

E²-Train: Training State-of-the-art CNNs with Over 80% Energy Savings

Yue Wang^{◇*}, Ziyu Jiang^{†*}, Xiaohan Chen^{†*}, Pengfei Xu[◇], Yang Zhao[◇],
Yingyan Lin^{◇‡} and Zhangyang Wang^{†‡}

[†]Department of Computer Science and Engineering, Texas A&M University

[◇]Department of Electrical and Computer Engineering, Rice University

[†]{jiangziyu, chernxh, atlaswang}@tamu.edu

[◇]{yw68, px5, zy34, yingyan.lin}@rice.edu

<http://rtml.eiclab.net/publications/e2-train>

Abstract

Convolutional neural networks (CNNs) have been increasingly deployed to edge devices. Hence, many efforts have been made towards efficient CNN *inference* in resource-constrained platforms. This paper attempts to explore an orthogonal direction: how to conduct more energy-efficient *training* of CNNs, so as to enable on-device training. We strive to reduce the energy cost during training, by dropping unnecessary computations from three complementary levels: stochastic mini-batch dropping on the **data level**; selective layer update on the **model level**; and sign prediction for low-cost, low-precision back-propagation, on the **algorithm level**. Extensive simulations and ablation studies, with *real energy measurements* from an FPGA board, confirm the superiority of our proposed strategies and demonstrate remarkable energy savings for training. For example, when training ResNet-74 on CIFAR-10, we achieve aggressive energy savings of >90% and >60%, while incurring a top-1 accuracy loss of only about 2% and 1.2%, respectively. When training ResNet-110 on CIFAR-100, an over 84% training energy saving is achieved without degrading inference accuracy.

1 Introduction

The increasing penetration of intelligent sensors has revolutionized how Internet of Things (IoT) works. For visual data analytics, we have witnessed the record-breaking predictive performance achieved by convolutional neural networks (CNNs) [1, 2, 3]. Although such high performance CNN models are initially learned in data centers and then deployed to IoT devices, we have witnessed an increasing necessity for the model to continue learning and updating itself *in situ*, such as for personalization for different users, or incremental/lifelong learning. Ideally, this learning/retraining process should take place on the device. Comparing to cloud-based retraining, training locally helps avoid transferring data back and forth between data centers and IoT devices, reduce communication cost/latency, and enhance privacy.

However, training on IoT devices is non-trivial and more time/resource-consuming yet much less explored than inference. IoT devices, such as smart phones and wearables, have limited computation and energy resources, which are even stringent for inference. Training CNNs consumes magnitudes higher number of computations than one inference. For example, training ResNet-50 for only one 224×224 image can take up to 12 GFLOPs (vs. 4GFLOPs for inference), which can easily drain a mobile phone battery when training with batch images [4]. This mismatch between the limited

*The first three authors (Yue Wang, Ziyu Jiang, Xiaohan Chen) contributed equally.

‡Correspondence should be addressed to: Yingyan Lin and Zhangyang Wang.

resources of IoT devices and the high complexity of CNNs is only getting worse because the network structures are getting more complex as they are designed to solve harder and larger-scale tasks [5].

This paper considers the most standard CNN training setting, assuming both the model structure and the dataset to be given. This “basic” training setting is not usually the realistic IoT case, but we address it as a starting point (with familiar benchmarks), and an opening door towards obtaining a toolbox that later may be extended to online/transfer learning too (see Section 5). Our goal is to **reduce the total energy cost of training**, which is complicated by a myriad of factors: from per-sample (mini-batch) complexity (both feed-forward and backward computations), to the empirical convergence rate (how many epochs it takes to converge), and more broadly, hardware/architecture factors such as data access and movements [6, 7, 8]. Despite a handful of works on efficient, accelerated CNN training [9, 10, 11, 12, 13], they mostly focus on reducing the total **training time** in *resource-rich* settings, such as by *distributed training* in large-scale GPU clusters. In contrast, our focus is to trim down the total **energy cost** of *in-situ, resource-constrained training*. This represents an orthogonal (and less studied) direction to [9, 10, 11, 12, 13, 14, 15], although the two can certainly be combined.

To unleash the potential of more energy-efficient in-situ training, we look at the full CNN training lifecycle closely. With the goal of “squeezing out” unnecessary costs, we raise three curious questions:

- *Q1: Are all samples always required throughout training:* is it necessary to use all training samples in all epochs?
- *Q2: Are all parts of the entire model equally important during training:* does every layer or filter have to be updated every time?
- *Q3: Are precise gradients indispensable for training:* can we efficiently compute and update the model with approximate gradients?

The above three questions only represent our “first stab” ideas for exploring energy-efficient training, whose full scope is much more profound. By no means do our questions above represent all possible directions. We envision that many other recipes can be blended too, such as training on lower bit precision or input resolution [16, 17]. We also recognize that energy-efficient CNN training should be jointly considered with hardware/architecture co-design [18, 19], which is beyond the current work.

Motivated by the above questions, this paper proposes a novel energy efficient CNN training framework dubbed **E²-Train**. It consists of three complementary aspects of efforts to trim down unnecessary training computations and data movements, each addressing one of the above questions:

- **Data-Level: Stochastic mini-batch dropping (SMD).** We show that CNN training could be accelerated by a “frustratingly easy” strategy: randomly skipping mini-batches with 0.5 probability throughout training. This could be interpreted as data sampling with (limited) replacements, and is found to incur minimal accuracy loss (and sometimes even increases accuracy).
- **Model-Level: Input-dependent selective layer update (SLU).** For each minibatch, we select a different subset of the CNN layers to be updated. The input-adaptive selection is based on a low-cost gating function jointly learned during training. While similar ideas were explored in efficient inference [20], for the first time, they are applied to and evaluated for training for the first time.
- **Algorithm-Level: Predictive sign gradient descent (PSG).** We explore the usage of an extremely low-precision gradient descent algorithm called SignSGD, which has recently found both theoretical and experimental grounds [21]. The original algorithm still requires the full gradient computation and therefore does not save energy. We create a novel “predictive” variant, that could obtain the sign without computing the full gradient via low-cost, bit-level prediction. Combined with mixed-precision design, it decreases both computation and data-movement costs.

Besides mainly experimental explorations, we find E²-Train has many interesting links to recent CNN training theories, e.g., [22, 23, 24, 25]. We evaluate E²-Train in comparison with its closest state-of-the-art competitors. To measure its actual performance, E²-Train is implemented and evaluated on **an FPGA board**. The results show that the CNN model applied with E²-Train consistently achieves higher training energy efficiency with marginal accuracy drops.

2 Related Works

Accelerated CNN training. A number of works have been devoted to accelerating training, in a resource-rich setting, by utilizing communication-efficient distributed optimization and larger mini-batch sizes [9, 10, 11, 12]. The latest work [13] combined distributed training with a mixed precision

framework, leading to training AlexNet within 4 minutes. However, their goals and settings are distinct from ours - while the distributed training strategy can reduce time, it will actually incur more total energy overhead, and is clearly not applicable to on-device resource-constrained training.

Low-precision training. It is well known that CNN training can be performed under substantially lower precision [14, 15, 16], rather than by using full-precision floats. Specifically, training with quantized gradients has been well studied in the distributed learning, whose main motivation is to reduce the communication cost during gradient aggregations between workers [21, 26, 27, 28, 29, 30]. A few works considered to only transmit the coordinates of large magnitudes [31, 32, 33]. Recently, the SignSGD algorithm [21, 26] even showed the feasibility of using one-bit gradients (signs) during training, without notably hampering the convergence rate or final result. However, most algorithms are optimized for distributed communication efficiency, rather than for reducing training energy costs. Many of them, including [21], need to first compute full-precision gradients and then quantize them.

Efficient CNN inference: Static and Dynamic. Compressing CNNs and speeding up their inference have attracted major research interests in recent years. Representative methods include weight pruning, weight sharing, layer factorization, and bit quantization, to name just a few [34, 35, 36, 37, 38].

While model compression presents “static” solutions for improving inference efficiency, a more interesting recent trend is to look at *dynamic inference* [20, 39, 40, 41, 42, 43] to reduce the latency, i.e., selectively executing subsets of layers in the network conditioned on each input. That sequential decision making process is usually controlled by low-cost gating or policy networks. This mechanism has also been applied to improve inference energy efficiency [44, 45].

In [46], a unique bit-level prediction framework called *PredictiveNet* was presented to accelerate CNN inference at a lower level. Since CNN layer-wise activations are usually highly sparse, the authors proposed to predict those zero locations using low-cost bit predictors, thereby bypassing a large fraction of energy-dominant convolutions without modifying the CNN structure.

Energy-efficient training is different from and more complicated than its inference counterpart. However, many insights gained from the latter can be lent to the former. For example, a recent work [47] showed that performing active channel pruning during training can accelerate the empirical convergence. Our proposed model-level SLU is inspired by [20]. The algorithm-level PSG also inherits the idea of bit-level low-cost prediction from [46].

3 The Proposed Framework

3.1 Data-Level: Stochastic mini-batch dropping (SMD)

We first adopt a straightforward, seemingly naive, yet surprisingly effective stochastic mini-batch dropping (SMD) strategy (see Fig. 1), to aggressively reduce the training cost by letting it see less mini-batches. At each epoch, SMD simply skips every mini-batch with a default probability of 0.5. All other training protocols, such as learning rate schedule, remain unchanged. Compared to the normal training, SMD can directly half the training cost, if both were trained with the same number of epochs. Yet amazingly, we observe in our experiments that SMD usually leads to negligible accuracy decrease, and sometimes even an increase (see Sec. 4). Why? We discuss possible explanations below.

SMD can be interpreted as sampling *with limited replacement*. To understand this, think of combining two consecutive SMD-enforced epochs into one, then it has the same number of mini-batches as one full epoch; but within it each training sample now has a 0.25, 0.5, and 0.25 probability, of being sampled 2, 1, and 0 times, respectively. The conventional wisdom is that for stochastic gradient descent (SGD), in each epoch, the mini-batches are sampled i.i.d. from data *without replacement* (i.e., each sample occurs exactly once per epoch) [48, 49, 50, 51, 52]. However, [22] proved that sampling mini-batches *with replacement* has a larger variance than sampling without replacement, and consequently SGD may have better regularization properties.

Alternatively, SMD could also be viewed as a special data augmentation way that injects more sampling noise to perturb training distribution every epoch. Past works [53, 54, 55] have shown that specific kinds of random noise aid convergence through escaping from saddle points or less generalizable minima. The structured sampling noise caused by SMD might aid this exploration.

Besides, [23, 56, 57] also showed that an importance sampling scheme that focuses on training more with “informative” examples leads to faster convergence under resource budgets. They implied that

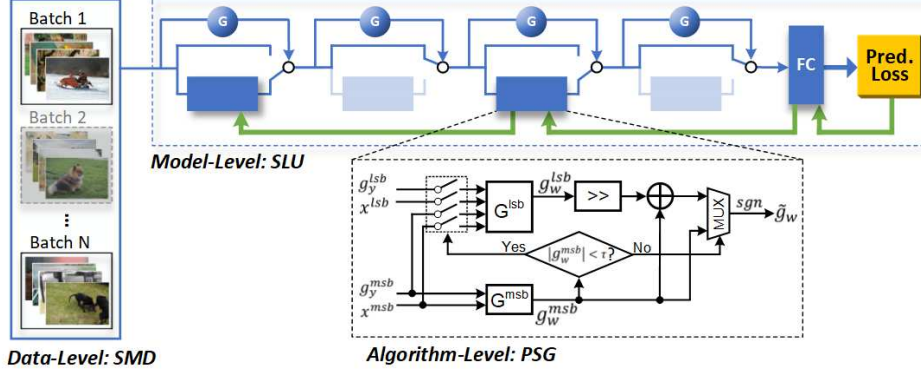


Figure 1: An illustration of the proposed framework. SLU: each blue circle G indicates an RNN gate and each blue square under G indicates one block of layers in the base model. Green arrows denote the backward propagation. The RNN gates generate strategies to select which layers to train for each input. In this specific example, the second and fourth blocks are “skipped” for **both feedforward and backward computations**. Only the first and third blocks are updated. The details of SMD and PSG are described in the main text.

the mini-batch dropping can be selective based on certain information criteria instead of stochastic. We use SMD because it has zero overhead, but more effective dropping options might be available if low-cost indicators of mini-batch importance can be identified: we leave this as future work.

3.2 Model-Level: Input-dependent selective layer update (SLU)

[20] proposed to dynamically skip a subset of layers for different inputs, in order to adaptively accelerate the feed-forward inference. However, [20] called for a post process after supervised training, i.e., to refine the dynamic skipping policy via reinforcement learning, thus causing undesired extra training overhead. We propose to extend the idea of dynamic inference to the training stage, i.e., dynamically skipping a subset of layers during **both feed-forward and back-propagation**. Crucially, we show that by adding an *auxiliary regularization*, such dynamic skipping can be learned from scratch and obtain satisfactory performance: **neither post refinement nor extra training iterations are required**. That is critical for dynamic layer skipping to be useful for energy-efficient training: we term this extended scheme as input-dependent selective layer update (SLU).

As depicted in Fig. 1, given a base CNN to be trained, we follow [20] to add a light-weight RNN gating network per layer block. Each gate takes the same input as its corresponding layer, and outputs soft-gating indicators between $[0,1]$, which are then used as the skipping probability, i.e., the higher the value is, the more likely that layer will be selected. Therefore, each layer will be adaptively selected or skipped, depending on the inputs. We will only select the layers activated by gates. The RNN gates cost less than 0.04% feed-forward FLOPs than the base models; hence their energy overheads are negligible. More details can be found in the supplement.

[20] first trained the gates in a supervised way together with the base model. Observing that such learned routing policies were often not sufficiently efficient, they used reinforcement learning post-processing to learn more aggressive skipping afterwards. While this is fine for the end goal of dynamic inference, we hope to get rid of the post-processing overhead. We incorporate the computational complexity regularization into the objective function to overcome this hurdle, defined as:

$$\min_{W, G} L(W, G) + \alpha C(W, G) \quad (1)$$

Here, α is a weighting coefficient of the computational complexity regularization, and W and G denote the parameters of the base model and the gating network, respectively. Also, $L(W, G)$ denotes the prediction loss, and $C(W, G)$ is calculated by accumulating the computational cost (FLOPs) of the layers that are selected. The regularization explicitly encourages learning more “parisious” selections throughout the training. We find that such SLU-regularized training leads to almost the same number of epochs to converge compared to standard training, i.e., SLU does not sacrifice empirical convergence speed. As a side effect, SLU will naturally yield CNNs with dynamic inference capability. Though not the focus of this paper, we find the CNN trained with SLU reaches a comparable accuracy-efficiency trade-off over the ones trained with the approach in [20].

The practice of SLU seems to align with several recent theories on CNN training. In [58], the authors suggested that “not all layers are created equal” for training. Specifically, some layers are

critical to be intensively updated for improving final predictions, while others are insensitive along training. There exist “non-critical” layers that barely change their weights throughout training: even resetting those layers in a trained model to their initial value has few negative consequences. The more recent work [25] further confirmed the phenomenon, though how to identify those non-critical model parts at the early training stage remains unclear. [59, 60] also observed that different samples might activate different sub-models. Those inspiring theories, combined with the dynamic inference practice, motivate us to propose SLU for more efficient training.

3.3 Algorithm-Level: Predictive sign gradient descent (PSG)

It is well recognized that low-precision fixed-point implementation is a very effective knob for achieving energy efficient CNNs, because both CNNs’ computational and data movement costs are approximately a quadratic function of their employed precision. For example, a state-of-the-art design [61] shows that adopting 8-bit precision for a multiplication, adder, and data movement can reduce energy cost by 95%, 97%, and 75%, respectively, as compared to that of a 32-bit floating point design when evaluated in a commercial 45nm CMOS technology.

The successful adoption of extremely low-precision (binary) gradients in SignSGD [21] is appealing, as it might lead to reducing both weight update computation and data movements. However, directly applying the original SignSGD algorithm for training will not save energy, because it actually computes the full-precision gradient first before taking the signs. We propose a novel predictive sign gradient descent (PSG) algorithm, which predicts the sign of gradients using low-cost bit-level predictors, thereby completely bypassing the costly full-gradient computation.

We next introduce how the gradients of weights are updated in PSG. Assume the following notations: the full precision and **most significant bits** (the latter, **MSB** part, is adopted as PSG’s low-cost predictors) of the input x and the gradient of the output g_y are denoted as (B_x, B_g) and $(B_x^{\text{msb}}, B_g^{\text{msb}})$, respectively, where the corresponding input and the gradient of the output for PSG’s predictors are denoted as x^{msb} and g_y^{msb} , respectively. As such, the quantization noise for the input and the gradient of the output are $q_x = x - x^{\text{msb}}$ and $q_{g_y} = g_y - g_y^{\text{msb}}$, respectively. Similarly, after back-propagation, we denote the full-precision and low-precision (i.e., taking the most significant bits (MSBs)) gradient of the weight as g_w and g_w^{msb} , respectively, the latter of which is computed using x^{msb} and g_y^{msb} . Then, with an empirically pre-selected threshold τ , PSG updates the i -th weight gradient as follows:

$$\tilde{g}_w[i] = \begin{cases} \text{sgn}(g_w^{\text{msb}}[i]) & , |g_w^{\text{msb}}[i]| \geq \tau \\ \text{sgn}(g_w[i]) & , \text{otherwise} \end{cases} \quad (2)$$

Note that in hardware implementation, the computation to obtain g_w^{msb} is embedded within that of g_w . Therefore, the PSG’s predictors do not incur energy overhead.

PSG for energy-efficient training. Recent work [16] has shown that most of the training process is robust to reduced precision (e.g., 8 bits instead of 32 bits), except for the weight gradient calculations and updates. Taking their learning, we similarly adopt a higher precision for the gradients than the inputs and weights, i.e., $B_{g_y} > B_x = B_w$. Specifically, when training with PSG, we first compute the predictors using B_x^{msb} (e.g., $B_x^{\text{msb}} = 4$) and $B_{g_y}^{\text{msb}}$ (e.g., $B_{g_y}^{\text{msb}} = 10$), and then update the weights’ gradients following Eq. (2). The further energy savings of training with PSG over the fixed-point training [16] results from the fact that the predictors computed using x^{msb} and g_y^{msb} require exponentially less computational and data movement energy.

Prediction guarantee of PSG. We analyze the probability of PSG’s prediction failure to discuss its performance guarantee. Specifically, if denoting the sign prediction failure produced by Eq. (2) as H , it can be proved that this probability is upbounded as follows,

$$P(H) \leq \Delta_x^2 E_1 + \Delta_{g_y}^2 E_2, \quad (3)$$

where $\Delta_x = 2^{-(B_x^{\text{msb}}-1)}$ and $\Delta_{g_y} = 2^{-(B_{g_y}^{\text{msb}}-1)}$ are the quantization noise step sizes of x^{msb} and g_y^{msb} , respectively, and E_1 and E_2 are given in the Appendix along with the proof of Eq. (3), which inherits the spirit in [47, 62]. Eq. (3) shows that the prediction failure probability of PSG is upbounded by a term that degrades exponentially with the precision assigned to the predictors, indicating that this failure probability can be very small if the predictors are designed properly.

Adaptive threshold. Training with PSG might lead to sign flips in the weight gradients as compared to that of the floating point one, which only occurs when the latter has a small magnitude and thus the

quantization noise of the predictors causes the sign flips. Therefore, it is important to properly select a threshold (e.g., τ in Eq.(2)) that can optimally balance this sign flip probability and the achieved energy savings. We adopt an adaptive threshold selection strategy because the dynamic range of gradients differ significantly from layers to layers: instead of using a fixed number, we will tune a ratio $\beta \in (0, 1)$ which yields the adaptive threshold as $\hat{\tau} = \beta \max_i \{g_w^{\text{msb}}[i]\}$.

4 Experiments

4.1 Experiment setup

Datasets: We evaluate our proposed techniques on two datasets: CIFAR-10 and CIFAR-100. Common data augmentation methods (e.g., mirroring/shifting) are adopted, and data are normalized as in [63]. **Models:** Three popular backbones, ResNet-74, ResNet-110 [64], and MobileNetV2 [65], are considered. For evaluating each of the three proposed techniques (i.e., SMD, SLU, and PSG), we consider various experimental settings using the ResNet-74 and CIFAR-10 dataset for ablation studies, as described in Sections 4.2-4.5. ResNet-110 and MobileNetV2 results are reported in Section 4.6. Top-1 accuracies are measured for CIFAR-10, and both top-1 and top-5 accuracies for CIFAR-100. **Training settings:** We adopt the training settings in [64] for the baseline default configurations. Specifically, we use an SGD with a momentum of 0.9 and a weight decaying factor of 0.0001, and the initialization introduced in [66]. Models are trained for 64k iterations. For experiments where PSG is used, the initial learning rate is adjusted to 0.03 as SignSGD[21] suggested small learning rates to benefit convergence. For others, the learning rate is initially set to be 0.1 and then decayed by 10 at the 32k and 48k iterations, respectively. We also employed the stochastic weight averaging (SWA) technique [67], which was found to notably stabilize training, when PSG is adopted.

Real energy measurements using FPGA: As the energy cost of CNN inference/training consists of both computational and data movement costs, the latter of which is often dominant but can not be captured by the commonly used metrics, such as the number of FLOPs [6], we evaluate the proposed techniques against the baselines in terms of accuracy and real measured energy consumption. Specifically, unless otherwise specified, all the energy or energy savings are obtained **through real measurements** by training the corresponding models and datasets in a state-of-the-art FPGA [68], which is a diligent ZedBoard Zynq-7000 ARM/FPGA SoC Development Board. Fig. 2 shows our FPGA measurement setup, in which the FPGA board is connected to a laptop through a serial port and a power meter. In particular, the training settings are downloaded from the laptop to the FPGA board, and the real-measured energy consumption is obtained via the power meter for the whole training process and then sent back to the laptop. **All energy results are measured from FPGA.**

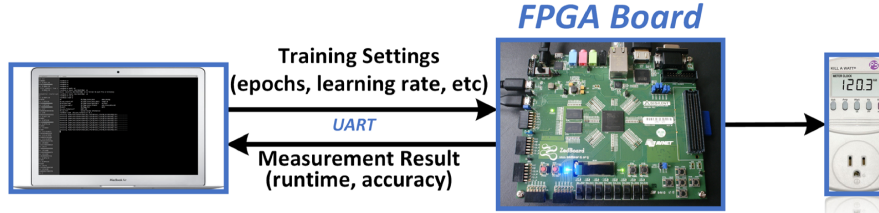


Figure 2: The energy measurement setup with (from left to right) a MAC Air laptop, a Xilinx FPGA board [68], and a power meter.

4.2 Evaluating stochastic mini-batch dropping

We first validate the energy savings achieved by SMD against a few “off-the-shelf” options: (1) can we train with the standard algorithm, using fewer iterations and otherwise the same training protocol? (2) can we train with the standard algorithm, using fewer iterations but properly increased learning rates? Two sets of carefully-designed experiments are presented below for addressing them.

Training with SMD vs. standard mini-batch (SMB): We first evaluate SMD over the standard mini-batch (SMB) training, which uses all (vs. 50% in SMD) mini-batch samples. As shown in Fig. 3a when the energy ratio is 1 (i.e., training with SMB + 64k iterations vs. SMD + 128k iterations), the proposed SMD technique is able to boost the inference accuracy by 0.39% over the standard way.

We next “naively” suppress the energy cost of SMB, by reducing training iterations. Specifically, we reduce the SMB training iterations to be $\{\frac{6}{12}, \frac{7}{12}, \frac{8}{12}, \frac{9}{12}, \frac{10}{12}, \frac{11}{12}, 1\}$ of the original one. Note the learning rate schedule (e.g., when to reduce learning rates) will be scaled proportionally with the total iteration number too. For comparison, we conduct experiments of training with SMD when

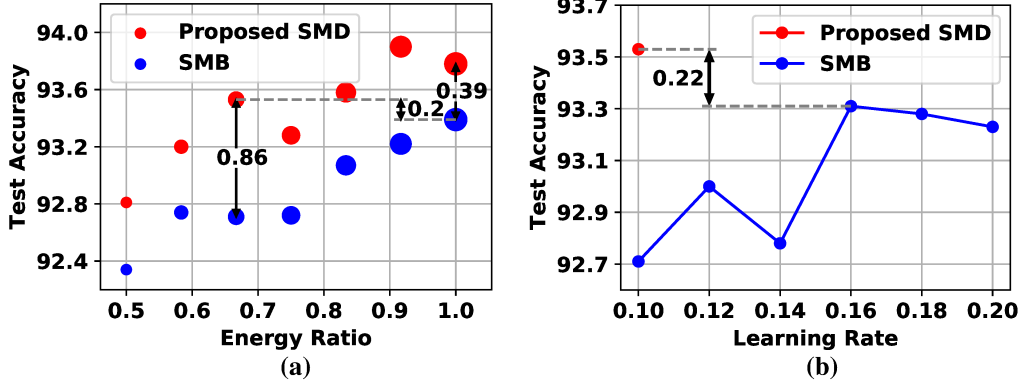


Figure 3: The top-1 accuracy of CIFAR-10 testing using the Resnet-74 model when: (a) training with SMD versus the standard mini-batch (SMB) method, with the training energy ratios of each ranging from 0.5 to 1. The size of markers is drawn proportionally with the measured training energy cost; and (b) training with SMD, and SMB with different increased learning rates – all under the same training energy budget.

the number of equivalent training iterations are the same as those of the SMB cases. Fig. 3a shows that training with SMD consistently achieves a higher inference accuracy than SMB with the margin ranging from 0.39% to 0.86%. Furthermore, we can see that training with SMD reduces the training energy cost by 0.33 while boosting the inference accuracy by 0.2% (see the cases of SMD under the energy ratio of 0.67 vs. SMB under the energy ratio of 1, respectively, in Fig. 3a), as compared to SMB. We adopt SMD under this energy ratio of 0.67 in all the remaining experiments.

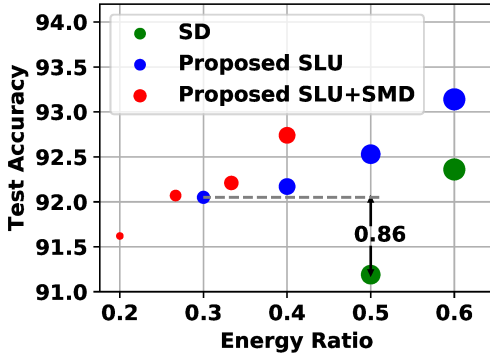


Figure 4: Inference accuracy vs. energy ratios, where the energy ratios are obtained by normalizing the corresponding energy over that of the original one (SMB + 64k iterations).

Table 1: The accuracy of SMD on other datasets and backbones (energy ratio 0.67).

Dataset	Backbone	Accuracy	
		SMB	SMD
CIFAR-10	ResNet-110	92.75%	93.05%
CIFAR-100	ResNet-74	71.11%	71.37%

4.3 Evaluating selective layer update

Our current SLU experiments are based on CNNs with residual connections, partially because they dominate in SOTA CNNs. We will extend SLU to other model structures in future work. We evaluate the proposed SLU by comparing it with stochastic depth (SD) [69], a technique originally developed for training very deep networks effectively, by updating only a random subset of layers at each mini-batch. It could be viewed as a “random” version of SLU (which uses learned layer selection). We follow all suggested settings in [69]. For a fair comparison, we adjust the hyper-parameter p_L [69], so that the SD dropping ratio is always the same as SLU.

From Fig. 4, training with SLU consistently achieves higher inference accuracy than SD when their training energy costs are the same. It is further encouraging to observe that training with SLU could

We repeated training ResNet-74 on CIFAR-10 using SMD for 10 iterations with different random initializations. The accuracy standard deviation is only 0.132%, showing high stability. We also conducted more experiments with different backbones and datasets. As shown in Tab. 1, SMD is consistently better than SMB.

Training with SMD vs. SMB + increased learning rates:

We further compare with SMB with tuned/larger learning rates, conjecturing that it would accelerate convergence by possibly reducing the training epochs needed. The results are summarized in Fig. 3b. Specifically, when the number of iterations are reduced by $\frac{1}{3}$, we do a grid search of learning rates, with a step size from 0.02 between [0.1, 0.2]. All compared methods are set with the same training energy budget. Fig. 3b demonstrates that while increasing learning rates seems to improve SMB’s energy efficiency over sticking to the original protocol, our proposed SMD still maintains a clear advantage of at least 0.22%.

even achieve higher accuracy sometimes in addition to saving energy. For example, comparing the cases when training with SLU + an energy ratio of 0.3 (i.e., 70% energy savings) and that of SD + an energy ratio of 0.5, the proposed SLU technique is able to reduce the training energy cost by 20% while boosting the inference accuracy by 0.86%. These results endorse the usage of data-driven gates instead of random dropping, in the context of energy-efficient training. Training with SLU + SMD combined further boosts the accuracy while reducing the energy cost. Furthermore, 20 trials of SLU experiments to ResNet38 on CIFAR-10 conclude that, with a 95% confidence level, that the confidence interval for the mean of the top-1 accuracy and the energy savings are [92.47%, 92.58%] (baseline: 92.50%) and [39.55%, 40.52%], respectively, verifying SLU’s trustworthy effectiveness.

4.4 Evaluating predictive sign gradient descent

We evaluate PSG against two alternatives: (1) the 8-bit fixed point training proposed in [16]; and (2) the original SignSGD [21]. For all experiments in Sections 4.4 and 4.5, we adopt 8-bit precision for the activations/weights and 16-bit for the gradients. The corresponding precision of the predictors are 4-bit and 10-bit, respectively. We use an adaptive threshold (see Section 3.3) of $\beta = 0.05$. More experiment details are in the Appendix.

Table 2: Comparing the inference accuracy and achieved energy savings (over the 32-bit floating point training) when training with SGD, 8-bit fixed point [16], SignSGD, and PSG using Resnet-74 and CIFAR-10.

Method	32-bit SGD	8-bit[16]	SignSGD[21]	PSG
Accuracy	93.52%	93.24%	92.54%	92.59%
Energy savings	-	38.62%	-	63.28%

Table 3: The inference accuracy and energy savings (over the 32-bit floating point training) of the proposed E²-Train under different (averaged) SLU skipping ratios and adaptive thresholds (i.e., β in Section 3.3) when using Resnet-74 and CIFAR-10.

Skipping	20%	40%	60%
Accuracy ($\beta=0.05$)	92.12%	91.84%	91.36%
Accuracy ($\beta=0.1$)	92.15%	91.72%	90.94%
Computational savings	80.27%	85.20%	90.13%
Energy savings	84.64%	88.72%	92.81%

As shown in Table 2, while the 8-bit fixed point training in [16] saves about 39% training energy (going from 32-bit to 8-bit in general leads to about 80% energy savings, which is compromised by their employed 32-bit gradients in this case) with a marginal accuracy loss of 0.28% as compared to the 32-bit SGD, the proposed PSG almost doubles the training energy savings (63% vs. 39% for [16]) with a negligible accuracy loss of 0.65% (93.24% vs. 92.59% for [16]). Interestingly, PSG slightly boosts the inference accuracy by 0.05% while saving 63% energy, i.e., $3\times$ better training energy efficiency with a slightly better inference accuracy, compared to SignSGD [21]. Besides, as we observed, the ratio of using g_w^{msb} for sign prediction typically remains at least 60% throughout the training process, given an adaptive threshold $\beta = 0.05$.

4.5 Evaluating E²-Train: SMD + SLU + PSG

We now evaluate the proposed E²-Train framework, which combines the SMD, SLU, and PSG techniques. As shown in Table 3, we can see that E²-Train: (1) indeed can further boost the performance as compared to training with SMD+SLU (e.g., E²-Train achieves a higher accuracy of 0.5% (92.1% vs. 91.6% (see Fig.4 at the energy ratio of 0.2) of training with SMD+SLU, when achieving 80% energy savings); and (2) can achieve an extremely aggressive energy savings of $> 90\%$ and $> 60\%$, while incurring an accuracy loss of only about 2.0% and 1.2%, respectively, as compared to that of the 32-bit floating point SGD (see Table 2), i.e., up to $9\times$ better training energy efficiency with small accuracy loss.

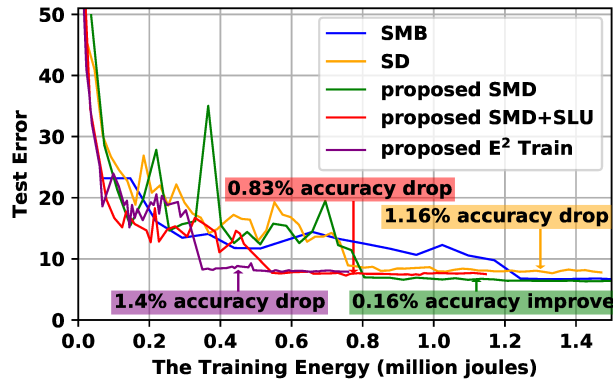


Figure 5: Inference accuracy vs. energy costs, at different stages of training, (i.e., the empirical convergence curves), when training with SMB, SD, merely SLU, SLU + SMD, and E²-Train on CIFAR-10 with ResNet-74.

Impact on empirical convergence speed. We plot the training convergence curves of different methods in Fig. 5, with the x-axis represented in the alternative form of training energy costs (up to

current iteration). We observe that E²-Train does not slow down the empirical convergence. In fact, it even makes the training loss decrease faster in the early stage.

Experiments on adapting a pre-trained model. We perform a proof-of-concept experiment for CNN fine-tuning by splitting CIFAR-10 training set into half, where each class was i.i.d. split evenly. We first pre-train ResNet-74 on the first half, then fine-tune it on the second half. During fine-tuning, we compare two energy-efficient options: (1) fine-tuning only the last FC layer using standard training; (2) fine-tuning all layers using E²-Train. With all hyperparameters being tuned to best efforts, the two fine-tuning methods improve over the pre-trained model top-1 accuracy by [0.30%, 1.37%] respectively, while (2) saves 61.58% more energy (FPGA-measured) than (1). This shows that E²-Train is the preferred option, leading to a higher accuracy and more energy savings

Table 4 evaluates E²-Train and its ablation baselines on various models and more datasets. The conclusions are aligned with the ResNet-74 cases. Remarkably, on CIFAR-10 with ResNet-110, E²-Train saves over 83% energy with only 0.56% accuracy loss. When saving over 91% (i.e., more than 10 $\times$), the accuracy drop is still less than 2%. On CIFAR-100 with ResNet-110, E²-Train can even surpass baseline on both top-1 and top5 accuracy while saving over 84% energy. More notably, E²-Train is also effective for even compact networks: it saves about 90% energy cost while achieving a comparable accuracy, when adopted for training MobileNetV2.

Table 4: Experiment results with ResNet-110 and MobileNetV2 on CIFAR-10/CIFAR-100.

Dataset	Method	Backbone	Computational Savings (FLOPs)	Energy Savings (FPGA measured)	Accuracy (top-1)	Accuracy (top-5)
CIFAR-10	SMB (original)	ResNet-110	-	-	93.57%	-
	SD[69]		50%	46.03%	91.51%	-
	SMB (original)	MobileNetV2[70]	-	-	92.47%	-
	E ² -Train (SMD+SLU+PSG)	ResNet-110	80.27%	83.40%	93.01%	-
			85.20%	87.42%	91.74%	-
			90.13%	91.34%	91.68%	-
		MobileNetV2[70]	75.34%	88.73%	92.06%	-
CIFAR-100	SMB (original)	ResNet-110	-	-	71.60%	91.50%
	SD[69]		50%	48.34%	70.40%	92.58%
	SMB (original)	MobileNetV2[70]	-	-	71.91%	-
	E ² -Train (SMD+SLU+PSG)	ResNet-110	80.27%	84.17%	71.63%	91.72%
			85.20%	88.72%	68.61%	89.84%
			90.13%	92.90 %	67.94%	89.06%
		MobileNetV2[70]	75.34%	88.17%	71.61%	-

5 Discussion of Limitations and Future Work

We propose the E²-Train framework to achieve energy-efficient CNN training in resource-constrained settings. Three complementary aspects of efforts to trim down training costs - from data, model and algorithm levels, respectively, are carefully designed, justified, and integrated. Experiments on both simulation and a real FPGA demonstrate the promise of E²-Train. Despite the preliminary success, we are aware of several limitations of E²-Train, which also points us to the future road map. For example, E²-Train is currently designed and evaluated for standard off-line CNN training, with all training data presented in batch, for simplicity. This is not scalable for many real-world IoT scenarios, where new training data arrives sequentially in a stream form, with limited or no data buffer/storage leading to the open challenge of “on-the-fly” CNN training [71]. In this case, while both SLU and PSG are still applicable, SMD needs to be modified, e.g., by one-pass active selection of stream-in data samples. Besides, SLU is not yet straightforward to be extended to plain CNNs without residual connections. We expect finer-grained selective model updates, such as online channel pruning [47], to be useful alternatives here. We also plan to optimize E²-Train for continuous adaptation or lifelong learning.

Acknowledgments

The work is in part supported by the NSF RTML grant (1937592, 1937588). The authors would like to thank all anonymous reviewers for their tremendously useful comments to help improve our work.

References

- [1] Alex Krizhevsky, Ilya Sutskever, and Geoffrey E Hinton. Imagenet classification with deep convolutional neural networks. In F. Pereira, C. J. C. Burges, L. Bottou, and K. Q. Weinberger, editors, *Advances in Neural Information Processing Systems 25*, pages 1097–1105. Curran Associates, Inc., 2012.

- [2] Ross B. Girshick, Jeff Donahue, Trevor Darrell, and Jitendra Malik. Rich feature hierarchies for accurate object detection and semantic segmentation. In Proceedings of the IEEE conference on computer vision and pattern recognition, volume abs/1311.2524, pages 580–587, 2014.
- [3] Yaniv Taigman, Ming Yang, Marc’ Aurelio Ranzato, and Lior Wolf. Deepface: Closing the gap to human-level performance in face verification. In Proceedings of the IEEE conference on computer vision and pattern recognition, pages 1701–1708, 2014.
- [4] Tien-Ju Yang, Yu-Hsin Chen, and Vivienne Sze. Designing energy-efficient convolutional neural networks using energy-aware pruning. 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Jul 2017.
- [5] Lukasz Kaiser, Aidan N. Gomez, Noam Shazeer, Ashish Vaswani, Niki Parmar, Llion Jones, and Jakob Uszkoreit. One model to learn them all. CoRR, abs/1706.05137, 2017.
- [6] Y. H. Chen, T. Krishna, J. S. Emer, and V. Sze. Eyeriss: An energy-efficient reconfigurable accelerator for deep convolutional neural networks. IEEE Journal of Solid-State Circuits, 52(1):127–138, Jan 2017.
- [7] Y. Lin, S. Zhang, and N. R. Shanbhag. Variation-tolerant architectures for convolutional neural networks in the near threshold voltage regime. In 2016 IEEE International Workshop on Signal Processing Systems (SiPS), pages 17–22, Oct 2016.
- [8] D. Bankman, L. Yang, B. Moons, M. Verhelst, and B. Murmann. An always-on 3.8 μ j/86processor with all memory on chip in 28-nm cmos. IEEE Journal of Solid-State Circuits, 54(1):158–172, Jan 2019.
- [9] Priya Goyal, Piotr Dollár, Ross B. Girshick, Pieter Noordhuis, Lukasz Wesolowski, Aapo Kyrola, Andrew Tulloch, Yangqing Jia, and Kaiming He. Accurate, large minibatch SGD: training imagenet in 1 hour. CoRR, abs/1706.02677, 2017.
- [10] Minsik Cho, Ulrich Finkler, Sameer Kumar, David S. Kung, Vaibhav Saxena, and Dheeraj Sreedhar. Powerai ddl. CoRR, abs/1708.02188, 2017.
- [11] Yang You, Zhao Zhang, Cho-Jui Hsieh, James Demmel, and Kurt Keutzer. Imagenet training in minutes. Proceedings of the 47th International Conference on Parallel Processing - ICPP 2018, 2018.
- [12] Takuya Akiba, Shuji Suzuki, and Keisuke Fukuda. Extremely large minibatch sgd: Training resnet-50 on imagenet in 15 minutes. CoRR, abs/1711.04325, 2017.
- [13] Xianyan Jia, Shutao Song, Wei He, Yangzihao Wang, Haidong Rong, Feihu Zhou, Liqiang Xie, Zhenyu Guo, Yuanzhou Yang, Liwei Yu, et al. Highly scalable deep learning training system with mixed-precision: Training imagenet in four minutes. arXiv preprint arXiv:1807.11205, 2018.
- [14] Suyog Gupta, Ankur Agrawal, Kailash Gopalakrishnan, and Pritish Narayanan. Deep learning with limited numerical precision. In International Conference on Machine Learning, pages 1737–1746, 2015.
- [15] Naigang Wang, Jungwook Choi, Daniel Brand, Chia-Yu Chen, and Kailash Gopalakrishnan. Training deep neural networks with 8-bit floating point numbers. In Advances in neural information processing systems, pages 7675–7684, 2018.
- [16] Ron Banner, Itay Hubara, Elad Hoffer, and Daniel Soudry. Scalable methods for 8-bit training of neural networks. In Advances in Neural Information Processing Systems, pages 5145–5153, 2018.
- [17] Ting-Wu Chin, Ruizhou Ding, and Diana Marculescu. Adascale: Towards real-time video object detection using adaptive scaling. arXiv preprint arXiv:1902.02910, 2019.
- [18] Shuang Wu, Guoqi Li, Lei Deng, Liu Liu, Dong Wu, Yuan Xie, and Luping Shi. L1-norm batch normalization for efficient training of deep neural networks. IEEE transactions on neural networks and learning systems, 2018.
- [19] Elad Hoffer, Ron Banner, Itay Golan, and Daniel Soudry. Norm matters: efficient and accurate normalization schemes in deep networks. In Advances in Neural Information Processing Systems, pages 2160–2170, 2018.
- [20] Xin Wang, Fisher Yu, Zi-Yi Dou, Trevor Darrell, and Joseph E Gonzalez. Skipnet: Learning dynamic routing in convolutional networks. In Proceedings of the European Conference on Computer Vision (ECCV), pages 409–424, 2018.
- [21] Jeremy Bernstein, Yu-Xiang Wang, Kamyar Azizzadenesheli, and Animashree Anandkumar. signSGD: Compressed Optimisation for Non-Convex Problems. In International Conference on Machine Learning (ICML-18), 2018.

- [22] Pratik Chaudhari and Stefano Soatto. Stochastic gradient descent performs variational inference, converges to limit cycles for deep networks. In 2018 Information Theory and Applications Workshop (ITA), pages 1–10. IEEE, 2018.
- [23] Angelos Katharopoulos and François Fleuret. Not all samples are created equal: Deep learning with importance sampling, 2018.
- [24] Chiyuan Zhang, Samy Bengio, and Yoram Singer. Are all layers created equal? arXiv preprint arXiv:1902.01996, 2019.
- [25] Jonathan Frankle and Michael Carbin. The lottery ticket hypothesis: Finding sparse, trainable neural networks. In ICLR 2019, 2019.
- [26] Frank Seide, Hao Fu, Jasha Droppo, Gang Li, and Dong Yu. 1-bit stochastic gradient descent and its application to data-parallel distributed training of speech dnns. In Fifteenth Annual Conference of the International Speech Communication Association, 2014.
- [27] Dan Alistarh, Jerry Li, Ryota Tomioka, and Milan Vojnovic. Qsgd: Randomized quantization for communication-optimal stochastic gradient descent. arXiv preprint arXiv:1610.02132, 2016.
- [28] Hantian Zhang, Jerry Li, Kaan Kara, Dan Alistarh, Ji Liu, and Ce Zhang. Zipml: Training linear models with end-to-end low precision, and a little bit of deep learning. In Proceedings of the 34th International Conference on Machine Learning-Volume 70, pages 4035–4043. JMLR. org, 2017.
- [29] Christopher De Sa, Matthew Feldman, Christopher Ré, and Kunle Olukotun. Understanding and optimizing asynchronous low-precision stochastic gradient descent. In ACM SIGARCH Computer Architecture News, volume 45, pages 561–574. ACM, 2017.
- [30] Wei Wen, Cong Xu, Feng Yan, Chunpeng Wu, Yandan Wang, Yiran Chen, and Hai Li. Terngrad: Ternary gradients to reduce communication in distributed deep learning. In I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, Advances in Neural Information Processing Systems 30, pages 1509–1519. Curran Associates, Inc., 2017.
- [31] Alham Fikri Aji and Kenneth Heafield. Sparse communication for distributed gradient descent. arXiv preprint arXiv:1704.05021, 2017.
- [32] Yujun Lin, Song Han, Huizi Mao, Yu Wang, and William J Dally. Deep gradient compression: Reducing the communication bandwidth for distributed training. arXiv preprint arXiv:1712.01887, 2017.
- [33] Jianqiao Wangni, Jiale Wang, Ji Liu, and Tong Zhang. Gradient sparsification for communication-efficient distributed optimization. In Advances in Neural Information Processing Systems, pages 1299–1309, 2018.
- [34] Song Han, Huizi Mao, and William J Dally. Deep compression: Compressing deep neural networks with pruning, trained quantization and huffman coding. In International Conference on Learning Representations, 2016.
- [35] Yihui He, Xiangyu Zhang, and Jian Sun. Channel pruning for accelerating very deep neural networks. In 2017 IEEE International Conference on Computer Vision (ICCV), pages 1398–1406. IEEE, 2017.
- [36] Jiecao Yu, Andrew Lukefahr, David Palframan, Ganesh Dasika, Reetuparna Das, and Scott Mahlke. Scalpel: Customizing dnn pruning to the underlying hardware parallelism. In ACM SIGARCH Computer Architecture News, volume 45, pages 548–560. ACM, 2017.
- [37] Junru Wu, Yue Wang, Zhenyu Wu, Zhangyang Wang, Ashok Veeraraghavan, and Yingyan Lin. Deep k-means: Re-training and parameter sharing with harder cluster assignments for compressing deep convolutions. In Jennifer Dy and Andreas Krause, editors, Proceedings of the 35th International Conference on Machine Learning, volume 80 of Proceedings of Machine Learning Research, pages 5363–5372, Stockholmsmässan, Stockholm Sweden, 10–15 Jul 2018. PMLR.
- [38] Wuyang Chen, Ziyu Jiang, Zhangyang Wang, Kexin Cui, and Xiaoning Qian. Collaborative global-local networks for memory-efficient segmentation of ultra-high resolution images. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, pages 8924–8933, 2019.
- [39] Zuxuan Wu, Tushar Nagarajan, Abhishek Kumar, Steven Rennie, Larry S. Davis, Kristen Grauman, and Rogerio Feris. Blockdrop: Dynamic inference paths in residual networks. 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition, Jun 2018.
- [40] Andreas Veit and Serge Belongie. Convolutional networks with adaptive inference graphs. Lecture Notes in Computer Science, page 3–18, 2018.

- [41] Ji Lin, Yongming Rao, Jiwen Lu, and Jie Zhou. Runtime neural pruning. In I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, Advances in Neural Information Processing Systems 30, pages 2181–2191. Curran Associates, Inc., 2017.
- [42] Zhourong Chen, Yang Li, Samy Bengio, and Si Si. Gaternet: Dynamic filter selection in convolutional neural network via a dedicated global gating network, 2018.
- [43] Jianghao Shen, Yue Wang, Pengfei Xu, Yonggan Fu, Zhangyang Wang, and Yingyan Lin. Fractional skipping: Toward finer-grained dynamic cnn inference. In Thirty-Fourth AAAI Conference on Artificial Intelligence (accepted), 2020.
- [44] Yue Wang, Tan Nguyen, Yang Zhao, Zhangyang Wang, Yingyan Lin, and Richard Baraniuk. Energynet: Energy-efficient dynamic inference. NeurIPS workshop, 2018.
- [45] Yue Wang, Jianghao Shen, Ting-Kuei Hu, Pengfei Xu, Tan Nguyen, Richard Baraniuk, Zhangyang Wang, and Yingyan Lin. Dual dynamic inference: Enabling more efficient, adaptive and controllable deep inference. arXiv preprint arXiv:1907.04523, 2019.
- [46] Yingyan Lin, Charbel Sakr, Yongjune Kim, and Naresh Shanbhag. Predictivenet: An energy-efficient convolutional neural network via zero prediction. In Proceedings of ISCAS, 2017.
- [47] Sangkug Lym, Esha Choukse, Siavash Zangeneh, Wei Wen, Sujay Sanghavi, and Mattan Erez. Prunetrain: Fast neural network training by dynamic sparse model reconfiguration, 2019.
- [48] Dimitri P Bertsekas. Incremental gradient, subgradient, and proximal methods for convex optimization: A survey. Optimization for Machine Learning, 2010(1-38):3, 2011.
- [49] Ohad Shamir. Without-replacement sampling for stochastic gradient methods. In Advances in neural information processing systems, pages 46–54, 2016.
- [50] Yoshua Bengio. Practical recommendations for gradient-based training of deep architectures. In Neural networks: Tricks of the trade, pages 437–478. Springer, 2012.
- [51] Benjamin Recht and Christopher Re. Beneath the valley of the noncommutative arithmetic-geometric mean inequality: conjectures, case-studies. Technical report, and consequences. Technical report, University of Wisconsin-Madison, 2012.
- [52] Mert Gürbüzbalaban, Asu Ozdaglar, and Pablo Parrilo. Why random reshuffling beats stochastic gradient descent. arXiv preprint arXiv:1510.08560, 2015.
- [53] Rong Ge, Furong Huang, Chi Jin, and Yang Yuan. Escaping from saddle points—online stochastic gradient for tensor decomposition. In Conference on Learning Theory, pages 797–842, 2015.
- [54] Nitish Shirish Keskar, Dheevatsa Mudigere, Jorge Nocedal, Mikhail Smelyanskiy, and Ping Tak Peter Tang. On large-batch training for deep learning: Generalization gap and sharp minima. arXiv preprint arXiv:1609.04836, 2016.
- [55] Hadi Daneshmand, Jonas Kohler, Aurelien Lucchi, and Thomas Hofmann. Escaping saddles with stochastic gradients. arXiv preprint arXiv:1803.05999, 2018.
- [56] Tyler B Johnson and Carlos Guestrin. Training deep models faster with robust, approximate importance sampling. In Advances in Neural Information Processing Systems, pages 7265–7275, 2018.
- [57] Cody Coleman, Stephen Mussmann, Baharan Mirzasoleiman, Peter Bailis, Percy Liang, Jure Leskovec, and Matei Zaharia. Select via proxy: Efficient data selection for training deep networks, 2019.
- [58] Chiyuan Zhang, Samy Bengio, and Yoram Singer. Are all layers created equal? CoRR, abs/1902.01996, 2019.
- [59] Andreas Veit, Michael Wilber, and Serge Belongie. Residual networks behave like ensembles of relatively shallow networks, 2016.
- [60] Klaus Greff, Rupesh K Srivastava, and Jürgen Schmidhuber. Highway and residual networks learn unrolled iterative estimation. ICLR, 2017.
- [61] Mark Horowitz. 1.1 computing’s energy problem (and what we can do about it). In 2014 IEEE international solid-state circuits conference digest of technical papers (ISSCC), pages 10–14. IEEE, 2014.

- [62] Charbel Sakr, Yongjune Kim, and Naresh Shanbhag. Analytical guarantees on numerical precision of deep neural networks. In Doina Precup and Yee Whye Teh, editors, Proceedings of the 34th International Conference on Machine Learning, volume 70 of Proceedings of Machine Learning Research, pages 3007–3016, International Convention Centre, Sydney, Australia, 06–11 Aug 2017. PMLR.
- [63] Alex Krizhevsky and Geoffrey Hinton. Learning multiple layers of features from tiny images. Technical report, Citeseer, 2009.
- [64] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In Proceedings of the IEEE conference on computer vision and pattern recognition, pages 770–778, 2016.
- [65] Mark Sandler, Andrew Howard, Menglong Zhu, Andrey Zhmoginov, and Liang-Chieh Chen. Mobilenetv2: Inverted residuals and linear bottlenecks, 2018.
- [66] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In Proceedings of the IEEE international conference on computer vision, pages 1026–1034, 2015.
- [67] Guandao Yang, Tianyi Zhang, Polina Kirichenko, Junwen Bai, Andrew Gordon Wilson, and Christopher De Sa. Swalp: Stochastic weight averaging in low-precision training. arXiv preprint arXiv:1904.11943, 2019.
- [68] Xilinx Inc. Digilent ZedBoard Zynq-7000 ARM/FPGA SoC Development Board. <https://www.xilinx.com/products/boards-and-kits/1-elhbt.html/>, 2019. [Online; accessed 20-May-2019].
- [69] Gao Huang, Yu Sun, Zhuang Liu, Daniel Sedra, and Kilian Q Weinberger. Deep networks with stochastic depth. In European conference on computer vision, pages 646–661. Springer, 2016.
- [70] Mark Sandler, Andrew Howard, Menglong Zhu, Andrey Zhmoginov, and Liang-Chieh Chen. Mobilenetv2: Inverted residuals and linear bottlenecks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, pages 4510–4520, 2018.
- [71] Doyen Sahoo, Quang Pham, Jing Lu, and Steven CH Hoi. Online deep learning: Learning deep neural networks on the fly. arXiv preprint arXiv:1711.03705, 2017.