
The Non-IID Data Quagmire of Decentralized Machine Learning

Kevin Hsieh^{1 2} Amar Phanishayee¹ Onur Mutlu^{3 2} Phillip B. Gibbons²

Abstract

Many large-scale machine learning (ML) applications need to perform *decentralized* learning over datasets generated at different devices and locations. Such datasets pose a significant challenge to decentralized learning because their different contexts result in significant data distribution skew across devices/locations. In this paper, we take a step toward better understanding this challenge by presenting a detailed experimental study of decentralized DNN training on a common type of data skew: skewed distribution of data labels across locations/devices. Our study shows that: (i) skewed data labels are a fundamental and pervasive problem for decentralized learning, causing significant accuracy loss across many ML applications, DNN models, training datasets, and decentralized learning algorithms; (ii) the problem is particularly challenging for DNN models with batch normalization; and (iii) the degree of skewness is a key determinant of the difficulty of the problem. Based on these findings, we present SkewScout, a system-level approach that adapts the communication frequency of decentralized learning algorithms to the (skew-induced) accuracy loss between data partitions. We also show that group normalization can recover much of the accuracy loss of batch normalization.

1. Introduction

The advancement of machine learning (ML) is heavily dependent on the processing of massive amounts of data. The most timely and relevant data are often generated at different devices all over the world, e.g., data collected by mobile phones and video cameras. Because of communication and privacy constraints, gathering all these data for centralized processing can be impractical/infeasible. For example, moving raw data across national borders is subject to data sovereignty law constraints (e.g., GDPR). Similar constraints apply to centralizing private data from phones.

¹Microsoft Research ²Carnegie Mellon University ³ETH Zürich. Correspondence to: Kevin Hsieh <kevin.hsieh@microsoft.com> and Phillip B. Gibbons <gibbons@cs.cmu.edu>.

These constraints motivate the need for ML training over widely distributed data (*decentralized learning*). For example, *geo-distributed learning* (Hsieh et al., 2017) trains a global ML model over data spread across geo-distributed data centers. Similarly, *federated learning* (McMahan et al., 2017) trains a centralized model over data from a large number of mobile devices. Federated learning has been an important topic both in academia (140+ papers in 2019) and industry (500+ million installations on Android devices).

Key Challenges in Decentralized Learning. There are two key challenges in decentralized learning. First, training a model over decentralized data using traditional training approaches (i.e., those designed for centralized data, often using a bulk synchronous parallel (BSP) approach (Valiant, 1990)) requires massive communication. This drastically slows down the training process because the communication is bottlenecked by the limited wide-area or mobile network bandwidth (Hsieh et al., 2017; McMahan et al., 2017). Second, decentralized data are typically generated at different contexts, which can lead to significant differences in the *distribution* of data across data partitions. For example, facial images collected by cameras will reflect the demographics of each camera’s location, and images of kangaroos will be collected only from cameras in Australia or zoos. Unfortunately, existing decentralized learning algorithms (e.g., (Hsieh et al., 2017; McMahan et al., 2017; Smith et al., 2017; Lin et al., 2018; Tang et al., 2018)) mostly focus on reducing communication, as they either (i) assume the data partitions are independent and identically distributed (IID) or (ii) conduct only a very limited study on non-IID data partitions. This leaves a key question mostly unanswered: *What happens to ML applications and decentralized learning algorithms when their data partitions are not IID?*

Our Goal and Key Findings. We aim to take a step to further the understanding of the above key question. In this work, we focus on a common type of non-IID data, widely used in prior work (e.g., (McMahan et al., 2017; Tang et al., 2018; Zhao et al., 2018)): skewed distribution of data labels across locations/devices. Such *skewed label partitions* arise frequently in the real world (see §2.2 for examples). Our study covers various DNN applications, DNNs, training datasets, decentralized learning algorithms, and degrees of label skewness. Our study reveals three key findings:

- Training over skewed label partitions is a fundamental and pervasive problem for decentralized learning. Three decentralized learning algorithms (Hsieh et al., 2017; McMahan et al., 2017; Lin et al., 2018) suffer from major model quality loss when run to convergence on skewed label partitions, across the applications, models, and training datasets in our study.
- DNNs with *batch normalization* (Ioffe & Szegedy, 2015) are particularly vulnerable to skewed label partitions, suffering significant model quality loss even under BSP, the most communication-heavy approach.
- The degree of skewness is a key determinant of the difficulty level of the problem.

These findings reveal that non-IID data is an important yet heavily understudied challenge in decentralized learning, worthy of extensive study. To facilitate further study on skewed label partitions, we release a real-world, geo-tagged dataset of common mammals on Flickr (Flickr) (§2.2).

Solutions. As two initial steps towards addressing the vast challenge of non-IID data, we first show that among the many proposed alternatives to batch normalization, *group normalization* (Wu & He, 2018) avoids the skew-induced accuracy loss of batch normalization under BSP. With this fix, all models in our study achieve high accuracy on skewed label partitions under (communication-heavy) BSP, and the problem can be viewed as a trade-off between accuracy and the amount of communication. Intuitively, there is a tug-of-war among different data partitions, with each partition pulling the model to reflect its data, and only close communication, tuned to the skew-induced accuracy loss, can save the overall model accuracy of the algorithms in our study. Accordingly, we present SkewScout, which periodically sends local models to remote data partitions and compares the model performance (e.g., validation accuracy) between local and remote partitions. Based on the accuracy loss, SkewScout adjusts the amount of communication among data partitions by controlling how *relaxed* the decentralized learning algorithms should be, such as controlling the threshold that determines which parameters are worthy of communication. Thus, SkewScout can seamlessly integrate with decentralized learning algorithms that provide such communication control. Our experimental results show that SkewScout’s adaptive approach automatically reduces communication by $9.6\times$ (under high skew) to $34.1\times$ (under mild skew) while retaining the accuracy of BSP.

Contributions. We make the following contributions. First, we conduct a detailed empirical study on the problem of skewed label partitions. We show that this problem is a fundamental and pervasive challenge for decentralized learning. Second, we build and release a large real-world dataset to facilitate future study on this challenge. Third, we make a new observation showing that this challenge is particularly

problematic for DNNs with batch normalization, even under BSP. We discuss the root cause of this problem and we find that it can be addressed by using an alternative normalization technique. Fourth, we show that the difficulty level of this problem varies with the data skew. Finally, we design and evaluate SkewScout, a system-level approach that adapts the communication frequency to reflect the skewness in the data, seeking to maximize communication savings while preserving model accuracy.

2. Background and Motivation

We first provide background on popular decentralized learning algorithms (§2.1). We then highlight a real-world example of skewed label partitions: geographical distribution of animal pictures on Flickr, among other examples (§2.2).

2.1. Decentralized Learning

In a decentralized learning setting, we aim to train a ML model w based on all the training data samples (x_i, y_i) that are generated and stored in one of the K partitions (denoted as P_k). The goal of the training is to fit w to all data samples. Typically, most decentralized learning algorithms assume the data samples are independent and identically distributed (IID) among different P_k , and we refer to such a setting as the *IID setting*. Conversely, we call it the *Non-IID setting* if such an assumption does not hold.

We evaluate three popular decentralized learning algorithms to see how they perform on different applications over the IID and Non-IID settings, using skewed label partitions. These algorithms can be used with a variety of stochastic gradient descent (SGD) approaches, and aim to reduce communication, either among data partitions (P_k) or between the data partitions and a centralized server.

- **Gaia** (Hsieh et al., 2017), a geo-distributed learning algorithm that dynamically eliminates insignificant communication among data partitions. Each partition P_k accumulates updates Δw_j to each model weight w_j locally, and communicates Δw_j to all other data partitions only when its relative magnitude exceeds a predefined threshold (Algorithm 1 in Appendix A¹).
- **FederatedAveraging** (McMahan et al., 2017), a popular algorithm for federated learning that combines local SGD on each client with model averaging. Specifically, **FederatedAveraging** selects a subset of the partitions P_k in each epoch, runs a pre-specified number of local SGD steps on each selected P_k , and communicates the resulting models back to a centralized server. The server averages all these models and uses the averaged w as the starting point for the next epoch (Algorithm 2 in Appendix A).

¹All Appendices are in the supplemental material.

- **DeepGradientCompression** (Lin et al., 2018), a popular algorithm that communicates only a pre-specified amount of gradients each training step, with various techniques to retain model quality such as momentum correction, gradient clipping (Pascanu et al., 2013), momentum factor masking, and warm-up training (Goyal et al., 2017) (Algorithm 3 in Appendix A).

In addition to these decentralized learning algorithms, we also show the results of using BSP (Valiant, 1990) over the IID and Non-IID settings. BSP is significantly slower than the above algorithms because it does not seek to reduce communication: all updates from each P_k are accumulated and shared among all data partitions after each training step. As noted earlier, for decentralized learning, there is a natural tension between the amount of communication and the quality of the resulting model: differing distributions among the P_k pull the model in different directions—more communication helps mitigate this “tug-of-war” in order that the model well-represents all the data. Thus, BSP, with its full communication, is used as the target for model quality.

2.2. Real-World Examples of Skewed Label Partitions

Non-IID data among devices/locations encompass many different forms. There can be skewed distribution of features ($\mathcal{P}(x)$), labels ($\mathcal{P}(y)$), or the relationship between features and labels (e.g., varying $\mathcal{P}(y|x)$ or $\mathcal{P}(x|y)$) among devices/locations (Kairouz et al., 2019) (see more discussions in Appendix I). In this work, we focus on skewed distribution of labels ($\mathcal{P}_{P_i}(y) \not\sim \mathcal{P}_{P_j}(y)$ for different data partitions P_i and P_j), which is also the setting considered by most prior work in this domain (e.g., (McMahan et al., 2017; Tang et al., 2018; Zhao et al., 2018)).

Skewed distribution of labels is common whenever data are generated from heterogeneous users or locations. For example, pedestrian and bicycles are more common in street cameras than highway cameras (Luo et al., 2019). In facial recognition tasks, most individuals will appear in only a few locations around the world. Certain types of clothing (mittens, cowboy boots, kimonos, etc.) will be nearly non-existent in many parts of the world. Similarly, certain animals (e.g., kangaroos) are far more likely to show up in certain locations (Australia). In the rest of this section, we highlight this phenomenon with a study of the geographical distribution of animal pictures on Flickr (Flickr).

Dataset Creation. We start with the 48 classes in the *mammal* subcategory of the 600 most common classes for bounding boxes in Open Images V4 (Kuznetsova et al., 2018). For each class label, we use Flickr’s API to search for relevant pictures. Due to noise in Flickr search results (e.g., “jaguar” returns the animal or the car), we clean the data with a state-of-the-art DNN, PNAS (Liu et al., 2018), which is pre-trained on ImageNet. As we can only clean classes

that exist in both Open Images and ImageNet, we end up with 42 mammal classes and 736,005 total pictures. The resulting dataset, which we call the *Flickr-Mammal* dataset, is openly available at <https://doi.org/10.5281/zenodo.3676081> (see Appendix B for more details).

Geographical Analysis. We map each Flickr picture’s geo-tag to its corresponding geographic regions based on the M49 Standard (United Nation Statistics Division, 2019). As we are mostly interested in $\mathcal{P}(y)$ among different regions, we normalize the number of samples across regions (unnormalized results are similar (Appendix B)). Table 1 illustrates the top-5 classes among first-level regions (continents) and their normalized share of samples in the world.

Region	Top 1	Top 2	Top 3	Top 4	Top 5
Africa	zebra (72%)	antelope (71%)	lion (68%)	cheetah (62%)	hippopotamus (59%)
Americas	mule (84%)	skunk (82%)	armadillo (73%)	harbor seal (65%)	squirrel (61%)
Asia	panda (64%)	hamster (59%)	monkey (58%)	camel (51%)	red panda (42%)
Europe	lynx (72%)	hedgehog (56%)	sheep (56%)	deer (43%)	otter (43%)
Oceania	kangaroo (92%)	koala (92%)	whale (44%)	sea lion (34%)	alpaca (32%)

Table 1. Top-5 mammals in each continent and their share of samples worldwide (e.g., 72% of zebra images are from Africa).

Skewed distribution of labels is a natural phenomenon.

As Table 1 shows, the top-5 mammals in each continent take 32%–92% of normalized sample share in the world (compared to 20% if the distribution were IID). As expected, the top mammals in each region reflect their population share in the world (e.g., kangaroos/koalas in Oceania and zebras/antelopes in Africa). Furthermore, there is *no overlap* for the top-5 classes among different continents, which suggests drastically different label distributions ($\mathcal{P}(y)$) among continents. We see similar phenomenon when the analysis is done based on second-level geographical regions (Appendix B). This demonstrates that in a decentralized learning setting, where such images would be collected and stored in their native regions, the distribution of labels across partitions would be highly skewed.

3. Setup

Our study consists of three dimensions: (i) ML applications/models, (ii) decentralized learning algorithms, and (iii) degree of skewness. We explore all three dimensions with rigorous experimental methodologies. In particular, we make sure the accuracy of our trained ML models on IID data matches the reported accuracy in corresponding papers.

Applications. We evaluate different deep learning applications, DNN model structures, and training datasets:

- **IMAGE CLASSIFICATION** with four DNN models: AlexNet (Krizhevsky et al., 2012), GoogLeNet (Szegedy et al., 2015), LeNet (LeCun et al., 1998), and ResNet (He et al., 2016). We use two datasets, CIFAR-10 (Krizhevsky, 2009) and ImageNet (Russakovsky et al., 2015). We use the default validation set in these two datasets to report the validation accuracy as our model quality metric. We use popular datasets in order to compare model accuracy with existing work, and we also report results with our *Flickr-Mammal* dataset.
- **FACE RECOGNITION** with the center-loss face model (Wen et al., 2016) over the CASIA-WebFace (Yi et al., 2014) dataset. We use verification accuracy on the LFW dataset (Huang et al., 2007) as the model quality metric.

For all applications, we tune the training parameters (e.g., learning rate, minibatch size, number of epochs, etc.) such that the baseline model (BSP in the IID setting) achieves the model quality of the original paper. We then use these training parameters in all other settings. We further ensure that training/validation accuracy has stopped improving by the end of all our experiments. Appendix C lists all major training parameters in our study.

Non-IID Data Partitions. In addition to studying *Flickr-Mammal*, we create non-IID data partitions by partitioning datasets using the data labels, i.e., using image classes for image classification and person identities for face recognition. We control the *skewness* by controlling the fraction of data that are non-IID. For example, 20% non-IID indicates 20% of the dataset are partitioned by labels, while the remaining 80% are partitioned randomly. In §4–§5 we focus on the 100% non-IID setting in which data are completely partitioned by label, while §6 studies the effect of varying the degree of label skew. As our goal is to train a global model, the model is tested on the entire validation set.

Hyper-Parameters Selection. The algorithms we study provide the following hyper-parameters (see Appendix A for the detail of these algorithms) to control the amount of communication (and hence the training time):

- Gaia uses T_0 , the initial threshold to determine if a Δw_j is significant. The significance threshold decreases whenever the learning rate decreases.
- FederatedAveraging uses $Iter_{Local}$ to control the number of local SGD steps on each selected P_k .
- DeepGradientCompression uses s to control the sparsity of updates (update magnitudes in top s percentile are exchanged). Following the original paper (Lin et al., 2018), s follows a warm-up scheduling: 75%, 93.75%, 98.4375%, 99.6%, 99.9%. We use a hyper-parameter E_{warm} , the number of epochs for each warm-up sparsity, to control the duration of the warm-up. For example, if $E_{warm} = 4$, s is 75% in epochs 1–4, 93.75% in epochs 5–8, and so on.

We select a hyper-parameter θ of each decentralized learning algorithms (T_0 , $Iter_{Local}$, E_{warm}) so that (i) θ achieves the same model quality as BSP in the IID setting and (ii) θ achieves similar communication savings across the three decentralized learning algorithms. We study the sensitivity of our findings to the choice of θ in §4.4.

4. Non-IID Study: Results Overview

This paper seeks to answer the question as to what happens to ML applications, ML models, and decentralized learning algorithms when their data label partitions are not IID. In this section, we provide an overview of our findings, showing that skewed label partitions cause *major model quality loss*, across many applications, models, and algorithms.

4.1. Image Classification

We first present the model quality with different decentralized learning algorithms over the IID and Non-IID settings for IMAGE CLASSIFICATION over the CIFAR-10 dataset. We use five partitions ($K = 5$) in this evaluation. As the CIFAR-10 dataset consists of ten object classes, each data partition has two object classes in the Non-IID setting. Figure 1 shows the results with four popular DNNs (AlexNet, GoogLeNet, LeNet, and ResNet). According to the hyper-parameter criteria in §3, we select $T_0 = 10\%$ for Gaia, $Iter_{Local} = 20$ for FederatedAveraging, and $E_{warm} = 8$ for DeepGradientCompression. We make two major observations.

1) Non-IID data is a pervasive problem. All three decentralized learning algorithms lose significant model quality for *all* four DNNs in the Non-IID setting. We see that while these algorithms retain the validation accuracy of BSP in the *IID setting* with $15\times$ – $20\times$ communication savings (agreeing with the results from the original papers for these algorithms), they lose 3% to 74% validation accuracy in the *Non-IID setting*. Simply running these algorithms for more epochs would not help because the training/validation accuracy has already stopped improving. Furthermore, the training is completely diverged in some cases, such as DeepGradientCompression with GoogLeNet and ResNet20 (DeepGradientCompression with ResNet20 also diverges in the IID setting). The pervasiveness of the problem is quite surprising, as we have a diverse set of decentralized learning algorithms and DNNs. This result shows that Non-IID data is a pervasive and challenging problem for decentralized learning, and this problem has been heavily understudied. §4.3 discusses the cause of this problem.

2) Even BSP cannot completely solve this problem. We see that even BSP, with its full communication every step, cannot retain model quality for some DNNs in the Non-IID setting. In particular, the validation accuracy of ResNet20 in

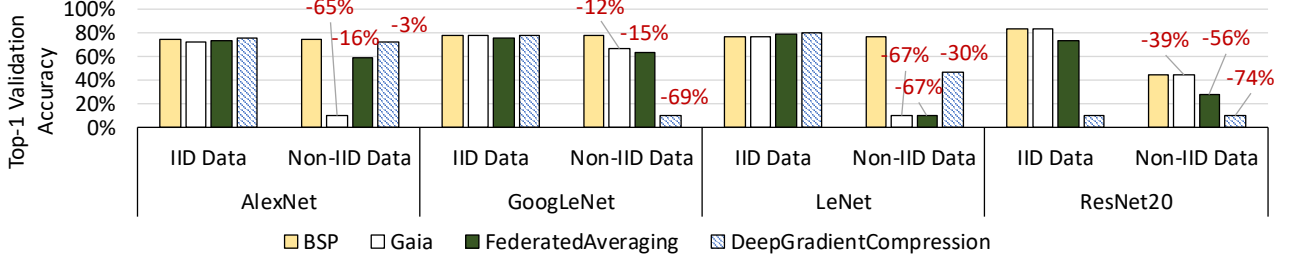


Figure 1. Top-1 validation accuracy for IMAGE CLASSIFICATION over the CIFAR-10 dataset. Each “-x%” label indicates the accuracy loss from BSP in the IID setting.

the Non-IID setting is 39% lower than that in the IID setting. This finding suggests that, for some DNNs, it *may not be possible* to solve the Non-IID data challenge by increasing communication between data partitions. We find that this problem not only exists in ResNet20, but in all DNNs we study with batch normalization (ResNet10, BN-LeNet (Ioffe & Szegedy, 2015) and Inception-v3 (Szegedy et al., 2016)). We discuss this problem and potential solutions in §5.

The same trend in a larger dataset. We conduct a similar study over the ImageNet dataset (Russakovsky et al., 2015) (1,000 image classes). We observe the same problems in the ImageNet dataset (e.g., an 8.1% to 61.7% accuracy loss on ResNet10), whose number of classes is two orders of magnitude more than the CIFAR-10 dataset. Appendix D discusses the experiment in detail.

The same problem in real-world datasets. We run similar experiments on our Flickr-Mammal dataset. We use five partitions ($K = 5$) in this experiment, one for each continent, where each partition has as its local training data precisely the images from its corresponding continent. Thus, we capture the real-world non-IID setting present in Flickr-Mammal. For comparison, we also consider an artificial IID setting, in which all the Flickr-Mammal images are randomly distributed among the five partitions. Figure 2 shows the results. We observe the same problems for decentralized learning algorithms on this real-world dataset. Specifically, Gaia and FederatedAveraging are able to retain the model quality in the (artificial) IID setting, but they lose 3.7% and 3.2% accuracy in the (real-world) Non-IID setting, respectively. (The loss is smaller than in Figure 1 in part because the dataset is not 100% non-IID.) This loss arises even with modest hyper-parameter settings, and is expected to be larger with less conservative settings that more greatly reduce communication. This is significant as the result suggests that skewed data labels arising in real-world settings are a major problem for decentralized learning.

4.2. Face Recognition

We further examine another popular ML application, FACE RECOGNITION, to see if the Non-IID data problem is a challenge across different applications. We use two parti-

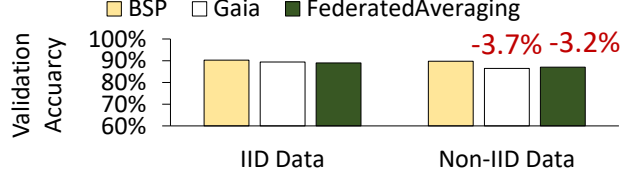


Figure 2. Top-1 validation accuracy for IMAGE CLASSIFICATION over the Flickr-Mammal dataset, where 5% data are randomly selected as the validation set. Non-IID Data is based on real-world data distribution among continents, and IID Data is the artificial setting in which training images are randomly assigned to partitions. We use GoogLeNet (Szegedy et al., 2015) in this experiment, and we select $T_0 = 10\%$ for Gaia and $Iter_{Local} = 20$ for FederatedAveraging based on the criteria in §3. Each “-x%” label indicates the accuracy loss from BSP in the IID setting. Note: y-axis starts at 60% accuracy.

tions in this evaluation. According to the hyper-parameter criteria in §3, we select $T_0=20\%$ for Gaia, $Iter_{Local}=50$ for FederatedAveraging, and $E_{warm}=8$ for Deep-GradientCompression. It is worth noting that the verification process of FACE RECOGNITION is fundamentally different from IMAGE CLASSIFICATION, as FACE RECOGNITION does *not* use the classification layer (and thus the training labels) at all in the verification process. Instead, for each pair of verification images, the DNN uses the distance between the feature vectors of these images to determine if the two images belong to the same person.

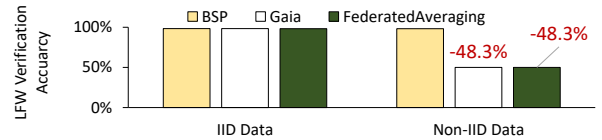


Figure 3. LFW verification accuracy for FACE RECOGNITION. Each label indicates the accuracy loss from BSP in the IID setting.

The same problem in different applications. Figure 3 presents the LFW verification accuracy using different decentralized learning algorithms. Again, the same problem happens in this application: the decentralized learning algorithms work well in the IID setting, but they lose signif-

icant accuracy in the Non-IID setting. In fact, both `Gaia` and `FederatedAveraging` cannot converge to a useful model in the Non-IID setting: their 50% accuracy is no better than random guessing for the binary questions. This result is interesting as FACE RECOGNITION uses a verification process that does not rely on the training labels, which are used to create the Non-IID setting to begin with.

4.3. The problem of decentralized algorithms

The above results show that three diverse decentralized learning algorithms all suffer drastic accuracy loss in the Non-IID setting. We find two reasons for the accuracy loss. First, for algorithms such as `Gaia` that save communication by allowing small model differences in each P_k , the Non-IID setting results in *completely different models among P_k* . The small differences give local models room for specializing for local data. Second, for algorithms that save communication by synchronizing sparsely (e.g., `FederatedAveraging` and `DeepGradient-Compression`), each P_k generates more diverged gradients in the Non-IID setting, which is not surprising as each P_k sees vastly different training data. When they are finally synchronized, they may have diverged so much from the global model that they lead the global model to the wrong direction. See Appendix E.

4.4. Algorithm Hyper-Parameters

We also study the sensitivity of the Non-IID problem to hyper-parameter choice among decentralized learning algorithms. We find that even relatively conservative hyper-parameter settings, which incur high communication costs, still suffer major accuracy loss in the Non-IID setting. In the IID setting, on the other hand, the *same* hyper-parameter achieves similar high accuracy as BSP. In other words, the the Non-IID problem is not specific to particular hyper-parameter choices. Appendix F shows all the results.

5. Batch Normalization: Problem and Solution

5.1. Batch Normalization in the Non-IID Setting

How BatchNorm works. Batch normalization (BatchNorm) (Ioffe & Szegedy, 2015) is one of the most popular mechanisms in deep learning (19,000+ citations as of June 2020). BatchNorm aims to stabilize a DNN by normalizing the input distribution to zero mean and unit variance. Because the *global* mean and variance is unattainable with stochastic training, BatchNorm uses *minibatch mean and variance* as an estimate of the global mean and variance. Specifically, for each minibatch $\mathcal{B}$, BatchNorm calculates the minibatch mean $\mu_{\mathcal{B}}$ and variance $\sigma_{\mathcal{B}}$, and then uses $\mu_{\mathcal{B}}$ and $\sigma_{\mathcal{B}}$ to normalize each input in $\mathcal{B}$. BatchNorm enables

faster and more stable DNN training because it enables larger learning rates (Bjorck et al., 2018; Santurkar et al., 2018).

BatchNorm and the Non-IID setting. While BatchNorm is effective in practice, its dependence on minibatch mean and variance ($\mu_{\mathcal{B}}$ and $\sigma_{\mathcal{B}}$) is known to be problematic in certain settings. This is because BatchNorm uses $\mu_{\mathcal{B}}$ and $\sigma_{\mathcal{B}}$ for training, but it typically uses an estimated global mean and variance (μ and σ) for validation. If there is a major mismatch between these means and variances, the validation accuracy is going to be low. This can happen if the minibatch size is small or the sampling of minibatches is not IID (Ioffe, 2017). The Non-IID setting in our study exacerbates this problem because each data partition P_k sees very different training samples. Hence, the $\mu_{\mathcal{B}}$ and $\sigma_{\mathcal{B}}$ in each P_k can vary significantly in the Non-IID setting, and the synchronized global model may not work for *any* set of data. Worse still, we cannot simply increase the minibatch size or do better minibatch sampling to solve this problem, because in the Non-IID setting the underlying dataset in each P_k does not represent the global dataset.

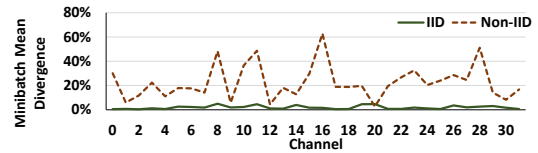


Figure 4. Minibatch mean divergence for the first layer of BN-LeNet over CIFAR-10 using two P_k .

We validate if there is indeed major divergence in $\mu_{\mathcal{B}}$ and $\sigma_{\mathcal{B}}$ among different P_k in the Non-IID setting. We calculate the divergence of $\mu_{\mathcal{B}}$ as the difference between $\mu_{\mathcal{B}}$ in different P_k over the average $\mu_{\mathcal{B}}$ (i.e., it is $\frac{||\mu_{\mathcal{B}, P_0} - \mu_{\mathcal{B}, P_1}||}{||AVG(\mu_{\mathcal{B}, P_0}, \mu_{\mathcal{B}, P_1})||}$ for two partitions P_0 and P_1). We use the average $\mu_{\mathcal{B}}$ over every 100 minibatches in each P_k so that we get better estimation. Figure 4 depicts the divergence of $\mu_{\mathcal{B}}$ for each channel of the first layer of BN-LeNet, which is constructed by inserting BatchNorm to LeNet after each convolutional layer. As we see, the divergence of $\mu_{\mathcal{B}}$ is significantly larger in the Non-IID setting (between 6% to 51%) than in the IID setting (between 1% to 5%). We also observe the same trend in minibatch variances $\sigma_{\mathcal{B}}$ (not shown). As this problem has nothing to do with the amount of communication among P_k , it explains why even BSP cannot retain model accuracy for BatchNorm in the Non-IID setting.

5.2. Alternatives to Batch Normalization

As the problem of BatchNorm in the Non-IID setting is due to its dependence on minibatches, the natural solution is to replace BatchNorm with alternative normalization mechanisms that are *not* dependent on minibatches. Unfortunately, most existing alternative normalization mechanisms (Weight

Normalization (Salimans & Kingma, 2016), Layer Normalization (Ba et al., 2016), Batch Renormalization (Ioffe, 2017)) have their own drawbacks (see Appendix G). Here, we discuss a particular mechanism that may be used instead.

Group Normalization. Group Normalization (GroupNorm) (Wu & He, 2018) is an alternative normalization mechanism that aims to overcome the shortcomings of BatchNorm and Layer Normalization (LayerNorm). GroupNorm divides adjacent channels into groups of a prespecified size $\mathcal{G}_{size}$, and computes the per-group mean and variance for *each input sample*. Hence, GroupNorm does not depend on minibatches for normalization (the shortcoming of BatchNorm), and GroupNorm does not assume all channels make equal contributions (the shortcoming of LayerNorm).

We evaluate GroupNorm with BN-LeNet over CIFAR-10. We carefully select $\mathcal{G}_{size} = 2$, which works best with this DNN. Figure 5 shows the Top-1 validation accuracy with GroupNorm and BatchNorm across decentralized learning algorithms. We make two major observations.

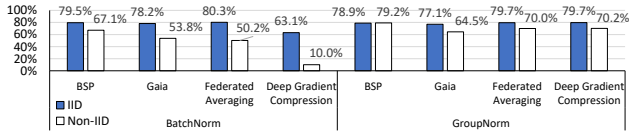


Figure 5. Top-1 validation accuracy with BatchNorm and GroupNorm for BN-LeNet over CIFAR-10 with 5 partitions.

First, GroupNorm successfully recovers the accuracy loss of BatchNorm with BSP in the Non-IID setting. As the figure shows, GroupNorm with BSP achieves 79.2% validation accuracy in the Non-IID setting, which is as good as the accuracy in the IID setting. This shows GroupNorm can be used as an alternative to BatchNorm to overcome the Non-IID data challenge for BSP. Second, GroupNorm dramatically helps the decentralized learning algorithms in the Non-IID setting as well. With GroupNorm, there is 14.4%, 8.9% and 8.7% accuracy loss for Gaia, FederatedAveraging and DeepGradientCompression, respectively. While the accuracy losses are still significant, they are better than their BatchNorm counterparts by an additive 10.7%, 19.8% and 60.2%, respectively.

Discussion. While our study shows that GroupNorm can be a good alternative to BatchNorm in the Non-IID setting, it is worth noting that BatchNorm is widely adopted in many DNNs, hence, more study should be done to see if GroupNorm can always replace BatchNorm for different applications and DNN models. As for other tasks such as recurrent (e.g., LSTM (Hochreiter & Schmidhuber, 1997)) and generative (e.g., GAN (Goodfellow et al., 2014)) models, other normalization techniques such as LayerNorm can be good options because (i) they are shown to be effective in these tasks and (ii) they are not dependent on minibatches.

6. Degree of Skewness

In §4–§5, we studied a strict case of skewed label partitions, where each label only exists in a data partition *exclusively*. While this case may be a reasonable approximation for some applications (e.g., for FACE RECOGNITION, a person’s face image may exist only in one data partition), it could be an extreme case for other applications (e.g., IMAGE CLASSIFICATION, as §2.2 shows). Here, we study how the problem changes with the degree of skewness by controlling the fraction of data that are non-IID (§3). Figure 6 illustrates the CIFAR-10 Top-1 validation accuracy of GN-LeNet (our name for BN-LeNet with GroupNorm replacing BatchNorm, Figure 5) in the 20%, 40%, 60% and 80% non-IID setting. We make two key observations.

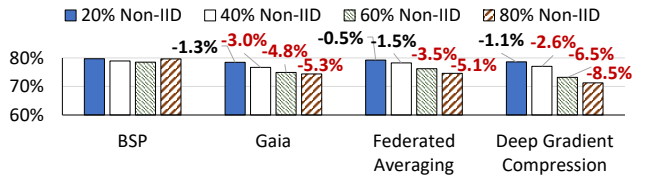


Figure 6. Top-1 validation accuracy for GN-LeNet over CIFAR-10, varying the degree of skewness. The “-x%” label on each bar indicates the accuracy loss from BSP in the IID setting. Note: y-axis starts at 60% accuracy.

1) Partial non-IID data is also problematic. We see that for all three decentralized learning algorithms, partial non-IID data still cause major accuracy loss. Even with a small degree of non-IID data such as 40%, we still see 1.5%–3.0% accuracy loss. Thus, the problem of non-IID data does not occur only with exclusive label partitioning, and hence, the problem exists in a vast majority of practical settings.

2) Degree of skew often determines the difficulty level of the problem. The model accuracy gets worse with higher degrees of skew, and the accuracy gap between 80% and 20% non-IID data can be as large as 7.4% (Deep-GradientCompression). In general, we see that the problem gets more difficult with higher skew.

7. Our Approach: SkewScout

To address the problem of skewed label partitions, we introduce SkewScout, a general approach that enables communication-efficient decentralized learning over *arbitrarily* skewed label partitions.

7.1. Overview of SkewScout

We design SkewScout as a general module that can be seamlessly integrated with different decentralized learning algorithms, ML training frameworks, and ML applications. Figure 7 overviews the SkewScout design.

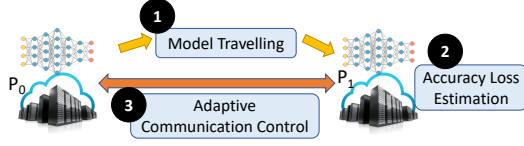


Figure 7. Overview of SkewScout

- **Estimate the degree of skew.** As §6 shows, knowing the degree of skew is very useful to determine an appropriate solution. To learn this information, SkewScout periodically moves the ML model from one data partition (P_k) to another during training (*model travelling*, ❶ in Figure 7). SkewScout then evaluates how well a model performs on a remote data partition by evaluating the model accuracy with a subset of training data on the remote node. As we already know the training accuracy of this model in its originated data partition, we can infer the *accuracy loss* in this remote data partition (❷).
- **Adaptive communication control (❸).** Based on the accuracy loss SkewScout learns from model travelling, SkewScout controls the amount of communication among data partitions to retain model quality. SkewScout controls the amount of communication by automatically tuning the hyper-parameters of the decentralized learning algorithm (§4.4). This tuning process is essentially solving an optimization problem that aims to minimize communication among data partitions while keeping accuracy loss within a reasonable threshold (§7.2 provides more details).

In essence, SkewScout handles non-IID data partitions by controlling communication based on accuracy loss. SkewScout is agnostic to the source of the loss, which may be from skewed label partitions or other forms of non-IID data (Appendix I). As long as increasing communication improves accuracy, SkewScout should be effective.

7.2. Mechanism Details

We now discuss the mechanisms of SkewScout in detail.

Accuracy Loss. The accuracy loss between data partitions represents the degree of model divergence. As §4.3 discusses, ML models in different data partitions tend to specialize for their training data, especially when we use decentralized learning algorithms to reduce communication.

We study accuracy loss under Gaia, for hyper-parameter choices $T_0=2\%, 5\%, 10\%, 20\%$, in the IID and non-IID settings. We find that accuracy loss changes drastically from the IID setting (0.4% on average) to the Non-IID setting (39.6% on average), and that lower T_0 results in smaller accuracy loss in the non-IID setting. See Appendix H for more details. Accordingly, we can use accuracy loss (i) to estimate how much the models diverge from each other (reflecting training data differences); and (ii) to serve as an

objective function for communication control. The computation overhead to evaluate accuracy loss is quite small because we only run inference with a small fraction of training data, and we only do so once in a while (we empirically find that once every 500 mini-batches is frequent enough).

Communication Control. The goal of communication control is to retain model quality while minimizing communication among data partitions. We achieve this by solving an optimization problem, which aims to minimize the communication while keeping the *accuracy loss* to a small threshold σ_{AL} so that we can control model divergence caused by non-IID data partitions. We solve this optimization problem periodically after we estimate the accuracy loss with model travelling. Specifically, our target function is:

$$\operatorname{argmin}_{\theta} \left(\lambda_{AL} (\max(0, AL(\theta) - \sigma_{AL})) + \lambda_C \frac{C(\theta)}{CM} \right) \quad (1)$$

where $AL(\theta)$ is the accuracy loss based on the previously selected hyper-parameter θ (we memoize the most recent value for each θ that has been explored), $C(\theta)$ is the amount of communication given θ , CM is the communication cost for the whole ML model, and λ_{AL} , λ_C are given parameters to determine the weights of accuracy loss and communication, respectively. We can employ various algorithms with Equation 1 to select θ such as hill climbing, stochastic hill climbing (Russell & Norvig, 2016), and simulated annealing (Van Laarhoven & Aarts, 1987).

7.3. Evaluation Results

We implement and evaluate SkewScout in a GPU parameter server system (Cui et al., 2016) based on Caffe (Jia et al., 2014). We evaluate several aforementioned auto-tuning algorithms and we find that hill climbing provides the best results. As our primary goal is to minimize accuracy loss, we set $\lambda_{AL} = 50$ and $\lambda_{AC} = 1$. We select $\sigma_{AL} = 5\%$ to tolerate an acceptable accuracy variation during training, which does not reduce the final validation accuracy.

We compare SkewScout with two other baselines: (1) **BSP**: the most communication-heavy approach that retains model quality in all Non-IID settings; and (2) **Oracle**: the ideal, yet unrealistic, approach that selects the most communication-efficient θ that retains model quality by *running all possible* θ in each setting prior to measured execution. Figure 8 shows the communication savings over BSP for both SkewScout and Oracle when training with Gaia. Note that all results achieve *the same* validation accuracy as BSP. We make two observations.

First, SkewScout is much more effective than BSP in handling Non-IID settings. Overall, SkewScout achieves 9.6–34.1 $\times$ communication savings over BSP in various Non-IID settings without sacrificing model accuracy. As expected, SkewScout saves more communication with less skewed data because SkewScout can safely loosen communication.

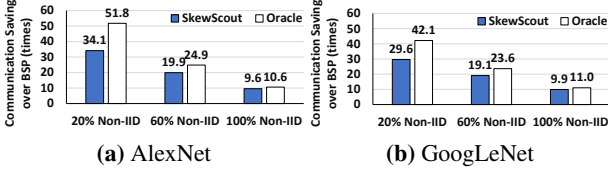


Figure 8. Communication savings over BSP with SkewScout and Oracle for training over CIFAR-10. All results achieve the same accuracy as BSP in the IID setting.

Second, SkewScout is not far from the ideal Oracle baseline. Overall, SkewScout requires only $1.1\text{--}1.5\times$ more communication than Oracle to achieve the same model accuracy. SkewScout cannot match the communication savings of Oracle because: (i) SkewScout needs to do model traveling periodically, which leads to some extra overheads; and (ii) for some θ , high accuracy loss at the beginning can still end up with a high quality model, which SkewScout cannot foresee. As Oracle requires *many* runs in practice, we conclude that SkewScout is an effective, one-pass solution for decentralized learning over non-IID data partitions.

8. Related Work

To our knowledge, this is the first study to show that skewed label partitions across devices/locations is a fundamental and pervasive problem for decentralized learning. Our study investigates various aspects of this problem, such as a real-world dataset, decentralized learning algorithms, batch normalization, and data skewness, as well as presenting our SkewScout approach. Here, we discuss related work.

Large-scale systems for centralized learning. There are many large-scale ML systems that aim to enable efficient ML training over *centralized* datasets using communication-efficient designs, such as relaxing synchronization requirements (Recht et al., 2011; Ho et al., 2013; Goyal et al., 2017) or sending fewer updates to parameter servers (Li et al., 2014a;b). These works assume the training data are centralized so they can be easily partitioned among the machines performing the training in an IID manner (e.g., by random shuffling). Hence, they are neither designed nor validated on non-IID data partitions.

Decentralized learning. Recent prior work proposes communication-efficient algorithms (e.g., (Hsieh et al., 2017; McMahan et al., 2017; Shokri & Shmatikov, 2015; Lin et al., 2018; Tang et al., 2018)) for ML training over *decentralized* datasets. However, as our study shows, these decentralized learning algorithms lose significant model quality in the Non-IID setting (§4). Some recent work studies the problem of non-IID data partitions. For example, instead of training a global model to fit non-IID data partitions, federated multi-task learning (Smith et al., 2017) trains local models for each data partition while leveraging other data partitions to improve model accuracy. How-

ever, this approach sidesteps the problem for global models, which are essential when a local model is unavailable (e.g., a brand new partition without training data) or ineffective (e.g., a partition with too few training examples for a class). Several recent works show significant accuracy loss for FederatedAveraging over non-IID data, and some propose algorithms to improve FederatedAveraging over non-IID data (Zhao et al., 2018; Li et al., 2019; Shoham et al., 2019; Karimireddy et al., 2019; Liang et al., 2019; Li et al., 2020a; Wang et al., 2020; Khaled et al., 2020). While the result of these works aligns with our observations, our study (i) broadens the problem scope to a variety of decentralized learning algorithms, ML applications, DNN models, and datasets, (ii) explores the problem of batch normalization and possible solutions, and (iii) designs and evaluates SkewScout, which can also complement these algorithms by controlling their hyper-parameters over arbitrarily skewed data partitions.

Non-IID dataset. Recent work offers non-IID datasets to facilitate the study of federated learning. For example, LEAF (Caldas et al., 2018) provides datasets that are partitioned in various ways. Luo et al. release 900 images collected from cameras in different locations, and they show severe skewed label distribution across cameras (Luo et al., 2019). Our study on geo-tagged mammals on Flickr shows the same problem at a much larger scale, and our dataset broadens the scope to include geo-distributed learning.

9. Conclusion

As most timely and relevant ML data are generated at different places, and often infeasible/impractical to collect centrally, decentralized learning provides an important path for ML applications to leverage these data. However, decentralized data are often generated at different contexts, which leads to a heavily understudied problem: *non-IID training data partitions*. We conduct a detailed empirical study of this problem for skewed label partitions, revealing three key findings. First, we show that training over skewed label partitions is a fundamental and pervasive problem for decentralized learning, as all decentralized learning algorithms in our study suffer major accuracy loss. Second, we find that DNNs with batch normalization are particularly vulnerable in the Non-IID setting, with even the most communication-heavy approach being unable to retain model quality. We further discuss the cause and potential solution to this problem. Third, we show that the difficulty level of this problem varies greatly with the degree of skew. Based on these findings, we present SkewScout, a general approach to minimizing communication while retaining model quality even for non-IID data. We hope that the findings and insights in this paper, as well as our open source code and dataset, will spur further research into the fundamental and important problem of non-IID data in decentralized learning.

Acknowledgments

We thank H. Brendan McMahan, Gregory R. Ganger, and the anonymous ICML reviewers for their valuable and constructive suggestions. Special thanks to Chenghao Zhang for his helps on setting up face recognition experiments. We also thank the members and companies of the PDL Consortium (Alibaba, Amazon, Datrium, Facebook, Google, Hewlett-Packard Enterprise, Hitachi, IBM, Intel, Microsoft, NetApp, Oracle, Salesforce, Samsung, Seagate, and Two Sigma) for their interest, insights, feedback, and support. We acknowledge the support of our industrial partners: Google, Huawei, Intel, Microsoft, and VMware. This work is supported in part by NSF. Most of the work was done when Kevin Hsieh was at Carnegie Mellon University.

References

- Ba, L. J., Kiros, J. R., and Hinton, G. E. Layer normalization. *CoRR*, abs/1607.06450, 2016.
- Bjorck, N., Gomes, C. P., Selman, B., and Weinberger, K. Q. Understanding batch normalization. In *NeurIPS*, 2018.
- Briggs, C., Fan, Z., and Andras, P. Federated learning with hierarchical clustering of local updates to improve training on non-IID data. *CoRR*, abs/2004.11791, 2020.
- Caldas, S., Wu, P., Li, T., Konečný, J., McMahan, H. B., Smith, V., and Talwalkar, A. LEAF: A benchmark for federated settings. *CoRR*, abs/1812.01097, 2018.
- Cui, H., Zhang, H., Ganger, G. R., Gibbons, P. B., and Xing, E. P. GeePS: Scalable deep learning on distributed GPUs with a GPU-specialized parameter server. In *EuroSys*, 2016. Software available at <https://github.com/cuihenggang/geeps>.
- Flickr. <https://www.flickr.com/>.
- Goodfellow, I. J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A. C., and Bengio, Y. Generative adversarial nets. In *NIPS*, 2014.
- Goyal, P., Dollár, P., Girshick, R. B., Noordhuis, P., Wesolowski, L., Kyrola, A., Tulloch, A., Jia, Y., and He, K. Accurate, large minibatch SGD: training ImageNet in 1 hour. *CoRR*, abs/1706.02677, 2017.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *CVPR*, 2016.
- Ho, Q., Cipar, J., Cui, H., Lee, S., Kim, J. K., Gibbons, P. B., Gibson, G. A., Ganger, G. R., and Xing, E. P. More effective distributed ML via a stale synchronous parallel parameter server. In *NIPS*, 2013.
- Hochreiter, S. and Schmidhuber, J. Long short-term memory. *Neural Computation*, 9(8), 1997.
- Hsieh, K., Harlap, A., Vijaykumar, N., Konomis, D., Ganger, G. R., Gibbons, P. B., and Mutlu, O. Gaia: Geo-distributed machine learning approaching LAN speeds. In *NSDI*, 2017.
- Huang, G. B., Ramesh, M., Berg, T., and Learned-Miller, E. Labeled faces in the wild: A database for studying face recognition in unconstrained environments. Technical Report 07-49, University of Massachusetts, Amherst, October 2007.
- Ioffe, S. Batch renormalization: Towards reducing minibatch dependence in batch-normalized models. In *NIPS*, 2017.
- Ioffe, S. and Szegedy, C. Batch normalization: Accelerating deep network training by reducing internal covariate shift. In *ICML*, 2015.
- Jia, Y., Shelhamer, E., Donahue, J., Karayev, S., Long, J., Girshick, R. B., Guadarrama, S., and Darrell, T. Caffe: Convolutional architecture for fast feature embedding. *CoRR*, 2014.
- Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., Bonawitz, K., Charles, Z., Cormode, G., Cummings, R., D'Oliveira, R. G. L., Rouayheb, S. E., Evans, D., Gardner, J., Garrett, Z., Gascón, A., Ghazi, B., Gibbons, P. B., Gruteser, M., Harchaoui, Z., He, C., He, L., Huo, Z., Hutchinson, B., Hsu, J., Jaggi, M., Javidi, T., Joshi, G., Khodak, M., Konečný, J., Korolova, A., Koushanfar, F., Koyejo, S., Lepoint, T., Liu, Y., Mittal, P., Mohri, M., Nock, R., Özgür, A., Pagh, R., Raykova, M., Qi, H., Ramage, D., Raskar, R., Song, D., Song, W., Stich, S. U., Sun, Z., Suresh, A. T., Tramèr, F., Vepakomma, P., Wang, J., Xiong, L., Xu, Z., Yang, Q., Yu, F. X., Yu, H., and Zhao, S. Advances and open problems in federated learning. *CoRR*, abs/1912.04977, 2019.
- Karimireddy, S. P., Kale, S., Mohri, M., Reddi, S. J., Stich, S. U., and Suresh, A. T. SCAFFOLD: stochastic controlled averaging for on-device federated learning. *CoRR*, abs/1910.06378, 2019.
- Khaled, A., Mishchenko, K., and Richtárik, P. Tighter theory for local SGD on identical and heterogeneous data. In *AISTATS*, 2020.
- Krizhevsky, A. Learning multiple layers of features from tiny images. Technical report, University of Toronto, 2009.
- Krizhevsky, A., Sutskever, I., and Hinton, G. E. ImageNet classification with deep convolutional neural networks. In *NIPS*, 2012.

- Kuznetsova, A., Rom, H., Alldrin, N., Uijlings, J. R. R., Krasin, I., Pont-Tuset, J., Kamali, S., Popov, S., Mallocci, M., Duerig, T., and Ferrari, V. The open images dataset V4: unified image classification, object detection, and visual relationship detection at scale. *CoRR*, abs/1811.00982, 2018.
- Laguel, Y., Pillutla, K., Malick, J., and Harchaoui, Z. Device heterogeneity in federated learning: A superquantile approach. *CoRR*, abs/2002.11223, 2020.
- LeCun, Y., Bottou, L., Bengio, Y., Haffner, P., et al. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11), 1998.
- Li, M., Andersen, D. G., Park, J. W., Smola, A. J., Ahmed, A., Josifovski, V., Long, J., Shekita, E. J., and Su, B. Scaling distributed machine learning with the parameter server. In *OSDI*, 2014a.
- Li, M., Andersen, D. G., Smola, A. J., and Yu, K. Communication efficient distributed machine learning with the parameter server. In *NIPS*, 2014b.
- Li, T., Sahu, A. K., Zaheer, M., Sanjabi, M., Talwalkar, A., and Smith, V. Federated optimization in heterogeneous networks. In *MLSys*, 2020a.
- Li, T., Sanjabi, M., Beirami, A., and Smith, V. Fair resource allocation in federated learning. In *ICLR*, 2020b.
- Li, X., Huang, K., Yang, W., Wang, S., and Zhang, Z. On the convergence of FedAvg on Non-IID data. *CoRR*, abs/1907.02189, 2019.
- Liang, X., Shen, S., Liu, J., Pan, Z., Chen, E., and Cheng, Y. Variance reduced local SGD with lower communication complexity. *CoRR*, abs/1912.12844, 2019.
- Lin, Y., Han, S., Mao, H., Wang, Y., and Dally, W. J. Deep gradient compression: Reducing the communication bandwidth for distributed training. In *ICLR*, 2018.
- Liu, C., Zoph, B., Neumann, M., Shlens, J., Hua, W., Li, L., Fei-Fei, L., Yuille, A. L., Huang, J., and Murphy, K. Progressive neural architecture search. In *ECCV*, 2018.
- Luo, J., Wu, X., Luo, Y., Huang, A., Huang, Y., Liu, Y., and Yang, Q. Real-world image datasets for federated learning. *CoRR*, abs/1910.11089, 2019.
- Mansour, Y., Mohri, M., Ro, J., and Suresh, A. T. Three approaches for personalization with applications to federated learning. *CoRR*, abs/2002.10619, 2020.
- McMahan, B., Moore, E., Ramage, D., Hampson, S., and y Arcas, B. A. Communication-efficient learning of deep networks from decentralized data. In *AISTATS*, 2017.
- Pascanu, R., Mikolov, T., and Bengio, Y. On the difficulty of training recurrent neural networks. In *ICML*, 2013.
- Recht, B., Ré, C., Wright, S. J., and Niu, F. Hogwild: A lock-free approach to parallelizing stochastic gradient descent. In *NIPS*, 2011.
- Russakovsky, O., Deng, J., Su, H., Krause, J., Satheesh, S., Ma, S., Huang, Z., Karpathy, A., Khosla, A., Bernstein, M., Berg, A. C., and Fei-Fei, L. ImageNet large scale visual recognition challenge. *IJCV*, 2015.
- Russell, S. J. and Norvig, P. *Artificial intelligence: a modern approach*. Malaysia; Pearson Education Limited,, 2016.
- Salimans, T. and Kingma, D. P. Weight normalization: A simple reparameterization to accelerate training of deep neural networks. In *NIPS*, 2016.
- Santurkar, S., Tsipras, D., Ilyas, A., and Madry, A. How does batch normalization help optimization? In *NeurIPS*, 2018.
- Shoham, N., Avidor, T., Keren, A., Israel, N., Benditkis, D., Mor-Yosef, L., and Zeitak, I. Overcoming forgetting in federated learning on non-iid data. *CoRR*, abs/1910.07796, 2019.
- Shokri, R. and Shmatikov, V. Privacy-preserving deep learning. In *CCS*, 2015.
- Smith, V., Chiang, C., Sanjabi, M., and Talwalkar, A. S. Federated multi-task learning. In *NIPS*, 2017.
- Szegedy, C., Liu, W., Jia, Y., Sermanet, P., Reed, S. E., Anguelov, D., Erhan, D., Vanhoucke, V., and Rabinovich, A. Going deeper with convolutions. In *CVPR*, 2015.
- Szegedy, C., Vanhoucke, V., Ioffe, S., Shlens, J., and Wojna, Z. Rethinking the inception architecture for computer vision. In *CVPR*, 2016.
- Tang, H., Lian, X., Yan, M., Zhang, C., and Liu, J. D²: Decentralized training over decentralized data. In *ICML*, 2018.
- United Nation Statistics Division. Methodology: Standard country or area codes for statistical use (m49), 2019.
- Valiant, L. G. A bridging model for parallel computation. *Commun. ACM*, 1990.
- Van Laarhoven, P. J. and Aarts, E. H. Simulated annealing. In *Simulated annealing: Theory and applications*, pp. 7–15. Springer, 1987.
- Wang, H., Yurochkin, M., Sun, Y., Papailiopoulous, D. S., and Khazaeni, Y. Federated learning with matched averaging. In *ICLR*, 2020.

Wen, Y., Zhang, K., Li, Z., and Qiao, Y. A discriminative feature learning approach for deep face recognition. In *ECCV*, 2016.

Wu, Y. and He, K. Group normalization. In *ECCV*, 2018.

Yi, D., Lei, Z., Liao, S., and Li, S. Z. Learning face representation from scratch. *CoRR*, abs/1411.7923, 2014.

Yu, T., Bagdasaryan, E., and Shmatikov, V. Salvaging federated learning by local adaptation. *CoRR*, abs/2002.04758, 2020.

Zhao, Y., Li, M., Lai, L., Suda, N., Civan, D., and Chandra, V. Federated learning with non-IID data. *CoRR*, abs/1806.00582, 2018.