}}_t\|_2^2 + \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2}, \quad (35)$$

$$\mathbb{E}_{\mathbf{z}_t} [F(\mathbf{w}_t) - F(\mathbf{w}_{t-1})] \leq -\eta_t \mathbf{g}_t^\top \tilde{\mathbf{g}}_t + \frac{l\eta_t^2}{2} \|\tilde{\mathbf{g}}_t\|_2^2 + \frac{l\eta_t^2}{2} \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2}. \quad (36)$$

Taking the expectation over stochastic data sampling, and based on Assumption 1, we get

$$\mathbb{E}_{\mathcal{L}_{n,t}} [\tilde{\mathbf{g}}_t] = \mathbb{E}_{\mathcal{L}_{n,t}} \left[\frac{\sum_{n \in \mathcal{B}_t} \tilde{\mathbf{g}}_{n,t}}{|\mathcal{B}_t|} \right] = \mathbf{g}_t, \quad (37)$$

$$\mathbb{E}_{\mathcal{L}_{n,t}} [\|\tilde{\mathbf{g}}_t\|_2^2] = \mathbb{E} \left[\left\| \frac{\sum_{n \in \mathcal{B}_t} \sum_{\mathbf{x} \in \mathcal{L}_{n,t}} \nabla f(\mathbf{w}_{t-1}, \mathbf{x})}{L_b |\mathcal{B}_t|} \right\|_2^2 \right] \leq \|\mathbf{g}_t\|_2^2 + \frac{G^2}{L_b |\mathcal{B}_t|}. \quad (38)$$

Finally, taking the expectation over noise and SGD on the left hand side of (36), and substituting its right hand side with (37) and (38), we have

$$\begin{aligned} \mathbb{E}[F(\mathbf{w}_t) - F(\mathbf{w}_{t-1})] &\leq -\eta_t \|\mathbf{g}_t\|_2^2 + \frac{l\eta_t^2}{2} \left(\|\mathbf{g}_t\|_2^2 + \frac{G^2}{L_b |\mathcal{B}_t|} \right) + \frac{l\eta_t^2}{2} \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2} \\ &= -\eta_t \left(1 - \frac{l\eta_t}{2} \right) \|\mathbf{g}_t\|_2^2 + \frac{l\eta_t^2}{2} \left(\frac{G^2}{L_b |\mathcal{B}_t|} + \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2} \right). \end{aligned} \quad (39)$$

Since $\mathbb{E}[F(\mathbf{w}_t) - F(\mathbf{w}_{t-1})] = \mathbb{E}[F(\mathbf{w}_t)] - \mathbb{E}[F(\mathbf{w}_{t-1})]$, Lemma 1 is proved.

APPENDIX B

PROOF OF THEOREM 1

By the μ -strong convexity of the loss functions (Assumption 3), the Polyak-Lojasiewicz inequality holds

$$\|\mathbf{g}_t\|_2^2 \geq 2\mu(F(\mathbf{w}_{t-1}) - F^*). \quad (40)$$

Substituting (40) into Lemma 1, and assuming that $\eta_t \leq \frac{1}{l}$ (thus $1 - \frac{l\eta_t}{2} \geq \frac{1}{2}$), we can obtain

$$\begin{aligned} \mathbb{E}[F(\mathbf{w}_t) - F(\mathbf{w}_{t-1})] &\leq -\eta_t \left(1 - \frac{l\eta_t}{2} \right) \|\mathbf{g}_t\|_2^2 + \frac{l\eta_t^2}{2} \left(\frac{G^2}{L_b |\mathcal{B}_t|} + \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2} \right) \\ &\leq -\eta_t \mu (\mathbb{E}[F(\mathbf{w}_{t-1})] - F^*) + \frac{\eta_t}{2} \left(\frac{G^2}{L_b |\mathcal{B}_t|} + \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2} \right). \end{aligned} \quad (41)$$

Let $A_t \triangleq \frac{\eta_t}{2} \left(\frac{G^2}{L_b |\mathcal{B}_t|} + \frac{\sigma_0^2 s}{\sigma_t^2 |\mathcal{B}_t|^2} \right)$, and thus (41) can be re-written as

$$\mathbb{E}[F(\mathbf{w}_t)] - F^* \leq (1 - \mu\eta_t)(\mathbb{E}[F(\mathbf{w}_{t-1})] - F^*) + A_t. \quad (42)$$

With recursion, we can prove Theorem 1:

$$\begin{aligned} \mathbb{E}[F(\mathbf{w}_t)] - F^* &\leq (1 - \mu\eta_t)(\mathbb{E}[F(\mathbf{w}_{t-1})] - F^*) + A_t \\ &\leq (1 - \mu\eta_t)(1 - \mu\eta_{t-1})(\mathbb{E}[F(\mathbf{w}_{t-2})] - F^*) + (1 - \mu\eta_t)A_{t-1} + A_t \\ &\leq \dots \leq (\mathbb{E}[F(\mathbf{w}_0)] - F^*) \prod_{i=1}^t (1 - \mu\eta_i) + \sum_{i=1}^{t-1} A_i \prod_{j=i+1}^t (1 - \mu\eta_j) + A_t. \end{aligned} \quad (43)$$

APPENDIX C

PROOF OF THEOREM 2

Let $y_{n,t} \triangleq \beta_{n,t}E_{n,t} - \frac{\bar{E}_n}{T}$, and $\tilde{y}_{n,t} \triangleq \beta_{n,t}\tilde{E}_{n,t} - \frac{\bar{E}_n}{T}$. Define the error of estimated energy consumption at device n in the t -th round as $\delta_{n,t} \triangleq \beta_{n,t}\tilde{E}_{n,t} - \beta_{n,t}E_{n,t} = \tilde{y}_{n,t} - y_{n,t}$, with maximum absolute value $\delta_0 \triangleq \max_{\{n,t\}} \left\{ \left| \tilde{E}_{n,t} - E_{n,t} \right| \right\}$. According to the evolution of the virtual queue, which is defined in (23), it is easy to prove that $q_{n,t+1}^2 \leq (q_{n,t} + y_{n,t})^2$ and $y_{n,t} \leq q_{n,t+1} - q_{n,t}$.

Define the Lyapunov function as $L(t) \triangleq \sum_{n=1}^N \frac{1}{2} q_{n,t}^2$, and the Lyapunov drift of a single round as $\Delta_1(t) \triangleq L(t+1) - L(t)$, which is given by

$$\begin{aligned} \Delta_1(t) &= L(t+1) - L(t) = \sum_{n=1}^N \left(\frac{1}{2} q_{n,t+1}^2 - \frac{1}{2} q_{n,t}^2 \right) \\ &\leq \sum_{n=1}^N \left(\frac{1}{2} y_{n,t}^2 + q_{n,t} y_{n,t} \right) \leq \theta_0 + \sum_{n=1}^N q_{n,t} y_{n,t}, \end{aligned} \quad (44)$$

where $\theta_0 \triangleq \sum_{n=1}^N \frac{1}{2} \theta_n^2$ and $\theta_n \triangleq \max_t \{|y_{n,t}|\}$. By adding VU_t on both sides of (44), an upper bound on the single-round drift-plus-penalty function is given by

$$\begin{aligned} \Delta_1(t) + VU_t &\leq \theta_0 + \sum_{n=1}^N q_{n,t} y_{n,t} + VU_t \\ &= \theta_0 + \sum_{n=1}^N q_{n,t} \left(\beta_{n,t} E_{n,t} - \frac{\bar{E}_n}{T} \right) + VU_t \end{aligned} \quad (45)$$

$$\begin{aligned} &= \theta_0 + \sum_{n=1}^N q_{n,t} (\tilde{y}_{n,t} - \delta_{n,t}) + VU_t \\ &= \theta_0 + \sum_{n=1}^N q_{n,t} \left(\beta_{n,t} \tilde{E}_{n,t} - \delta_{n,t} - \frac{\bar{E}_n}{T} \right) + VU_t. \end{aligned} \quad (46)$$

The classical drift-plus-penalty algorithm of Lyapunov optimization aims to minimize the upper bound of $\Delta_1(t) + VU_t$, as shown in (45). Since we do not have the exact value of $E_{n,t}$, we instead minimize the estimated-drift-plus-penalty, as shown in (46).

Define the T -round drift as $\Delta_T \triangleq L(T+1) - L(1) = \sum_{n=1}^N \frac{1}{2} q_{n,T+1}^2$. Then the T -round drift-plus-penalty function can be bounded by:

$$\begin{aligned} \Delta_T + V \sum_{t=1}^T U_t &\leq \sum_{t=1}^T \left(\theta_0 + \sum_{n=1}^N q_{n,t} (\tilde{y}_{n,t} - \delta_{n,t}) \right) + V \sum_{t=1}^T U_t \\ &= \theta_0 T + \sum_{t=1}^T \left(\sum_{n=1}^N q_{n,t} \tilde{y}_{n,t} + VU_t - \sum_{n=1}^N q_{n,t} \delta_{n,t} \right) \end{aligned} \quad (47)$$

We use superscript $*$ to represent the optimal offline solution of $\mathcal{P}4$ (σ_t is not an optimization variable), superscript $\dagger$ to represent the classical drift-plus-penalty algorithm, i.e., $\min_{\{\beta_{n,t}\}} VU_t + \sum_{n=1}^N \beta_{n,t} q_{n,t} E_{n,t}$, and $\ddagger$ to represent our proposed estimated-drift-plus-penalty algorithm that solves $\mathcal{P}6$.

The T -round drift-plus-penalty is bounded by:

$$\begin{aligned} \Delta_T^\ddagger + V \sum_{t=1}^T U_t^\ddagger &\leq \theta_0 T + \sum_{t=1}^T \left(\sum_{n=1}^N q_{n,t} \tilde{y}_{n,t}^\ddagger + VU_t^\ddagger - \sum_{n=1}^N q_{n,t} \delta_{n,t}^\ddagger \right) \\ &\stackrel{(a)}{\leq} \theta_0 T + \sum_{t=1}^T \left(\sum_{n=1}^N q_{n,t} \tilde{y}_n^\dagger(t) + VU_t^\dagger - \sum_{n=1}^N q_{n,t} \delta_{n,t}^\ddagger \right) \\ &= \theta_0 T + \sum_{t=1}^T \left(\sum_{n=1}^N q_{n,t} (y_{n,t}^\dagger + \delta_{n,t}^\dagger) + VU_t^\dagger - \sum_{n=1}^N q_{n,t} \delta_{n,t}^\ddagger \right) \\ &= \theta_0 T + \sum_{t=1}^T \left(\sum_{n=1}^N q_{n,t} y_{n,t}^\dagger + VU_t^\dagger + \sum_{n=1}^N q_{n,t} (\delta_{n,t}^\dagger - \delta_{n,t}^\ddagger) \right) \\ &\stackrel{(b)}{\leq} \theta_0 T + \sum_{t=1}^T \left(\sum_{n=1}^N q_{n,t} y_{n,t}^* + VU_t^* + 2\delta_0 \sum_{n=1}^N q_{n,t} \right). \end{aligned} \quad (48)$$

Inequality (a) holds because optimally solving $\mathcal{P}6$ yields a minimum value $\sum_{n=1}^N q_{n,t} \tilde{y}_{n,t}^\ddagger + VU_t^\ddagger$ for each t . Inequality (b) holds since the drift-plus-penalty algorithm achieves the minimum value of $\sum_{n=1}^N q_{n,t} y_{n,t} + VU_t$, and thus plugging in the optimal offline policy on the right-hand-side increases the value.

Now we bound the right-hand-side of (48). Note that $q_{n,t+1} - q_{n,t} \leq \theta_n, \forall t, n$, and thus

$$q_{n,t} = q_{n,t} - q_{n,1} = \sum_{\tau=1}^{t-1} (q_{n,\tau+1} - q_{n,\tau}) \leq (t-1)\theta_n, \quad (49)$$

$$q_{n,t}y_{n,t}^* = (q_{n,t} - q_{n,1})y_{n,t}^* \leq (t-1)\theta_n^2. \quad (50)$$

Substituting (49) and (50) into (48) yields

$$\begin{aligned} \Delta_T^\dagger + V \sum_{t=1}^T U_t^\dagger &\leq \theta_0 T + V \sum_{t=1}^T U_t^* + \sum_{t=1}^T \sum_{n=1}^N (t-1)\theta_n^2 + 2\delta_0 \sum_{t=1}^T \sum_{n=1}^N (t-1)\theta_n \\ &= \theta_0 T + V \sum_{t=1}^T U_t^* + \theta_0 T(T-1) + T(T-1)\delta_0 \sum_{n=1}^N \theta_n \\ &= V \sum_{t=1}^T U_t^* + \theta_0 T^2 + T(T-1)\delta_0 \sum_{n=1}^N \theta_n. \end{aligned} \quad (51)$$

Notice that $\Delta_T^\dagger \geq 0$, (30) in Theorem 2 can be derived from (51) by dividing both sides by V .

As $U_t > 0$, and for $\forall n$, $\frac{1}{2}q_{n,T+1}^2 \leq \Delta_T$, we get

$$\begin{aligned} \sum_{t=1}^T y_{n,t} &= \sum_{t=1}^T \beta_{n,t} E_{n,t} - \bar{E}_n \leq \sum_{t=1}^T q_{n,t+1} - q_{n,t} = q_{n,T+1} \\ &\leq \sqrt{2\Delta_T} = \sqrt{2V \sum_{t=1}^T U_t^* + 2\theta_0 T^2 + 2T(T-1)\delta_0 \sum_{n=1}^N \theta_n}. \end{aligned} \quad (52)$$

Thus eq. (31) in Theorem 2 is proved.

REFERENCES

- [1] Y. Sun, S. Zhou, and D. Gunduz, "Energy-aware analog aggregation for federated learning with redundant data," in *Proc. IEEE Int. Conf. Commun. (ICC)*, Dublin, Ireland, Jun. 2020.
- [2] C. Jiang, H. Zhang, Y. Ren, Z. Han, K. -C. Chen, and L. Hanzo, "Machine learning paradigms for next-generation wireless networks," *IEEE Wireless Commun.*, vol. 24, no. 2, pp. 98-105, Apr. 2017.
- [3] D. Gunduz, P. de Kerret, N. D. Sidiropoulos, D. Gesbert, C. R. Murthy and M. van der Schaar, "Machine learning in the air," *IEEE J. Sel. Areas Commun.*, vol. 37, no. 10, pp. 2184-2199, Oct. 2019.
- [4] J. Park, S. Samarakoon, M. Bennis, and M. Debbah, "Wireless network intelligence at the edge," in *Proceedings of the IEEE*, vol. 107, no. 11, pp. 2204-2239, Nov. 2019.
- [5] J. Konecny, H. B. McMahan, F. X. Yu, P. Richtárik, A. T. Suresh, and D. Bacon, "Federated learning: Strategies for improving communication efficiency," *NIPS Workshop on Private Multi-Party Machine Learning*, Oct. 2016.
- [6] B. McMahan, E. Moore, D. Ramage, et al. "Communication-efficient learning of deep networks from decentralized data," in *Proc. Artificial Intelligence and Statistics (AISTats)*, Apr. 2017.
- [7] K. Bonawitz, et al. "Towards federated learning at scale: System design," in *Proc. Conf. Systems and Machine Learning*, Stanford, CA, USA, Apr. 2019.
- [8] T. Li, A. K. Sahu, A. Talwalkar, and V. Smith, "Federated learning: Challenges, methods, and future directions," in *IEEE Signal Process. Mag.*, vol. 37, no. 3, pp. 50-60, 2020.

- [9] Y. Zhao, M. Li, L. Lai, N. Suda, D. Civin, and V. Chandra, "Federated learning with non-iid data," *arXiv preprint arXiv:1806.00582*, Jun. 2018.
- [10] W. Y. B. Lim, et al., "Federated learning in mobile edge networks: A comprehensive survey," in *IEEE Commun. Surveys Tuts.*, vol. 22, no. 3, pp. 2031-2063, thirdquarter 2020.
- [11] M. Chen, D. Gunduz, K. Huang, W. Saad, M. Bennis, A. V. Feljan, and H. V. Poor, "Distributed learning in wireless networks: Recent progress and future challenges," *arXiv preprint arXiv:2104.02151*, Apr. 2021.
- [12] D. Alistarh, D. Grubic, J. Li, R. Tomioka, and M. Vojnovic, "QSGD: Communication-efficient SGD via gradient quantization and encoding," in *Proc. Advances in Neural Information Processing Systems (NIPS)*, Dec. 2017.
- [13] J. Wangni, J. Wang, J. Liu, and T. Zhang, "Gradient sparsification for communication-efficient distributed optimization," in *Advances in Neural Information Processing Systems (NIPS)*, Dec. 2018.
- [14] Y. Du, S. Yang, and K. Huang, "High-dimensional stochastic gradient quantization for communication-efficient edge learning," *IEEE Trans. Signal Process.*, vol. 68, pp. 2128-2142, Mar. 2020.
- [15] A. Reisizadeh, A. Mokhtari, H. Hassani, A. Jadbabaie, R. Pedarsani, "FedPAQ: A communicationefficient federated learning method with periodic averaging and quantization," in *International Conference on Artificial Intelligence and Statistics (AISTats)*, Jun. 2020.
- [16] H. H. Yang, Z. Liu, T. Q. S. Quek, and H. V. Poor, "Scheduling policies for federated learning in wireless networks," *IEEE Trans. Commun.*, vol. 68, no. 1, pp. 317-333, Jan. 2020.
- [17] M. Mohammadi Amiri, D. Gunduz, S. R. Kulkarni, and H. V. Poor, "Convergence of update aware device scheduling for federated learning at the wireless edge," in *IEEE Trans. Wireless Commun.*, early access, Jan. 2021.
- [18] S. Wang, T. Tuor, T. Salonidis, K. K. Leung, C. Makaya, T. He and K. Chan, "Adaptive federated learning in resource constrained edge computing systems," *IEEE J. Sel. Areas Commun.*, vol. 37, no. 6, pp. 1205-1221, Jun. 2019.
- [19] N. Yoshida, T. Nishio, M. Morikura, K. Yamamoto and R. Yonetani, "Hybrid-FL for wireless networks: Cooperative learning mechanism using non-iid data," in *Proc. IEEE Int. Conf. Commun. (ICC)*, Dublin, Ireland, Jun. 2020.
- [20] Q. Zeng, Y. Du, K. K. Leung, and K. Huang, "Energy-efficient radio resource allocation for federated edge learning," in *Proc. IEEE Int. Conf. Commun. Workshops*, Dublin, Ireland, Jun. 2020.
- [21] J. Xu and H. Wang, "Client selection and bandwidth allocation in wireless federated learning networks: A long-term perspective," in *IEEE Trans. Wireless Commun.*, early access, Oct. 2020.
- [22] Y. Sun, W. Shi, X. Huang, S. Zhou and Z. Niu, "Edge learning with timeliness constraints: Challenges and solutions," in *IEEE Commun. Mag.*, vol. 58, no. 12, pp. 27-33, Dec. 2020.
- [23] M. Chen, H. V. Poor, W. Saad and S. Cui, "Convergence time optimization for federated learning over wireless networks," in *IEEE Trans. Wireless Commun.*, early access, Dec. 2020.
- [24] W. Shi, S. Zhou, Z. Niu, M. Jiang and L. Geng, "Joint device scheduling and resource allocation for latency constrained wireless federated learning," in *IEEE Trans. Wireless Commun.*, vol. 20, no. 1, pp. 453-467, Jan. 2021.
- [25] J. Ren, Y. He, D. Wen, G. Yu, K. Huang and D. Guo, "Scheduling for cellular federated edge learning with importance and channel awareness," in *IEEE Trans. Wireless Commun.*, vol. 19, no. 11, pp. 7690-7703, Nov. 2020.
- [26] M. S. H. Abad, E. Ozfatura, D. Gunduz, and O. Ercetin, "Hierarchical federated learning across heterogeneous cellular networks," *IEEE Int. Conf. Acoustics, Speech and Signal Process. (ICASSP)*, Barcelona, Spain, May 2020.
- [27] N. H. Tran, W. Bao, A. Zomaya, M. N. H. Nguyen and C. S. Hong, "Federated learning over wireless networks: Optimization model design and analysis," in *Proc. IEEE INFOCOM*, Paris, France, May 2019.

- [28] Z. Yang, M. Chen, W. Saad, C. S. Hong and M. Shikh-Bahaei, "Energy efficient federated learning over wireless communication networks," in *IEEE Trans. Wireless Commun.*, early access, Nov. 2020.
- [29] X. Mo and J. Xu, "Energy-efficient federated edge learning with joint communication and computation design," *arXiv preprint arXiv:2003.00199*, Mar. 2020.
- [30] D. Gunduz, D. B. Kurka, M. Jankowski, M. Mohammadi Amiri, E. Ozfatura and S. Sreekumar, "Communicate to learn at the edge," in *IEEE Commun. Mag.*, vol. 58, no. 12, pp. 14-19, Dec. 2020.
- [31] G. Zhu, D. Liu, Y. Du, C. You, J. Zhang and K. Huang, "Toward an intelligent edge: Wireless communication meets machine learning," in *IEEE Commun. Mag.*, vol. 58, no. 1, pp. 19-25, Jan. 2020.
- [32] M. Mohammadi Amiri and D. Gunduz, "Machine learning at the wireless edge: Distributed stochastic gradient descent over-the-air," *IEEE Trans. Signal Process.*, vol. 68, pp. 2155-2169, Apr. 2020.
- [33] M. Mohammadi Amiri and D. Gunduz, "Federated learning over wireless fading channels," in *IEEE Trans. Wireless Commun.*, vol. 19, no. 5, pp. 3546-3557, May 2020.
- [34] G. Zhu, Y. Wang, and K. Huang, "Low-latency broadband analog aggregation for federated edge learning," in *IEEE Trans. Wireless Commun.*, vol. 19, no. 1, pp. 491-506, Jan. 2020.
- [35] G. Zhu, J. Xu and K. Huang, "Over-the-air computing for 6G: Turning air into a computer," *arXiv preprint arXiv:2009.02181* Sept. 2020.
- [36] H. Guo, A. Liu and V. K. N. Lau, "Analog gradient aggregation for federated learning over wireless networks: Customized design and convergence analysis," in *IEEE Internet Things J.*, vol. 8, no. 1, pp. 197-210, Jan. 2021.
- [37] X. Wei and C. Shen, "Federated learning over noisy channels: Convergence analysis and design examples," *arXiv preprint arXiv:2101.02198*, Jan. 2021.
- [38] M. M. Amiri, D. Gunduz, S. R. Kulkarni, and H. V. Poor, "Convergence of federated learning over a noisy downlink," *arXiv preprint arXiv:2008.11141*, Aug. 2020.
- [39] N. Zhang and M. Tao, "Gradient statistics aware power control for over-the-air federated learning in fading channels," *arXiv preprint arXiv:2003.02089*, Mar. 2020.
- [40] T. Sery, N. Shlezinger, K. Cohen, and Y. C. Eldar, "Over-the-air federated learning from heterogeneous data," *arXiv preprint arXiv:2009.12787*, Sept. 2020.
- [41] M. Mohammadi Amiri, T. M. Duman, D. Gunduz, S. R. Kulkarni, H. V. Poor, "Blind federated edge learning," *arXiv preprint arXiv:2010.10030*, Oct. 2020.
- [42] Y. Shao, D. Gunduz, S. C. Liew, "Federated edge learning with misaligned over-the-air computation," *arXiv preprint arXiv:2102.13604*, Feb. 2021.
- [43] G. Zhu, Y. Du, D. Gunduz and K. Huang, "One-bit over-the-air aggregation for communication-efficient federated edge learning: Design and convergence analysis," in *IEEE Trans. Wireless Commun.*, Nov. 2020.
- [44] M. J. Neely, *Stochastic Network Optimization With Application to Communication and Queueing Systems*. San Rafael, CA, USA: Morgan & Claypool, 2010.
- [45] Y. Sun, S. Zhou, and J. Xu, "EMM: Energy-aware mobility management for mobile edge computing in ultra dense networks," in *IEEE J. Sel. Areas Commun.*, vol. 35, no. 11, pp. 2637-2646, Nov. 2017.