



Recurrence of optimum for training weight and activation quantized networks



Ziang Long^{a,*}, Penghang Yin^b, Jack Xin^a

^a University of California, Irvine, United States of America

^b University at Albany, SUNY, United States of America

ARTICLE INFO

Article history:

Received 19 January 2022

Received in revised form 15 July 2022

Accepted 26 July 2022

Available online 4 August 2022

Communicated by Radu Balan

Keywords:

Quantization

Neural networks

Discrete optimization

Straight-through estimator

ABSTRACT

Deep neural networks (DNNs) are quantized for efficient inference on resource-constrained platforms. However, training deep learning models with low-precision weights and activations involves a demanding optimization task, which calls for minimizing a stage-wise loss function subject to a discrete set-constraint. While numerous training methods have been proposed, existing studies for full quantization of DNNs are mostly empirical. From a theoretical point of view, we study practical techniques for overcoming the combinatorial nature of network quantization. Specifically, we investigate a simple yet powerful projected gradient-like algorithm for quantizing two-layer convolutional networks, by repeatedly moving one step at float weights in the negative direction of a heuristic *fake* gradient of the loss function (so-called coarse gradient) evaluated at quantized weights. For the first time, we prove that under mild conditions, the sequence of quantized weights recurrently visit the global optimum of the discrete minimization problem for training a fully quantized network. We also show numerical evidence of the recurrence phenomenon of weight evolution in training quantized deep networks.

© 2022 The Author(s). Published by Elsevier Inc. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Deep neural networks (DNNs) have been profoundly transforming machine learning, in applications of computer vision, reinforcement learning, and natural language processing, and so on. While achieving human level or even super-human performances, DNNs typically have tremendous number of weights with high resource consumption at inference time, which poses a challenge for their deployment on mobile devices used in our daily lives. To address this challenge, research efforts have been made to the quantizing weights and activations of DNNs while maintaining their superior performance. Quantization methods train DNNs with the weights and activation values being constrained to low-precision arithmetic rather than the conventional floating-point representation in full-precision. [11,27,3,26,17,28], which offer the feasibility of

* Corresponding author.

E-mail address: zlong6@uci.edu (Z. Long).

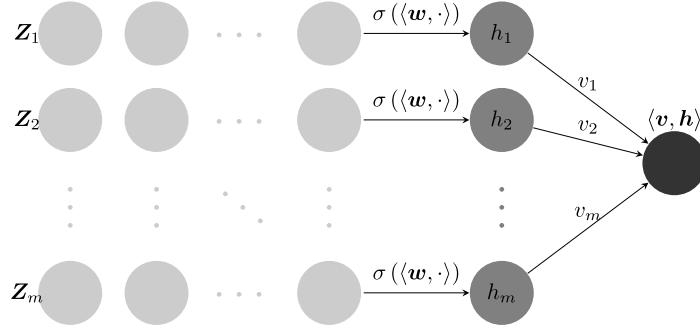


Fig. 1. One-hidden-layer neural network. The first linear layer resembles a convolutional layer with each $\mathbf{Z}_i$ being a patch of size n and $\mathbf{w}$ being the shared weights or filter. The second linear layer serves as the classifier.

running DNNs on edge devices with limited memory storage and battery power. For example, DNNs for ImageNet recognition [6] typically requires hundreds of megabytes storage and billions of floating-point operations at inference time. In contrast, the XNOR-Net [18] with binary weights and activations can achieve $58\times$ faster convolutional operations and $32\times$ memory savings, compared with the float counterpart.

Mathematically, training a fully quantized DNN requires solving a challenging optimization problem with a piecewise constant (and non-convex) training loss function and a discrete set-constraint. That is, one considers the following constrained population risk minimization problem:

$$\min_{\mathbf{w}} f(\mathbf{w}) := \mathbb{E}_{\mathbf{x} \sim p(\mathbf{x})} [\ell(\mathbf{w}; \mathbf{x})] \quad \text{subject to} \quad \mathbf{w} \in \mathcal{Q}$$

where $p(\mathbf{x})$ is the probability distribution of data; $\ell(\mathbf{w}; \mathbf{x})$ is the sample loss function for input $\mathbf{x}$; $\mathcal{Q}$ abstractly denotes the discrete set of quantized weights.

1.1. Problem setup

In this paper, we consider the training of a one-hidden-layer model that outputs the prediction for any input sample $\mathbf{Z} \in \mathbb{R}^{m \times n}$:

$$y(\mathbf{Z}; \mathbf{w}) := \sum_{i=1}^m v_i \sigma(\mathbf{Z}_i^\top \mathbf{w}) = \mathbf{v}^\top \sigma(\mathbf{Z} \mathbf{w}) \quad (1)$$

where $\mathbf{Z}_i^\top$ denotes the i -th row vector of $\mathbf{Z}$; $\mathbf{w} \in \mathbb{R}^n$ is the trainable weights in the first linear layer, and $\mathbf{v} \in \mathbb{R}^m$ the weights in the second linear layer which are assumed to be known and fixed during the training process; the activation function $\sigma(x) = \mathbb{1}_{\{x > 0\}}$ is *binary*, acting component-wise on the vector $\mathbf{Z} \mathbf{w}$. The label is generated according to $y_{\mathbf{Z}}^* := y(\mathbf{Z}; \mathbf{w}^*)$ for some unknown true parameters $\mathbf{w}^* \in \mathbb{R}^n$ (see Fig. 1).

We fit the described model with quantized weights $\mathbf{w} \in \mathcal{Q}$ and binary activation function $\sigma(x) = \mathbb{1}_{\{x > 0\}}$ on the i.i.d. Gaussian data $\{(\mathbf{Z}, y_{\mathbf{Z}}^*)\}_{\mathbf{Z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})}$. In this paper, we will focus on the cases of binary and ternary weights. In the binary case, every quantized weight in $\mathbf{w}$ is either α or $-\alpha$ for some universal real-valued constant $\alpha > 0$, or equivalently, $\mathcal{Q} = \mathbb{R}_+ \times \{\pm 1\}^n$; this setup of binary weights is widely adopted in the literature; for example, [18]. Similarly in the ternary case, we take $\mathcal{Q} = \mathbb{R}_+ \times \{0, \pm 1\}^n$; see [14, 25] for examples.

Furthermore, we use the squared loss to measure the discrepancy between the model output $y(\mathbf{Z}; \mathbf{w})$ and the true label $y_{\mathbf{Z}}^*$:

$$\begin{aligned}\ell(\mathbf{w}; \mathbf{Z}) &:= \frac{1}{2} (y(\mathbf{Z}; \mathbf{w}) - y_{\mathbf{Z}}^*)^2 \\ &= \frac{1}{2} (\mathbf{v}^\top \sigma(\mathbf{Z}\mathbf{w}) - \mathbf{v}^\top \sigma(\mathbf{Z}\mathbf{w}^*))^2,\end{aligned}\tag{2}$$

and cast the learning task as the following population loss minimization problem:

$$\min_{\mathbf{w} \in \mathbb{R}^n} f(\mathbf{w}) := \mathbb{E}_{\mathbf{Z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})} [\ell(\mathbf{w}; \mathbf{Z})] \quad \text{subject to} \quad \mathbf{w} \in \mathcal{Q} \tag{3}$$

where the sample loss function $\ell(\mathbf{w}; \mathbf{Z})$ is given by (2). Hereby it is not hard to show that the gradient of $\ell(\mathbf{w}; \mathbf{Z})$ w.r.t. $\mathbf{w}$ is

$$\nabla_{\mathbf{w}} \ell(\mathbf{w}; \mathbf{Z}) = \mathbf{Z}^\top (\sigma'(\mathbf{Z}\mathbf{w}) \odot \mathbf{v}) (y(\mathbf{Z}; \mathbf{w}) - y_{\mathbf{Z}}^*). \tag{4}$$

Note that σ' is zero a.e., which makes $\nabla_{\mathbf{w}} \ell(\mathbf{w}; \mathbf{Z})$ zero a.e. and thus inapplicable to the training.

In this setting, we study the following iterative algorithm for training fully quantized networks

$$\begin{cases} \mathbf{y}^{t+1} = \mathbf{y}^t - \eta_t \mathbb{E}[\tilde{\nabla}_{\mathbf{w}} \ell(\mathbf{w}^t; \mathbf{x})] \\ \mathbf{w}^{t+1} = \text{proj}_{\mathcal{Q}}(\mathbf{y}^{t+1}), \end{cases} \tag{QUANT}$$

where $\tilde{\nabla}_{\mathbf{w}} \ell$ denotes some heuristic modification of the vanished gradient $\nabla_{\mathbf{w}} \ell$ in (4) based on the so-called straight-through estimator (STE) [1,9], rendering a valid search direction; see section 2.1 for details. Following [24], we shall refer to this fake ‘gradient’ induced by STE as coarse gradient throughout this paper.

1.2. Related works

For the best possible performance under quantization, the pre-trained full-precision networks need to be re-trained. In the regime of weight quantization, the BinaryConnect scheme:

$$\begin{cases} \mathbf{y}^{t+1} = \mathbf{y}^t - \eta_t \mathbb{E}[\nabla_{\mathbf{w}} \ell(\mathbf{w}^t; \mathbf{x})] \\ \mathbf{w}^{t+1} = \text{proj}_{\mathcal{Q}}(\mathbf{y}^{t+1}) \end{cases} \tag{5}$$

was first proposed in [5] for training DNNs with binary (1-bit) weights. It is similar to QUANT, but simply uses the standard gradient $\nabla_{\mathbf{w}} \ell$ as the activation values were not quantized. The method was then extended to multi-bit weight quantization such as ternary weight networks [14]. On the theoretical side, [15] analyzed the convergence of BinaryConnect scheme for weight quantization, and proved that $\{\mathbf{w}^t\}$ converge to an error floor region of the optimal quantized weights under strong convexity and smoothness assumptions on f . Recently, [16] used an algorithm called “error feedback” for pruning networks [7,22]. It is basically the same as BinaryConnect, except that the weight quantization step $\text{proj}_{\mathcal{Q}}$ is replaced with weight pruning/thresholding which can also be viewed as a projection. The authors showed the convergence to a neighborhood of optimal solution under strong convexity and smoothness assumptions whose radius is $O(\sqrt{d})$ with d being the number of model parameters. Moreover, it remains unclear whether the global optimum can actually be reached in this setting.

The idea of STE has been extensively used for efficiently handling discrete-valued functions arising in machine learning problems. A STE, used in the backward pass only, is a heuristic proxy that substitutes the a.e. zero derivative of discrete component composited in the loss function when computing the gradient under chain rule. Its applications include, but are not limited to, network quantization [10,3,27,4,11,21,2], neural architecture search [20], knowledge graphs [23], discrete latent representations [12]. For networks with binary

activations (and real-valued weights), [24] showed that STE-based gradient (called coarse gradient) methods converge only when a proper STE like ReLU STE [3] is used. And they proved that the negative of the resulting coarse gradient points to a descent direction that reduces the training loss. For quantization of both weights and activations, [10,11,3,4,27] utilized QUANT scheme which is the combination of BinaryConnect and STE, and achieved state-of-the-art classification accuracies. Yet to our best knowledge, no convergence results of QUANT have been established to date.

1.3. Main contributions

In this paper, we examine the quantization of one-hidden-layer networks with binary activation and binary or ternary weights using the QUANT algorithm. Surprisingly, the sequence of quantized weights $\{\mathbf{w}^t\}$ generated by QUANT is generically divergent. Our key contributions are the *first* theoretical results on the dynamics of QUANT algorithm for learning fully quantized neural nets: (1) we prove the generic divergence if the teacher parameters are not in a quantized state, and give an explicit example of oscillatory divergence behavior (the sequence $\{\mathbf{w}^t\}$ has period 3 and jumps between sub-optimal quantized states; see Example 1). (2) We explicitly point out, in the ternary case, the n (out of $3^n - 1$) sub-optimal quantized states that $\{\mathbf{w}^t\}$ could visit infinite many times; see Remark 1 and Lemma 8. (3) We prove that $\{\mathbf{w}^t\}$ oscillates around the global optimum of quantization problem. Under conditions that teacher parameters and their quantized values are close enough (see Theorem 1), $\{\mathbf{w}^t\}$ visits the quantized teacher parameters (the optimum) infinitely often (*recurrence*). Compared with theoretical results for BinaryConnect [15,16], our analysis is more precise and in depth in order to overcome a biased gradient modification in QUANT based on straight-through estimator (STE) [9,1]. Our result is stronger in that the recurrence behavior at global minimum holds *without global convexity assumption of the loss function*.

Organization. In section 2, we introduce the concept of coarse gradient and present some useful preliminary results about the QUANT algorithm. In section 3, we summarize the main results regarding the recurrence behavior of QUANT algorithm. More technical details and sketch of proofs are presented in section 4.

2. Preliminaries

We investigate the convergence behavior of QUANT described in Algorithm 1 for solving the quantization problem (3). In Algorithm 1, $\tilde{\nabla} f := \mathbb{E}[\tilde{\nabla}_{\mathbf{w}} \ell(\mathbf{w}^t; \mathbf{x})]$ stands for coarse gradient [24] in expectation specified in section 2.1 below, which side-steps the vanished gradient issue. Since the loss function $\ell(\mathbf{w}; \mathbf{Z})$ defined in (2) is scale-invariant, i.e., $\ell(\mathbf{Z}; \mathbf{w}) = \ell(\mathbf{Z}; \mathbf{w}/c)$ for any scalar $c > 0$, without loss of generality, we assume that $\|\mathbf{w}^*\| = 1$ is unit-normed.

Algorithm 1: QUANT algorithm for solving (3).

```

Input: number of iterations  $T$ , learning rate  $\eta_t$ , weight bits  $b$ ;
Initialize: auxiliary real-valued weights  $\mathbf{y}^0 \in \mathbb{R}^n$ , iteration number  $t = 1$ ;
while  $t \leq T$  do
     $\mathbf{y}^t = \mathbf{y}^{t-1} - \eta_t \tilde{\nabla} f(\mathbf{w}^{t-1})$ ;
     $\mathbf{w}^t = \text{proj}_{\mathcal{Q}}(\mathbf{y}^t)$ ;
     $t = t + 1$ ;
end

```

In addition, throughout this paper we make the following assumptions on the learning rate $\eta_t > 0$:

1. $\sum_{t=1}^{\infty} \eta_t = \infty$.
2. η_t is upper bounded by some positive constant η .

2.1. Coarse gradient

In this part, we specify the coarse gradient $\tilde{\nabla}f(\mathbf{w})$ in Algorithm 1. As shown in (4), the standard gradient of $\ell(\mathbf{w}; \mathbf{Z})$ w.r.t. $\mathbf{w}$ is given by

$$\nabla_{\mathbf{w}}\ell(\mathbf{w}; \mathbf{Z}) = \mathbf{Z}^\top (\sigma'(\mathbf{Z}\mathbf{w}) \odot \mathbf{v}) (y(\mathbf{Z}; \mathbf{w}) - y_{\mathbf{Z}}^*).$$

The associated coarse gradient w.r.t. $\mathbf{w}$ associated with the sample $(\mathbf{Z}, y_{\mathbf{Z}}^*)$ is given by replacing σ' with a surrogate derivative, known as straight-through estimator (STE) [1,24]. Here we consider the derivative of ReLU function $\mu(x) = \max\{x, 0\}$ which is a widely used STE for quantization, namely, we modify the original gradient $\nabla_{\mathbf{w}}\ell(\mathbf{w}; \mathbf{Z})$ as follows:

$$\tilde{\nabla}_{\mathbf{w}}\ell(\mathbf{w}; \mathbf{Z}) = \mathbf{Z}^\top (\mu'(\mathbf{Z}\mathbf{w}) \odot \mathbf{v}) (y(\mathbf{Z}; \mathbf{w}) - y_{\mathbf{Z}}^*).$$

The coarse gradient induced by ReLU STE μ' is just the expectation of $\tilde{\nabla}_{\mathbf{w}}\ell(\mathbf{w}; \mathbf{Z})$ over $\mathbf{Z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. We evaluate the coarse gradient $\tilde{\nabla}f(\mathbf{w})$ used in Algorithm 1:

Lemma 1. *The expected coarse gradient of $\ell(\mathbf{w}; \mathbf{Z})$ w.r.t. $\mathbf{w}$ is*

$$\begin{aligned} \tilde{\nabla}f(\mathbf{w}) &:= \mathbb{E}_{\mathbf{Z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})} [\tilde{\nabla}_{\mathbf{w}}\ell(\mathbf{w}; \mathbf{Z})] \\ &= \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left(\frac{\mathbf{w}}{\|\mathbf{w}\|} - \mathbf{w}^* \right). \end{aligned} \quad (6)$$

2.2. Characterization of optimal solutions

To study the convergence of Algorithm 1, we first obtain the closed-form expression of the objective function in (2), which only depends on the angle between quantized weight vector $\mathbf{w}$ and the true weight vector $\mathbf{w}^*$. This helps us find the expression of global minimum to (1).

Lemma 2. *Let $\mathbf{w} \neq \mathbf{0}$ be nonzero vector.*

- the training loss in (3) is given by

$$f(\mathbf{w}) = \frac{\|\mathbf{v}\|^2}{2\pi} \arccos \left(\frac{\mathbf{w}^\top \mathbf{w}^*}{\|\mathbf{w}\|} \right)$$

- For any $\delta > 0$, $\mathbf{w} = \delta \cdot \text{proj}_{\mathcal{Q}}(\mathbf{w}^*)$ is a global optimum of quantization problem (3).

The above result can be easily derived from Lemma 1 of [24], so we omit the proof. Lemma 2 states that the optimal quantized weights are just the projection of $\mathbf{w}^*$ onto $\mathcal{Q}$, i.e., the direct quantization of teacher parameters $\mathbf{w}^*$. Note that the projection/quantization may not be unique, we refer to $\text{proj}_{\mathcal{Q}}(\mathbf{y})$ as any choice of the projection of $\mathbf{y}$ onto $\mathcal{Q}$.

2.3. Weight quantization step

The following two lemmas give the closed-form formulas of the projection/quantization $\text{proj}_{\mathcal{Q}}(\cdot)$ in Algorithm 1 in the binary and ternary cases, respectively.

Lemma 3 (Binary Case). For any non-zero $\mathbf{y} \in \mathbb{R}^n$, the projection of $\mathbf{y}$ onto $\mathcal{Q} = \mathbb{R}_+ \times \{\pm 1\}^n$ is

$$\text{proj}_{\mathcal{Q}}(\mathbf{y}) = \frac{\|\mathbf{y}\|_1}{n} \widetilde{\text{sign}}(\mathbf{y}),$$

where the sign function acts element-wise

$$\widetilde{\text{sign}}(\mathbf{y})_i = \begin{cases} 1 & \text{if } y_i \geq 0 \\ -1 & \text{if } y_i < 0. \end{cases}$$

The above lemma is due to [18]. In the ternary case, [25] gives the following result:

Lemma 4 (Ternary Case). For any non-zero $\mathbf{y} \in \mathbb{R}^n$, the projection of $\mathbf{y}$ on $\mathcal{Q} = \mathbb{R}_+ \times \{0, \pm 1\}^n$ is

$$\text{proj}_{\mathcal{Q}}(\mathbf{y}) = \frac{\|\mathbf{y}_{[j^*]}\|_1}{j^*} \text{sign}(\mathbf{y}_{[j^*]})$$

where $j^* = \arg \max_{1 \leq j \leq n} \frac{\|\mathbf{y}_{[j]}\|_1^2}{j}$, and $\mathbf{y}_{[j]} \in \mathbb{R}^n$ extracts the first j largest entries in magnitude of $\mathbf{y}$ and enforces 0 elsewhere. Here,

$$\text{sign}(\mathbf{y})_i = \begin{cases} 1 & \text{if } y_i > 0 \\ 0 & \text{if } y_i = 0 \\ -1 & \text{if } y_i < 0. \end{cases}$$

3. Main results

By Lemma 2, we assume, for the ease of presentation, that the iterates $\{\mathbf{w}^t\}$ are normalized, that is, we re-define $\mathbf{w}^t$ in Algorithm 1 by

$$\mathbf{w}^t = \widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{y}^t) := \frac{\text{proj}_{\mathcal{Q}}(\mathbf{y}^t)}{\|\text{proj}_{\mathcal{Q}}(\mathbf{y}^t)\|}$$

Our results extend trivially to the original QUANT without normalization as the value of $f(\mathbf{w})$ does not depend on $\|\mathbf{w}\|$. Furthermore, we denote by $\widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{w}^*)$ the normalization of the quantization/projection of $\mathbf{w}^*$, $\text{proj}_{\mathcal{Q}}(\mathbf{w}^*)$, which is a global minimum according to Lemma 2. Our main results show that the optimum $\widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{w}^*)$ is recurrent as long as $\mathbf{w}^*$ is close to its normalized quantization.

Theorem 1. Consider the setup of quantization problem (3). Let $\mathcal{Q}$ be either $\mathbb{R}_+ \times \{\pm 1\}^n$ (binary case) or $\mathbb{R}_+ \times \{0, \pm 1\}^n$ (ternary case). There exists constant $\epsilon > 0$ that depends on the weight bit-width and dimension n only, such that for any $\mathbf{w}^*$ with

$$0 < \|\mathbf{w}^* - \widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{w}^*)\| < \epsilon,$$

we have $\mathbf{w}^t = \widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{w}^*)$ for infinitely many t values, where $\{\mathbf{w}^t\}$ is the sequence generated by Algorithm 1 with any initialization.

Intuitively, ternary weights should work better than binary weights. The following remark confirms this intuition by showing that the number of points where $\mathbf{w}^t$ visits infinitely many times is limited.

Remark 1. In the ternary case, we can further prove that the sequence $\{\mathbf{w}^t\}$ generated by Algorithm 1 has at most n sub-sequential limits.

4. Proof sketch

On one hand, the binary case is rather simple. We show that part of the coordinates is stable while others have oscillating sign. We further prove that the set of oscillating coordinates is not empty as long as $\mathbf{w}^* \notin \mathcal{Q} = \mathbb{R}_+ \times \{\pm 1\}^n$ is not quantized.

On the other hand, the proof of the ternary case follows the following steps. Our first step shows the sequence $\mathbf{y}^t$ generated by Algorithm 1 is bounded away from the origin for all but finitely many t values. Then, our second step shows each coordinate of $\mathbf{y}^t$ is of the same sign of $\mathbf{w}^*$ for all but finitely many t values. This forces $\mathbf{y}^t$ to stay in the same orthant to which $\mathbf{w}^*$ belongs. As a matter of fact, an n -dimensional space has in total 2^n orthants, which means $\mathbf{y}^t$ can only stay in a small region near $\mathbf{w}^*$. After that, our third step furthermore cuts the orthant into $n!$ congruent cones and argues $\mathbf{y}^t$ must stay in the same cone where $\mathbf{w}^*$ is for all but finitely many t values. In the last step, we prove the ternary case of Theorem 1, which asserts that as long as the underlying true parameter $\mathbf{w}^*$ is close to quantized state $\mathcal{Q} = \mathbb{R}_+ \times \{0, \pm 1\}^n$, i.e., any vertex of the cone it belongs to, the optimum is guaranteed to be recurrent.

4.1. Binary weight

In view of Lemmas 2 and 3, we have that the normalized optimum of (3) is $\frac{1}{\sqrt{n}} \widetilde{\text{sign}}(\mathbf{w}^*)$. The Lemma below shows that some coordinates of $\mathbf{w}^t$ generated by Algorithm 1 have oscillating signs.

Proposition 1. Let $\mathbf{w}^t$ be any infinite sequence generated by Algorithm 1. If $|w_j^*| < \frac{1}{\sqrt{n}}$, then there exist infinitely many t_1 and t_2 such that $w_j^{t_1} = \frac{1}{\sqrt{n}}$ and $w_j^{t_2} = -\frac{1}{\sqrt{n}}$.

The above lemma clearly implies that $\mathbf{w}^t$ does not converge, as long as $\mathbf{w}^* \notin \mathcal{Q}$.

Corollary 1. If $\mathbf{w}^* \notin \mathcal{Q}$, then any sequence $\{\mathbf{w}^t\}$ generated by Algorithm 1 does not converge.

Since Algorithm 1 does not have a limit unless the weights in the network are already quantized, we ask a natural question: Can we guarantee the optimum to be visited infinitely many times? The general answer is no. We have the following example *demonstrating that the optimum may never be achieved*. We refer the proof of the following example to the appendix.

Example 1. Let $\mathbf{w}^* = \left(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2}\sqrt{\frac{11}{3}}\right)$ so that the best the optimum $\widetilde{\text{proj}}_{\mathcal{Q}} \mathbf{w}^* = \left(\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2}\right)$. Let $\eta_t = \eta$, $\lambda = \frac{\eta \|\mathbf{v}\|^2}{6\sqrt{2\pi}}$ and

$$\begin{cases} y_1^0 \in (-\lambda, 0) \\ y_2^0 \in (0, \lambda) \\ y_3^0 \in (\lambda, 2\lambda) \\ y_4^0 \in (0, \infty) \end{cases}$$

the sequence $\{\mathbf{w}^t\}$ generated by Algorithm 1 with initialization $\mathbf{y}^0$ satisfies $\mathbf{w}^{t+3} = \mathbf{w}^t$ and $\mathbf{w}^t \neq \widetilde{\text{proj}}_{\mathcal{Q}} \mathbf{w}^*$ for all t .

In the following, we give a sufficient condition for the optimum to be recurrent. The condition requires $\mathbf{w}^*$ to be close to $\mathcal{Q}$. The following result is for the binary case of Theorem 1.

Theorem 1 (Binary Case). *If the optimum $\hat{\mathbf{w}} := \widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{w}^*) = \frac{1}{\sqrt{n}} \widetilde{\text{sign}}(\mathbf{w}^*)$ of (3) satisfies $0 < \sum_{|w_j^*| < \frac{1}{\sqrt{n}}} |w_j^* - \hat{w}_j| < \frac{2}{\sqrt{n}}$ then there exist infinitely many t values for any sequence $\{\mathbf{w}^t\}$ generated by Algorithm 1 such that $\mathbf{w}^t = \widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{w}^*)$.*

4.2. Ternary weights

The first result shows that $\mathbf{w}^t$ generated by Algorithm 1 is generally divergent, and it converges only when the true parameters $\mathbf{w}^* \in \mathcal{Q} = \mathbb{R}_+ \times \{0, \pm 1\}^n$.

Proposition 2 (Ternary Case). *Let $\{\mathbf{w}^t\}$ be any sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q} = \mathbb{R}_+ \times \{0, \pm 1\}^n$, then $\{\mathbf{w}^t\}$ is not a convergent sequence.*

In what follows, we detail the proof of convergence behavior of Algorithm 1.

Our first step is to rule out an exceptional case that the direction of $\mathbf{y}^t$ changes significantly in only one iteration. As shown in Lemma 1, the coarse gradient is bounded by a constant depending only on the fixed weight vector $\mathbf{v}$. So it suffices to show that $\|\mathbf{y}^t\|$ is bounded away from zero for all but finitely many t values.

Lemma 5. *Let $\{\mathbf{y}^t\}$ be any auxiliary sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q}$, then $\|\mathbf{y}^t\|_1$ converges to infinity as t increases.*

Lemma 5 shows that for any positive constant $c > 0$, we have $\|\mathbf{y}^t\|_1 > c$ for all but finitely many t values.

Since Lemma 5 guarantees that the direction of $\mathbf{y}^t$ will not change significantly, we cut down the region that $\mathbf{y}^t$ can belong to in two steps. To describe our first cut down, we need the following definition to make our statement precise.

Definition 1. For any $\mathbf{x} \in \mathbb{R}^n$, we define the orthant of $\mathbf{x}$ as

$$\mathbf{O}(\mathbf{x}) := \{\mathbf{y} \in \mathbb{R}^n : \text{sign}(\mathbf{y}) = \text{sign}(\mathbf{x})\},$$

where $\text{sign}(\cdot)$ acts coordinate-wise. Furthermore, we say $\mathbf{O}(\mathbf{x})$ is regular if any coordinate of $\mathbf{x}$ is not zero.

We state some basic properties of the defined orthant.

Proposition 3. *For any $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, the following statements are true:*

1. *Either $\mathbf{O}(\mathbf{x}) = \mathbf{O}(\mathbf{y})$ or $\mathbf{O}(\mathbf{x}) \cap \mathbf{O}(\mathbf{y}) = \emptyset$.*
2. *$\mathbf{x} \in \mathbf{O}(\mathbf{x})$.*
3. *$\cup_{\mathbf{x} \in \mathbb{R}^n} \mathbf{O}(\mathbf{x}) = \mathbb{R}^n$.*
4. *There are in total 3^n orthants.*
5. *There are in total 2^n regular orthants.*

Lemma 6. *Let $\{\mathbf{y}^t\}$ be any auxiliary sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q}^n$, then any subsequential limit of $\hat{\mathbf{y}}^t := \frac{\mathbf{y}^t}{\|\mathbf{y}^t\|}$ belongs to the closure of $\mathbf{O}(\mathbf{w}^*)$. Furthermore, if $\mathbf{O}(\mathbf{w}^*)$ is regular, then $\mathbf{y}^t$ lies in $\mathbf{O}(\mathbf{w}^*)$ for all but finitely many t values.*

In our previous step, we have partitioned $\mathbb{R}^n$ into orthants and showed that $\mathbf{y}^t$ enter into a small neighborhood of the orthant where $\mathbf{w}^*$ stays. Now, we prove a stronger result based on the conclusion of our previous step. We would like to cut each orthant into several congruent cones which we shall define later and

argue $\mathbf{y}^t$ will move and stay in close neighborhood of the cone where $\mathbf{w}^*$ stays. This step makes a stronger statement because we manage to shrink the size of the region where $\mathbf{y}^t$ can stay.

Definition 2. For any non-zero vector $\mathbf{x} \in \mathbb{R}^n$, we define the cone of $\mathbf{x}$ to be

$$\text{Cone}(\mathbf{x}) := \left\{ \mathbf{y} \in \mathcal{O}(\mathbf{x}) : \right. \\ \left. \text{sign}(|y_j| - |y_i|) = \text{sign}(|x_j| - |x_i|) \text{ for } \forall i, j \in [n] \right\}.$$

Moreover, we say $\text{Cone}(\mathbf{x})$ is regular if $\mathcal{O}(\mathbf{x})$ is regular and any $|x_j| \neq |x_i|$ for all $j \neq i$.

Proposition 4. For any $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, the following statements are true:

1. Either $\text{Cone}(\mathbf{x}) = \text{Cone}(\mathbf{y})$ or $\text{Cone}(\mathbf{x}) \cap \text{Cone}(\mathbf{y}) = \emptyset$.
2. $\mathbf{x} \in \text{Cone}(\mathbf{x})$.
3. If $\mathbf{y} \in \text{Cone}(\mathbf{x})$, then $\text{Cone}(\mathbf{y}) = \text{Cone}(\mathbf{x})$.
4. $\cup_{\mathbf{y} \in \mathcal{O}(\mathbf{x})} \text{Cone}(\mathbf{y}) = \mathcal{O}(\mathbf{x})$.
5. Any regular orthant contains $n!$ regular cones.

Lemma 7. Let $\{\mathbf{y}^t\}$ be any auxiliary real-valued sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q}$, then any sub-sequential limit of $\tilde{\mathbf{y}}^t := \frac{\mathbf{y}^t}{\|\mathbf{y}^t\|}$ belongs to the closure of $\text{Cone}(\mathbf{w}^*)$. Moreover, if $\text{Cone}(\mathbf{w}^*)$ is regular, then $\mathbf{y}^t \in \text{Cone}(\mathbf{w}^*)$ for all but finitely many t values.

The auxiliary weight vector $\mathbf{y}^t$ can only stay in a small region around $\mathbf{w}^*$ for large t values.

Definition 3. For any point $\mathbf{x} \in \mathbb{R}^n$, assume $(j_1, j_2, \dots, j_n)$ is a permutation of $[n]$ such that

$$|x_{j_1}| \geq |x_{j_2}| \geq \dots \geq |x_{j_n}|$$

We define the set of vertexes of $\mathbf{x}$ to be

$$\Lambda(\mathbf{x}) := \left\{ \frac{1}{\sqrt{k}} \sum_{i=1}^k \text{sign}(x_{j_i}) \mathbf{e}_{j_i} : x_{j_{k+1}} \neq x_{j_k} \text{ are nonzeros} \right\}.$$

Below are some basic facts about connection between vertexes and cones.

Proposition 5. For any $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ let $k := |\Lambda(\mathbf{x})|$, the following statements are true:

1. $0 \leq k \leq n$.
2. $\Lambda(\mathbf{x})$ is empty if and only if $\mathbf{x} = \mathbf{0}$.
3. $\Lambda(\mathbf{x})$ is a subset of the boundary of $\text{Cone}(\mathbf{x})$.
4. $\text{Cone}(\mathbf{x}) = \text{Cone}(\mathbf{y})$ if and only if $\Lambda(\mathbf{x}) = \Lambda(\mathbf{y})$.
5. $\widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{x}) \in \Lambda(\mathbf{x})$.
6. $\mathbf{y}$ lies in $\text{Cone}(\mathbf{x})$ if and only if there exists k positive numbers $\{\mu_{\mathbf{z}}(\mathbf{y}) : \mathbf{z} \in \Lambda(\mathbf{x})\}$ such that $\mathbf{y} = \sum_{\mathbf{z} \in \Lambda(\mathbf{x})} \mu_{\mathbf{z}}(\mathbf{y}) \mathbf{z}$.
7. $\mathbf{y}$ lies in the closure of $\text{Cone}(\mathbf{x})$ if and only if there exists k non-negative numbers $\{\mu_{\mathbf{z}}(\mathbf{y}) : \mathbf{z} \in \Lambda(\mathbf{x})\}$ such that $\mathbf{y} = \sum_{\mathbf{z} \in \Lambda(\mathbf{x})} \mu_{\mathbf{z}}(\mathbf{y}) \mathbf{z}$.

Table 1

Validation Accuracy of LeNet-5 on MNIST and ResNet-20/VGG-11 on CIFAR-10.

	float	binary	ternary
LeNet-5	99.37	99.33	99.34
ResNet-20	92.33	89.42	90.86
VGG-11	92.15	89.47	90.91

$$8. \cup_{\mathbf{x} \in \mathbb{R}^n} \Lambda(\mathbf{x}) = \{\mathbf{x} \in \mathcal{Q} : \|\mathbf{x}\| = 1\}.$$

Lemma 8. Let $\{\mathbf{w}^t\}$ be the sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q} = \mathbb{R}_+ \times \{0, \pm 1\}^n$, then $\mathbf{w}^t \in \Lambda(\mathbf{w}^*)$ for all but finitely many t values.

The following result is the ternary case of Theorem 1 stated in section 3.

Theorem 1 (Ternary Case). Let $\{\mathbf{z}_j\}_{j=1}^k = \Lambda(\mathbf{w}^*)$ where $\mathbf{z}_1 = \widetilde{\text{proj}}_{\mathcal{Q}} \mathbf{w}^*$ is the optimum and $\mathbf{w}^* = \sum_{j=1}^k \lambda_j \mathbf{z}_j$. If $0 < \sum_{j=2}^k \lambda_j < 1$, we have $\mathbf{w}^t = \widetilde{\text{proj}}_{\mathcal{Q}} \mathbf{w}^*$ for infinitely many t values, where $\mathbf{w}^t$ is any infinite sequence generated by Algorithm 1 with any initialization.

Intuitively, the parameter λ_j in Theorem 1 stands for the proportion of time that $\{\mathbf{w}^t\}$ stays at $\mathbf{z}_j$. For instance, if $\lambda_j \approx 1$, then most of $\{\mathbf{w}^t\}$ stay at $\mathbf{z}_j$ so that the oscillation has a longer ‘period’ and is harder to observe. On the contrary, if all λ_j ’s are almost the same then $\{\mathbf{w}^t\}$ behaves like uniform distribution and oscillation becomes more obvious. Beside λ_j ’s, a smaller learning rate can render $\mathbf{y}^t$ moves slower which can also slow down the oscillation. Although there are ways to stabilize the training process, both our theorem and the experiments in the next section suggests the oscillation behavior is inevitable.

5. Experiments

In this section, we implement QUANT algorithm on both synthetic data and MNIST/CIFAR image data. Our goals are (1) to validate our theoretical findings and (2) to show the appearance of the oscillation behavior in more complicated setups. **With that said, we emphasize that we did not extensively tune the hyper-parameters or use ad-hoc tricks to achieve the best possible validation accuracy.** More comprehensive experimental results for QUANT-based approaches can be found in, for examples, [3,4,11,27]. Here we report the validation accuracies on MNIST and CIFAR-10 for fully quantized networks in Table 1. For both synthetic and image data sets, we observed the oscillation behavior.

5.1. Synthetic data

We take $m = 4$, $n = 8$ in (1) and construct $\mathbf{v} \sim N(\mathbf{0}, \mathbf{I}_m)$ and $\mathbf{w}^* \sim N(\mathbf{0}, \mathbf{I}_n)$ be random vectors. For each run, we fix $\mathbf{v}$ and $\mathbf{w}^*$ and train the neural network (1) by algorithm (1) for 200 iterations with a learning rate being 0.1. Fig. 2 show the evolution of binary/ternary weight of $\mathbf{w}^t$ in the last 100 iterations. Each block of size 8×100 corresponds to the evolution of $\mathbf{w}^t$ during the 100 iterations. The (quantized) global minimum $\text{proj}_{\mathcal{Q}} \mathbf{w}^*$ for each run is shown on the right side of the corresponding subplot in Fig. 2.

5.2. MNIST

We train LeNet-5 with binary/ternary weights and 4-bit activations using QUANT algorithm. For deep networks, the (quantized) global optimum is generally unknown, we instead show the oscillating behavior around local optimum. Note that Fig. 4 shows the training loss no longer drop significantly during the last

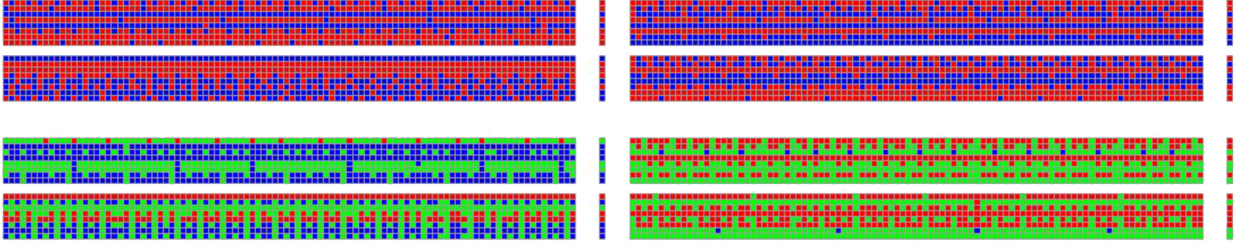


Fig. 2. Evolution of Weight signs of synthetic network described in (1). Each of the 8 large blocks is a colored display of weight sign values via 8×100 matrix (i.e., 8 filter weight signs evolved over the last 100 iterations). The bars to the right of blocks are the corresponding optima. **Top two rows:** Binary weight signs, red /blue for $1/-1$. **Bottom two rows:** Ternary weight signs, red/green/blue for $1/0/-1$.

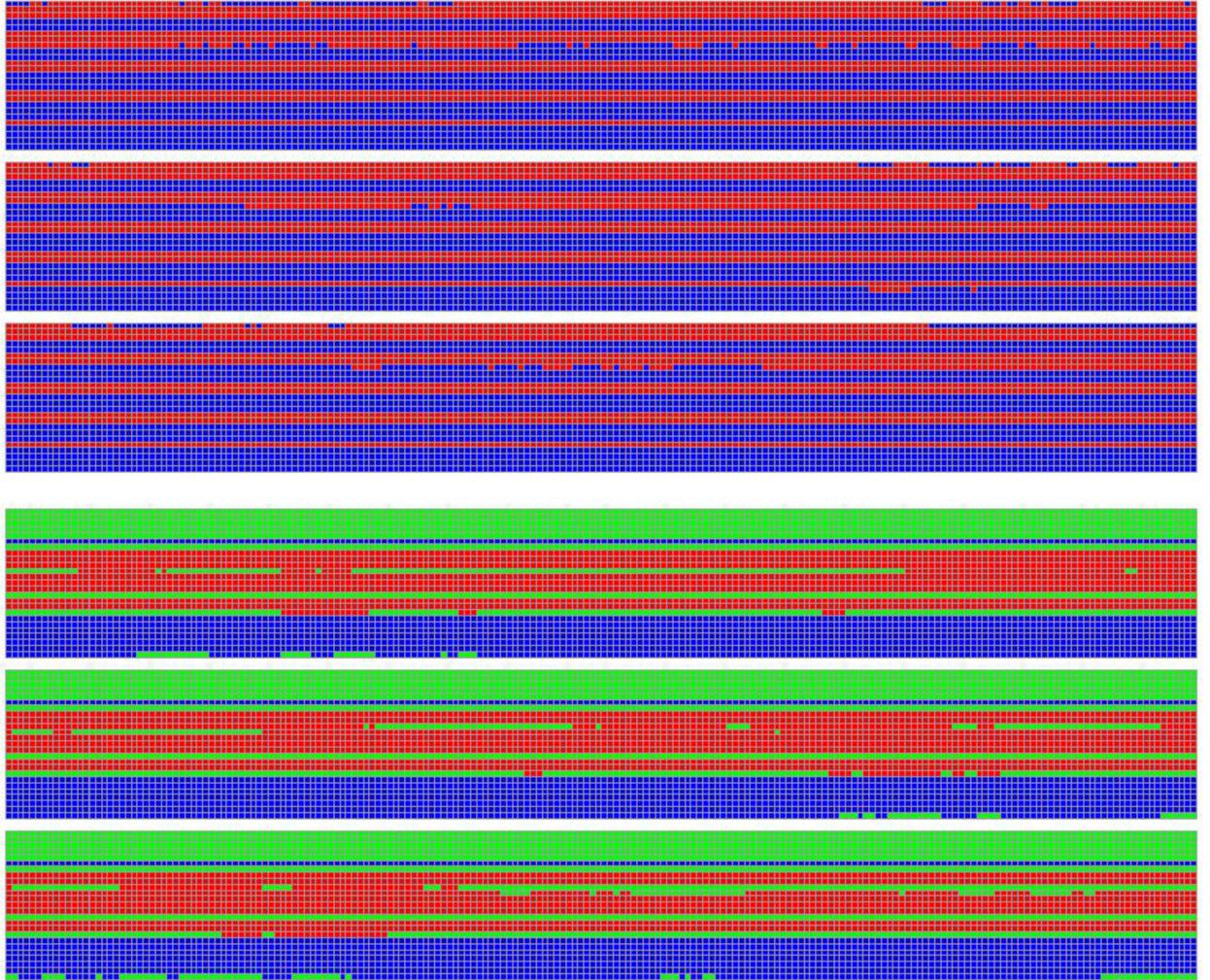


Fig. 3. Evolution of signs of weight filters in the last training epoch (or 600 iterations) of LeNet-5. Each of the six 25×200 blocks corresponds to evolution of the 5×5 convolutional filter over 200 iterations. **Top three rows:** Binary weights over the last 600 iterations of training, red/blue for sign values $1/-1$. **Bottom three rows:** Ternary weights over the last 600 iterations of training, red/green/blue for sign values $1/0/-1$.

30 epochs (50 in total). This suggests the network parameters have reached a local valley. However, Fig. 3 shows the iterating sequence of model parameters still have oscillating signs towards the end of training.

Fig. 3 shows the evolution of the quantized weights of one convolutional filter in the first convolutional layer during the last 600 iterations. To visualize the weights, each quantized filter is reshaped into a 25-

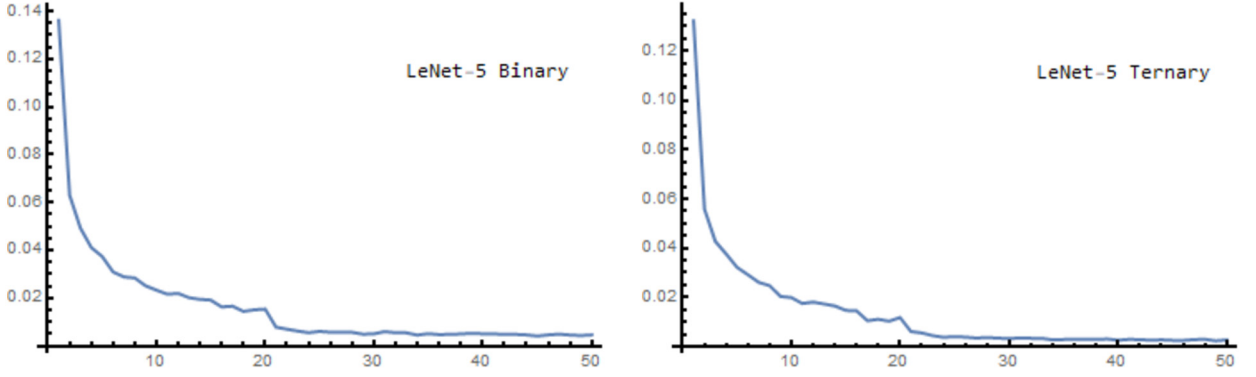


Fig. 4. LeNet-5 Training Loss v.s. Epoch. **Left:** Binary weights. **Bottom:** Ternary weights.

dimensional column vector. Each block (3 in a group) of size 25×200 corresponds to the evolution of the one filter during 200 iterations. As we can see from these two figures, a proportion of the weights do not converge to a limit but rather have oscillating signs.

5.3. CIFAR-10

We repeat the experiments on CIFAR-10 [13] with ResNet-20/VGG-11. We train ResNet-20 [8]/VGG-11 [19] with binary/ternary weights and 4-bits activation using QUANT for 200 epochs. We refer to the appendix for some figures that show similar oscillation behavior. Towards the end of training, although there has been no noticeable decay of training loss, we can see a clearer pattern of the oscillating signs of the weights.

6. Concluding remarks

We studied the convergence behavior of widely used QUANT algorithm [10,3,4,27] for the quantization of one-hidden-layer networks. We showed that the sequence of quantized weights $\{w^t\}$ generated by QUANT is generically divergent if the teacher parameters are not in a quantized state, and constructed an explicit example of oscillatory divergence behavior. Under conditions that teacher parameters and their quantized values are close enough, we proved the recurrence of QUANT algorithm at the global minimum.

Acknowledgment

This work was partially supported by NSF grants IIS-1632935, DMS-1854434, DMS-1924548, DMS-1924935, and DMS-1952644.

Appendix A

Lemma 1. *The expected coarse gradient of $\ell(w; Z)$ w.r.t. w is*

$$\tilde{\nabla} f(w) = \frac{\|v\|^2}{2\sqrt{2\pi}} \left(\frac{w}{\|w\|} - w^* \right). \quad (6)$$

Proof or Lemma 1. [24] gives

$$\tilde{\nabla} f(w) = \frac{\|v^*\|^2}{\sqrt{2\pi}} \left(\frac{w}{\|w\|} - \cos\left(\frac{\theta}{2}\right) \frac{\frac{w}{\|w\|} + w^*}{\left\| \frac{w}{\|w\|} + w^* \right\|} \right).$$

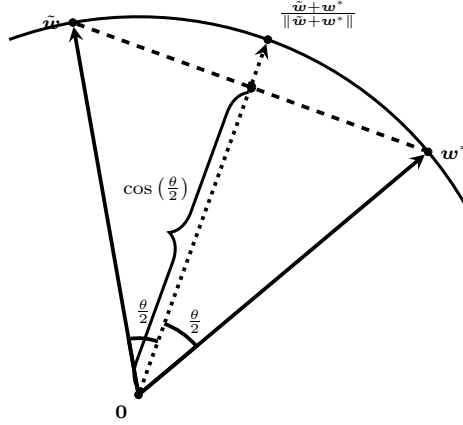


Fig. 5. 2-dim section of $\mathbb{R}^n$ spanned by $\tilde{\mathbf{w}}$ and $\mathbf{w}^*$.

Let $\tilde{\mathbf{w}} = \frac{\mathbf{w}}{\|\mathbf{w}\|}$, we can easily see from Fig. 5 that the coarse gradient can be further simplified as (6). $\square$

Proposition 1. Let $\mathbf{w}^t$ be any infinite sequence generated by Algorithm 1. If $|w_j^*| < \frac{1}{\sqrt{n}}$, then there exist infinitely many t_1 and t_2 values such that $w_j^{t_1} = \frac{1}{\sqrt{n}}$ and $w_j^{t_2} = -\frac{1}{\sqrt{n}}$.

Proof of Lemma 1. For notational simplicity, since $\|\mathbf{w}_j^*\| < \frac{1}{\sqrt{n}}$, we have

$$\alpha := \frac{1}{\sqrt{n}} - w_j^* > 0 \quad \text{and} \quad \beta := \frac{1}{\sqrt{n}} + w_j^* > 0.$$

Using Lemma 3 in Algorithm 1, we see that

$$\begin{aligned} y_j^{t+1} &= y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) \\ &= y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left(w_j^* - \frac{1}{\sqrt{n}} \widetilde{\text{sign}}(y_j^t) \right), \end{aligned}$$

and thus

$$y_j^{t+1} = \begin{cases} y_j^t - \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \alpha & \text{if } y_j^t \geq 0 \\ y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \beta & \text{if } y_j^t < 0 \end{cases}$$

Since y_j^t is bounded for each fixed $t \geq 0$ and $j \in [n]$, our desired result follows from our assumptions on learning rate η_t . $\square$

Corollary 1. If $\mathbf{w}^* \notin \mathcal{Q}$, any sequence $\{\mathbf{w}^t\}$ generated by Algorithm 1 does not converge.

Proof of Corollary 1. Since $\mathbf{w}^* \notin \tilde{\mathcal{Q}}_1^n$, we know there must exist some $j \in [n]$ such that $|w_j^*| < \frac{1}{\sqrt{n}}$ and Proposition 1 gives our desired result. $\square$

Example 1. Let $\mathbf{w}^* = \left(\frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{2}\sqrt{\frac{11}{3}}\right)$ so that the best the optimum $\widetilde{\text{proj}}_{\mathcal{Q}} \mathbf{w}^* = \left(\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2}\right)$. Let $\eta_t = \eta$, $\lambda = \frac{\eta \|\mathbf{v}\|^2}{6\sqrt{2\pi}}$ and

$$\begin{cases} y_1^0 \in (-\lambda, 0) \\ y_2^0 \in (0, \lambda) \\ y_3^0 \in (\lambda, 2\lambda) \\ y_4^0 \in (0, \infty) \end{cases}$$

the sequence $\{\mathbf{w}^t\}$ generated by Algorithm 1 with initialization $\mathbf{y}^0$ satisfies $\mathbf{w}^{t+3} = \mathbf{w}^t$ and $\mathbf{w}^t \neq \widetilde{\text{proj}}_{\mathcal{Q}} \mathbf{w}^*$ for all t .

Proof of Example 1. In order to show the periodicity, it suffices to show $w_j^{t+3} = w_j^t$. Note that $\tilde{\partial}_{w_4} f(\mathbf{w}) < 0$ we have $y_4^t > 0$ for all t since $w_4^0 > 0$. It follows that $w_4^t = w_4^0 = \frac{1}{2}$. Next, we would like to show the periodicity of w_j^t for $j \in [3]$. Note that

$$y_j^{t+1} = \begin{cases} y_j^t + \frac{\eta \|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left(w_j^* + \frac{1}{2} \right) & \text{if } y_j^t < 0 \\ y_j^t + \frac{\eta \|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left(-w_j^* + \frac{1}{2} \right) & \text{if } y_j^t \geq 0 \end{cases}$$

we choose $\mathbf{w}_j^* = \frac{1}{6}$ so that with

$$\lambda = \frac{\eta \|\mathbf{v}\|^2}{6\sqrt{2\pi}}$$

we have

$$y_j^{t+1} = \begin{cases} y_j^t + 2\lambda & \text{if } y_j^t < 0 \\ y_j^t - \lambda & \text{if } y_j^t \geq 0 \end{cases}$$

Hence, we have

$$\mathbf{w}^t = \begin{cases} \left(-\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2} \right) & \text{if } t \equiv 0(\text{mod } 3) \\ \left(\frac{1}{2}, -\frac{1}{2}, \frac{1}{2}, \frac{1}{2} \right) & \text{if } t \equiv 1(\text{mod } 3) \\ \left(\frac{1}{2}, \frac{1}{2}, -\frac{1}{2}, \frac{1}{2} \right) & \text{if } t \equiv 2(\text{mod } 3) \end{cases} \quad \square$$

Theorem 1 (Binary Case). If the optimum $\hat{\mathbf{w}} := \widetilde{\text{proj}}_{\mathcal{Q}^n} \mathbf{w}^*$ of (3) satisfies

$$\sum_{|w_j^*| < \frac{1}{\sqrt{n}}} |w_j^* - \hat{w}_j| < \frac{2}{\sqrt{n}}$$

then there exists infinitely many t values for any sequence $\{\mathbf{w}^t\}$ generated by Algorithm 1 such that $\mathbf{w}^t = \widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{w}^*)$.

Proof of Theorem 1 on $b = 1$. Without loss of generality, we can assume $w_j^* \geq 0$ for all $j \in [n]$ so that $\hat{w}_j = \frac{1}{\sqrt{n}}$ for all j .

Firstly, if $w_j^* > \frac{1}{\sqrt{n}}$, we know

$$y_j^{t+1} = y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) \geq y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left(w_j^* - \frac{1}{\sqrt{n}} \right),$$

so that

$$y_j^t \geq y_j^0 + \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left(\sum_{s=0}^{t-1} \eta_s \right) \left(w_j^* - \frac{1}{\sqrt{n}} \right)$$

where the right hand side goes to infinity and thus $w_j^t = \hat{w}_j$ for all but finitely many t values.

Secondly, if $w_j^* = \frac{1}{\sqrt{n}}$, we know when $w_j^t < 0$:

$$y_j^{t+1} = y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) = y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \frac{2}{\sqrt{n}}$$

holds so that there must exist some t such that $y_j^t > 0$. Once $y_j^t > 0$ we have $w_j^* = w_j^t$ so that $y_j^{t+1} = y_j^t$ and hence $w_j^t = \hat{w}_j$ for all but finitely many t values.

Third, if $w_j^* < \frac{1}{\sqrt{n}}$, we have $y_j^t \cdot \tilde{\partial}_j f(\mathbf{w}^t) > 0$ so that y_j^t is increasing when $y_j^t < 0$ and decreasing when $y_j^t > 0$. This tells us y_j^t is bounded uniformly in t . Furthermore,

$$y_j^t = y_j^0 + \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left[\left(\sum_{s=0}^{t-1} \mathbb{1}_{\{w_j^s > 0\}} \eta_s \right) \left(w_j^* - \frac{1}{\sqrt{n}} \right) + \left(\sum_{s=0}^{t-1} \mathbb{1}_{\{w_j^s < 0\}} \eta_s \right) \left(w_j^* + \frac{1}{\sqrt{n}} \right) \right].$$

For notation simplicity, we let

$$\alpha_j = \frac{1}{\sqrt{n}} - w_j^* > 0 \quad \text{and} \quad \beta_j = w_j^* + \frac{1}{\sqrt{n}} > 0,$$

$$a_j^t = \frac{1}{t} \sum_{s=0}^{t-1} \mathbb{1}_{\{w_j^s > 0\}} \eta_s \quad \text{and} \quad b_j^t = \frac{1}{t} \sum_{s=0}^{t-1} \mathbb{1}_{\{w_j^s < 0\}} \eta_s.$$

Now, we have

$$\frac{y_j^t - y_j^0}{t} = \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (-\alpha_j a_j^t + \beta_j b_j^t).$$

Since y_j^t is bounded for all $w_j^* < \frac{1}{\sqrt{n}}$, we let $t \rightarrow \infty$ so that left hand side vanishes and

$$\lim_{t \rightarrow \infty} \frac{b_j^t}{a_j^t + b_j^t} = \frac{\alpha_j}{\alpha_j + \beta_j}.$$

By assumption, we have

$$\lim_{t \rightarrow \infty} \sum_{j=1}^n \frac{b_j^t}{a_j^t + b_j^t} = \sum_{j=1}^n \frac{\alpha_j}{\alpha_j + \beta_j} < 1.$$

Hence, we know

$$\lim_{t \rightarrow \infty} \sum_{s=0}^{t-1} \mathbb{1}_{\{\mathbf{w}^s = \hat{\mathbf{w}}^*\}} \eta_s \geq \lim_{t \rightarrow \infty} \left[\left(1 - \sum_{j=1}^n \frac{b_j^t}{a_j^t + b_j^t} \right) \sum_{s=0}^{t-1} \eta_s \right] = \infty,$$

where we used the assumption $\sum_{t=0}^{\infty} \eta_t = \infty$. Now, the desired result follows. $\square$

Proposition 2 (Ternary Case). *Let $\mathbf{w}^t$ be any sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q}$, then $\{\mathbf{w}^t\}$ is not a converging sequence.*

Proof of Proposition 2. We prove by contradiction. Observe that $\mathcal{Q} \cap \mathcal{S}^{n-1}$ is a finite set, we know $\mathbf{w}^t$ converges to $\mathbf{w}^\infty$ is equivalent to $\mathbf{w}^t = \mathbf{w}^\infty$ for all but finitely many t values. Assume $\mathbf{w}^t = \mathbf{w}^\infty$ for all but finitely many t values, we know there exists some $T \geq 0$ such that $\mathbf{w}^t = \mathbf{w}^\infty$ for all $t \geq T$. Thus,

$$\begin{aligned} \mathbf{y}^{T+t} &= \mathbf{y}^T - \sum_{s=0}^{t-1} \eta_{T+s} \tilde{\nabla} f(\mathbf{w}^{T+s}) \\ &= \mathbf{y}^T - \left(\sum_{s=0}^{t-1} \eta_{T+s} \right) \tilde{\nabla} f(\mathbf{w}^\infty) \\ &= \mathbf{y}^t + \left(\sum_{s=0}^{t-1} \eta_{T+s} \right) \frac{\|\mathbf{v}\|^2}{2\sqrt{2}\pi} (\mathbf{w}^* - \mathbf{w}^\infty). \end{aligned}$$

Now, we have

$$\langle \mathbf{y}^{T+t}, \mathbf{w}^\infty \rangle = \langle \mathbf{y}^T, \mathbf{w}^\infty \rangle + \left(\sum_{s=0}^{t-1} \eta_{T+s} \right) \frac{\|\mathbf{v}\|^2}{2\sqrt{2}\pi} \langle \mathbf{w}^* - \mathbf{w}^\infty, \mathbf{w}^\infty \rangle$$

where

$$\langle \mathbf{w}^* - \mathbf{w}^\infty, \mathbf{w}^\infty \rangle = \langle \mathbf{w}^*, \mathbf{w}^\infty \rangle - 1 < 0.$$

Note that $\sum_{s=0}^{\infty} \eta_{T+s} = \infty$, there exists some $T_1(T)$, such that for all $t > T_1(T)$

$$\langle \mathbf{y}^t, \mathbf{w}^\infty \rangle < 0.$$

This contradicts Lemma 4 and our desired result follows. $\square$

Lemma 5. *Let $\{\mathbf{y}^t\}$ be any auxiliary sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q}$, then $\|\mathbf{y}^t\|_1$ converges to infinity as t increases.*

Proof of Lemma 5. $\mathcal{Q} \cap \mathcal{S}^{n-1}$ is a compact set because it is finite. Also, since $\mathcal{Q}$ is symmetric, $\mathbf{w}^* \notin \mathcal{Q}$ also implies $-\mathbf{w}^* \notin \mathcal{Q}$. It follows that

$$\alpha := \inf_{\mathbf{w} \in \mathcal{Q} \cap \mathcal{S}^{n-1}} \theta(\mathbf{w}^*, \mathbf{w}) \in (0, \pi).$$

Hence, for any $\mathbf{w} \in \mathcal{Q} \cap \mathcal{S}^{n-1}$ we have

$$\langle -\tilde{\nabla} f(\mathbf{w}), \mathbf{w}^* \rangle = \frac{\|\mathbf{v}\|^2}{2\sqrt{2}\pi} \langle \mathbf{w}^* - \mathbf{w}, \mathbf{w}^* \rangle \geq \frac{\|\mathbf{v}\|^2}{2\sqrt{2}\pi} (1 - \cos \alpha).$$

Now, we know

$$\begin{aligned}\langle \mathbf{y}^T, \mathbf{w}^* \rangle &= \langle \mathbf{y}^0, \mathbf{w}^* \rangle + \sum_{t=0}^{T-1} \eta_t \langle -\tilde{\nabla} f(\mathbf{w}^t), \mathbf{w}^* \rangle \\ &\geq \langle \mathbf{y}^0, \mathbf{w}^* \rangle + \left(\sum_{t=0}^{T-1} \eta_t \right) \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \cdot (1 - \cos \alpha).\end{aligned}$$

Let $T \rightarrow \infty$, we see that $\lim_{t \rightarrow \infty} \|\mathbf{y}^t\| = \infty$ which is equivalent to $\lim_{t \rightarrow \infty} \|\mathbf{y}^t\|_1 = \infty$. $\square$

Lemma 9. Let $\mathbf{w} = \text{proj}_{\mathcal{Q}}(\mathbf{y})$, then $|y_j| < \frac{1}{5n} \|\mathbf{y}\|_1$ implies $w_j = 0$.

Proof of Lemma 9. Without loss of generality, we assume $y_i \geq 0$ for all $i \in [n]$ and $y_j < \frac{1}{5n} \|\mathbf{y}\|_1$ for a fixed $j \in [n]$. Let $\delta = \frac{1}{5n} \|\mathbf{y}\|_1$ and

$$j_\delta := |\{i \in [n] : |y_i| \geq \delta\}|$$

we know $j_\delta \geq 1$ by the principle of drawer. Now, with

$$j^* = \arg \max_j \frac{\|\mathbf{y}_{[j]}\|_1^2}{j}$$

for any $1 \leq k \leq n - j_\delta$

$$\begin{aligned}& \frac{\|\mathbf{y}_{[j^*]}\|_1^2}{j^*} - \frac{\|\mathbf{y}_{[j_\delta+k]}\|_1^2}{j_\delta+k} \\ & \geq \frac{\|\mathbf{y}_{[j_\delta]}\|_1^2}{j_\delta} - \frac{\|\mathbf{y}_{[j_\delta+k]}\|_1^2}{j_\delta+k} \\ & = \frac{(j_\delta+k) \|\mathbf{y}_{[j_\delta]}\|_1^2 - j_\delta \|\mathbf{y}_{[j_\delta+k]}\|_1^2}{j_\delta(j_\delta+k)},\end{aligned}$$

where the numerator is

$$\begin{aligned}& k \|\mathbf{y}_{[j_\delta]}\|_1^2 - j_\delta \left(\|\mathbf{y}_{[j_\delta+k]}\|_1^2 - \|\mathbf{y}_{[j_\delta]}\|_1^2 \right) \\ & \geq k \left[\|\mathbf{y}_{[j_\delta]}\|_1^2 - j_\delta \delta \left(\|\mathbf{y}_{[j_\delta+k]}\|_1 + \|\mathbf{y}_{[j_\delta]}\|_1 \right) \right].\end{aligned}$$

With $\tau = \frac{\|\mathbf{y}_{[j_\delta+k]}\|_1}{n\delta}$, we have

$$\begin{aligned}& k \left[\|\mathbf{y}_{[j_\delta]}\|_1^2 - j_\delta \delta \left(\|\mathbf{y}_{[j_\delta+k]}\|_1 + \|\mathbf{y}_{[j_\delta]}\|_1 \right) \right] \\ & \geq k \left[\left(\|\mathbf{y}_{[j_\delta+k]}\|_1 - k\delta \right)^2 - 2n\delta \|\mathbf{y}_{[j_\delta+k]}\|_1 \right] \\ & = k(n\delta)^2 (\tau^2 - 4\tau + 1).\end{aligned}$$

Note that

$$\tau = \frac{\|\mathbf{y}_{[j_\delta+k]}\|_1}{n\delta} \geq \frac{\|\mathbf{y}\|_1 - n\delta}{n\delta} \geq 4,$$

we conclude that

$$\frac{\|\mathbf{y}_{[j^*]}\|_1^2}{j^*} > \frac{\|\mathbf{y}_{[j_\delta+k]}\|_1^2}{j_\delta+k}$$

and hence $j^* \leq j_\delta$. Now, Lemma 4 gives $w_j = 0$. $\square$

Lemma 10. *Let $\{\mathbf{w}^t\}$ and $\{\mathbf{y}^t\}$ be the sequence and the auxiliary sequence generated by Algorithm 1. Assume $\mathbf{w}^* \notin \mathcal{Q}$, the following statements hold.*

- If $w_j^* = 0$, then y_j^t is bounded and $w_j^t = 0$ for all but finitely many t values.
- If $w_j^* \neq 0$, then $\text{sign}(y_j^t) = \text{sign}(w_j^*)$ for all but finitely many t values.

Proof of Lemma 10. On the one hand, we consider the case $w_j^* = 0$, so that

$$y_j^{t+1} = y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) = y_j^t - \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} w_j^t.$$

Note that Lemma 4 shows y_j^t and w_j^t are of the same sign if $w_j^t \neq 0$, we know y_j^t is bounded by $C_j := \max\{|y_j^0|, \eta \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}}\}$. Moreover Lemma 5 shows $\|\mathbf{y}^t\|_1 > 5nC_j$ for all but finitely many t values. Finally, we see from Lemma 9 that $w_j^t = 0$ for all but finitely many t values.

On the other hand, consider the case $w_j^* \neq 0$. Without loss of generality, we can assume $w_j^* > 0$. Note that whenever $y_j^t \leq 0$, we also have $w_j^t \leq 0$ so that

$$y_j^{t+1} = y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) \geq y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} w_j^*.$$

From the above inequality, we see that y_j^t is increasing where the increment is bounded from below by $\eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} w_j^* > 0$ where $\sum \eta_t = \infty$, so that there must exist some $T_j > 0$ such that $y_j^{T_j} > 0$. With Lemma 5, we can without loss of generality assume that $\|\mathbf{y}^t\|_1 \geq 5n\eta \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}}$ for all $t \geq T_j$. For ease of notation, we let $\delta = \eta \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}}$ so that $\|\mathbf{y}^t\|_1 \geq 5n\delta$ for all $t \geq T_j$. We shall next prove that $y_j^t \geq 0$ for all $t \geq T_j$. We prove by induction, assume $y_j^t > 0$ for some $t > T_j$ and show $y_j^{t+1} > 0$.

1. If $y_j^t > \delta$,

$$y_j^{t+1} = y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) \geq y_j^t - \delta > 0.$$

2. If $0 < y_j^t \leq \delta$, since $\|\mathbf{y}^t\|_1 \geq 5n\delta$, Lemma 9 shows $w_j^t = 0$ so that

$$y_j^{t+1} = y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) = y_j^t + \frac{\eta_t \delta}{\eta} w_j^* > y_j^t > 0.$$

Combining the above two cases, we get our desired result. $\square$

Lemma 6. *Let $\{\mathbf{y}^t\}$ be any auxiliary sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q}$, then any sub-sequential limit of $\tilde{\mathbf{y}}^t := \frac{\mathbf{y}^t}{\|\mathbf{y}^t\|}$ belongs to the closure of $\mathcal{O}(\mathbf{w}^*)$. Furthermore, if $\mathcal{O}(\mathbf{w}^*)$ is regular, then $\mathbf{y}^t$ lies in $\mathcal{O}(\mathbf{w}^*)$ for all but finitely many t values.*

Proof of Lemma 6. By Lemma 10, we see that $\text{sign}(y_j^t) = \text{sign}(w_j^*)$ for all $w_j^* \neq 0$. We only need to prove $w_j^* = 0$ implies $\lim_{t \rightarrow \infty} \tilde{y}_j^t = 0$. Indeed, by Lemma 10, we know that y_j^t is bounded by C_j while Lemma 5 tells us $\|\mathbf{y}^t\|$ goes to infinity. Thus, $\lim_{t \rightarrow \infty} \tilde{y}_j^t = \frac{y_j^t}{\|\mathbf{y}^t\|} = 0$. $\square$

Lemma 11. *Let $\{\mathbf{w}^*\}$ and $\{\mathbf{y}^t\}$ be any sequence and auxiliary sequence generated by Algorithm 1. Assuming that $\mathbf{w}^* \notin \mathcal{Q}_2^n$, we have the following fact.*

1. *If $|w_j^*| > |w_i^*|$, then $|y_j^t| > |y_i^t|$ for all but finitely many t values.*
2. *If $|w_j^*| = |w_i^*|$, then $||y_j^t| - |y_i^t||$ is bounded and $|w_j^t| = |w_i^t|$ for all but finitely many t values.*

Proof of Lemma 11. Without loss of generality, we can assume $w_1^* \geq w_2^* \geq \dots \geq w_n^* \geq 0$.

For the first statement, we only need to show that $w_j^* > w_{j+1}^*$ implies $y_j^t > y_{j+1}^t$ for all but finitely many t values. Note that whenever $y_j^t < y_{j+1}^t$, then Lemma 4 implies $w_j^t \leq w_{j+1}^t$, hence

$$\begin{aligned} & y_j^{t+1} - y_{j+1}^{t+1} \\ &= \left(y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) \right) - \left(y_{j+1}^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_{j+1}^* - w_{j+1}^t) \right) \\ &= (y_j^t - y_{j+1}^t) + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} [(w_j^* - w_{j+1}^*) + (w_{j+1}^t - w_j^t)] \\ &\geq (y_j^t - y_{j+1}^t) + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_{j+1}^*). \end{aligned}$$

Now that we know $y_j^t - y_{j+1}^t$ is increasing as long as it is negative and $\sum \eta_t = \infty$. Therefore, we conclude that there exist infinitely many t values such that $y_j^t - y_{j+1}^t > 0$. We can therefore assume $y_j^T - y_{j+1}^T > 0$, where T is the constant in Lemma 5 such that $\|\mathbf{y}^t\|_1 \geq 5n\sqrt{2}\epsilon$ for all $t \geq T$ where we set $\epsilon = \frac{\eta\|\mathbf{v}^*\|^2}{\sqrt{2\pi n}}$. Next, we would like to show $y_j^t - y_{j+1}^t > 0$ for all $t \geq T$ by induction.

Next, assuming $y_j^t - y_{j+1}^t > 0$, we want to show $y_j^{t+1} - y_{j+1}^{t+1} > 0$.

On the one hand, if $y_j^t - y_{j+1}^t \geq \epsilon$, we have

$$\begin{aligned} & y_j^{t+1} - y_{j+1}^{t+1} \\ &= (y_j^t - y_{j+1}^t) + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} [(w_j^* - w_{j+1}^*) + (w_{j+1}^t - w_j^t)] \\ &> (y_j^t - y_{j+1}^t) + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_{j+1}^t - w_j^t) \\ &\geq (y_j^t - y_{j+1}^t) - \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \frac{1}{\sqrt{n}} \geq (y_j^t - y_{j+1}^t) - \epsilon \geq 0. \end{aligned}$$

On the other hand, if $y_j^t - y_{j+1}^t < \epsilon$, we still have

$$y_j^{t+1} - y_{j+1}^{t+1} > (y_j^t - y_{j+1}^t) - \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^t - w_{j+1}^t),$$

so that it suffices to show $w_j^t = w_{j+1}^t$. From Lemma 4, we see that with

$$j^* = \arg \max_{j \in [n]} \frac{\|\mathbf{y}_{[j]}^t\|_1^2}{j}, \quad (7)$$

we only need to show $j \neq j^*$. We prove by contradiction, assuming $j = j^*$ so that $w_j^t > 0$ and $w_{j+1}^t = 0$. Lemma 9 shows $y_j^t \geq \frac{1}{5n} \|\mathbf{y}^t\|_1$. Also, (7) gives

$$\frac{\|\mathbf{y}_{[j-1]}^t\|_1^2}{j-1} \leq \frac{\|\mathbf{y}_{[j]}^t\|_1^2}{j} = \frac{(\|\mathbf{y}_{[j-1]}^t\|_1 + y_j^t)^2}{j}. \quad (8)$$

Simplifying the above inequality, we get

$$\left(\frac{\|\mathbf{y}_{[j-1]}^t\|_1}{y_j^t} \right)^2 - 2(j-1) \left(\frac{\|\mathbf{y}_{[j-1]}^t\|_1}{y_j^t} \right) - (j-1) \leq 0.$$

Left hand side is a quadratic function of $\left(\frac{\|\mathbf{y}_{[j-1]}^t\|_1}{y_j^t} \right)$, we know

$$\frac{\|\mathbf{y}_{[j-1]}^t\|_1}{y_j^t} \leq j-1 + \sqrt{j(j-1)} \leq n-1 + \sqrt{n(n-1)} < 2n. \quad (9)$$

We write equation (8) in a different way and get

$$j \geq \frac{(\|\mathbf{y}_{[j-1]}^t\|_1 + y_j^t)^2}{y_j^t (2\|\mathbf{y}_{[j-1]}^t\|_1 + y_j^t)}. \quad (10)$$

Now, we use $j = j^*$ again, to get

$$\frac{\|\mathbf{y}_{[j]}^t\|_1^2}{j} \leq \frac{\|\mathbf{y}_{[j+1]}^t\|_1^2}{j+1}. \quad (11)$$

Rewriting the above inequality, we get

$$j \leq \frac{(\|\mathbf{y}_{[j-1]}^t\|_1 + y_j^t)^2}{y_{j+1}^t (2\|\mathbf{y}_{[j-1]}^t\|_1 + 2y_j^t + y_{j+1}^t)}. \quad (12)$$

Combining (10) and (12), we get

$$(y_j^t - y_{j+1}^t)^2 - (2\|\mathbf{y}_{[j-1]}^t\|_1 + 4y_j^t)(y_j^t - y_{j+1}^t) + 2(y_j^t)^2 \leq 0.$$

Solving the above inequality, we get

$$\begin{aligned}
y_j^t - y_{j+1}^t &\geq \left\| \mathbf{y}_{[j-1]}^t \right\|_1 + 2y_j^t \\
&\quad - \sqrt{\left\| \mathbf{y}_{[j-1]}^t \right\|_1^2 + 4 \left\| \mathbf{y}_{[j-1]}^t \right\|_1 y_j^t + 2(y_j^t)^2}.
\end{aligned} \tag{13}$$

Combining (9) and (13), we get

$$y_j^t - y_{j+1}^t \geq (y_j^t)^2 \left(2n + 2 - \sqrt{4n^2 + 4n + 2} \right) > \frac{(y_j^t)^2}{2}. \tag{14}$$

Recalling that $y_j^t \geq \frac{1}{5n} \|\mathbf{y}^t\|_1 \geq \sqrt{2\epsilon}$, we have

$$y_j^t - y_{j+1}^t > \frac{1}{2} \left(\frac{\|\mathbf{y}^t\|_1}{5n} \right)^2 > \epsilon.$$

This contradiction shows $j \neq j^*$, and hence $w_j^t = w_{j+1}^t$ and it follows that $y_j^{t+1} > y_{j+1}^{t+1}$. Now, we have proved our first statement.

For the second statement, since $w_j^* = w_i^*$, we have

$$\begin{aligned}
&y_j^{t+1} - y_i^{t+1} \\
&= \left(y_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^* - w_j^t) \right) - \left(y_i^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_i^* - w_i^t) \right) \\
&= y_j^t - y_i^t - \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (w_j^t - w_i^t) = y_j^t - y_i^t - 2\frac{\eta_t \epsilon}{\eta} (w_j^t - w_i^t).
\end{aligned}$$

Hence, we know that $|y_j^t - y_i^t|$ is bounded by

$$C_{i,j} := \max \left\{ |y_j^0 - y_i^0|, \frac{\eta \|\mathbf{v}^*\|^2}{\sqrt{2\pi}} \right\}.$$

Without loss of generality, we can assume $j < i$ and $\min \{y_j^t, y_i^t\} \geq 0$ by Lemma 10. Recalling (14), we have $w_j^t \neq w_i^t$ implying that

$$|y_j^t - y_i^t| > \frac{\max \{y_j^t, y_i^t\}}{2} \geq \frac{1}{2} \left(\frac{\|\mathbf{y}^t\|}{5n} \right)^2$$

where the right hand side goes to infinity. This contradicts the boundedness of $|y_j^t - y_i^t|$ if there are infinitely many t values such that $w_j^t \neq w_i^t$. $\square$

Lemma 7. *Let $\{\mathbf{y}^t\}$ be any auxiliary sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q}$, then any sub-sequential limit of $\tilde{\mathbf{y}}^t := \frac{\mathbf{y}^t}{\|\mathbf{y}^t\|}$ belongs to the closure of $\text{Cone}(\mathbf{w}^*)$. Moreover, if $\text{Cone}(\mathbf{w}^*)$ is regular, then $\mathbf{y}^t \in \text{Cone}(\mathbf{w}^*)$ for all but finitely many t values.*

Proof of Lemma 7. Note that we already have Lemma 6, we only need to show for any sub-sequential limit $\mathbf{y}$ of $\tilde{\mathbf{y}}^t$, we have $\text{sign}(|y_j| - |y_i|) = \text{sign}(|w_j^*| - |w_i^*|)$. The first statement of Lemma 11 tells us that it is true for all $\text{sign}(|w_j^*| - |w_i^*|) \neq 0$. Thus, it suffices to show that $|w_j^*| = |w_i^*|$ implies $|y_j| = |y_i|$.

Note that the second statement of Lemma 11 says that $||y_j| - |y_i||$ is bounded by $C_{i,j}$, while Lemma 5 gives $\lim_{t \rightarrow \infty} \|\mathbf{y}^t\| = \infty$, we see that

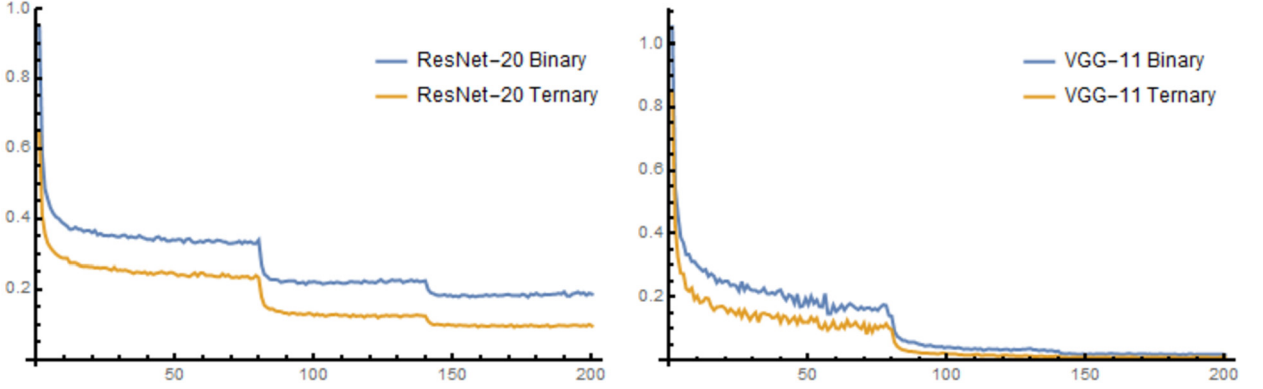


Fig. 6. Training Loss of CIFAR-10. **Left:** Binary/Ternary weight ResNet-20. **Right:** Binary/Ternary weight VGG-11.

$$|\tilde{y}_j| = \lim_{k \rightarrow \infty} \frac{|y_j^{t_k}|}{\|\mathbf{y}^{t_k}\|} = \lim_{k \rightarrow \infty} \frac{|y_i^{t_k}|}{\|\mathbf{y}^{t_k}\|} = |\tilde{y}_i|. \quad \square$$

Lemma 8. Let $\{\mathbf{w}^t\}$ be the sequence generated by Algorithm 1. If $\mathbf{w}^* \notin \mathcal{Q}$, then $\mathbf{w}^t \in \Lambda(\mathbf{w}^*)$ for all but finitely many t values.

Proof of Lemma 8. First, by Proposition 4, $\mathbf{y}^t \in \text{Cone}(\mathbf{w}^*)$ implies $\widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{y}^t) \in \Lambda(\mathbf{w}^*)$.

Second, let $\tilde{\partial}\text{Cone}(\mathbf{w}^*) = \overline{\text{Cone}(\mathbf{w}^*)} - \text{Cone}(\mathbf{w}^*)$. Now, a non-zero $\mathbf{y}^t \in \tilde{\partial}\text{Cone}(\mathbf{w}^*)$ implies $\text{Cone}(\mathbf{y}^t) \subset \tilde{\partial}\text{Cone}(\mathbf{w}^*)$ so that we also have $\widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{y}^t) \in \Lambda(\mathbf{y}^t) \subset \Lambda(\mathbf{w}^*)$.

Third, by compactness of $\overline{\text{Cone}(\mathbf{w}^*)} \cap \mathcal{S}^{n-1}$, we know there exists some $\epsilon > 0$ such that $\tilde{\mathbf{y}}^t := \frac{\mathbf{y}^t}{\|\mathbf{y}^t\|}$ lies in ϵ -neighborhood of $\text{Cone}(\mathbf{w}^*) \cap \mathcal{S}^{n-1}$ implying $\widetilde{\text{proj}}_{\mathcal{Q}}(\mathbf{y}^t) \in \Lambda(\mathbf{w}^*)$.

Finally, Lemma 7 suggests $\tilde{\mathbf{y}}^t$ lies in ϵ -neighborhood of $\text{Cone}(\mathbf{w}^*)$ for all but finitely many t values. We get our desired result. $\square$

Theorem 1 (Ternary Case). Let $\{\mathbf{z}_j\}_{j=1}^k = \Lambda(\mathbf{w}^*)$ where $\mathbf{z}_1 = \widetilde{\text{proj}}_{\mathcal{Q}}\mathbf{w}^*$ is the optimum and $\mathbf{w}^* = \sum_{j=1}^k \lambda_j \mathbf{z}_j$. If $0 < \sum_{j=2}^k \lambda_j < 1$, we have $\mathbf{w}^t = \widetilde{\text{proj}}_{\mathcal{Q}}\mathbf{w}^*$ for infinite many t values, where $\mathbf{w}^t$ is any infinite sequence generated by Algorithm 1 with any initialization.

Proof of Theorem 1 (Ternary Case). Note that Lemma 7 suggests $\tilde{\mathbf{y}}^t = \frac{\mathbf{y}^t}{\|\mathbf{y}^t\|}$ lies in ϵ -neighborhood of $\text{Cone}(\mathbf{w}^*)$ for all but finitely many t values. Let $\Lambda(\mathbf{w}^*) = \{\mathbf{z}_1, \dots, \mathbf{z}_k\}$ and define μ_j^t be the constants such that

$$\mathbf{y}^t = \sum_{j=1}^k \mu_j^t \mathbf{z}_j$$

which is determined uniquely by $\mathbf{y}^t$.

Let $\mathbf{w}^t = \mathbf{z}_{j_t}$, we know from Algorithm 1 that

$$\mathbf{y}^{t+1} - \mathbf{y}^t = \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} (\mathbf{w}^* - \mathbf{z}_{j_t}).$$

Thus

$$\sum_{j=2}^k \mu_j^{t+1} = \sum_{j=2}^k \mu_j^t + \eta_t \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left[\left(\sum_{j=2}^k \lambda_j \right) - 1 \right].$$

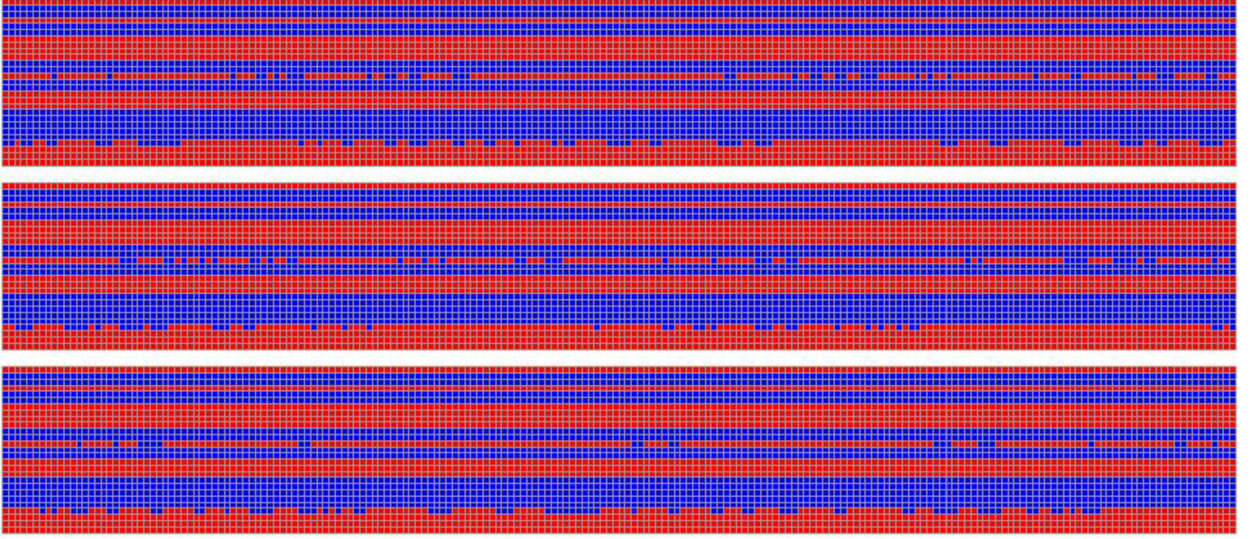


Fig. 7. Evolution of signs of weight filters in the last training epoch (or 600 iterations) of ResNet-20. Each of the three 27×200 blocks corresponds to evolution of the $3 \times 3 \times 3$ convolutional filter over 200 iterations. Binary weights over the last 600 iterations of training, red/blue for sign values $1/-1$.

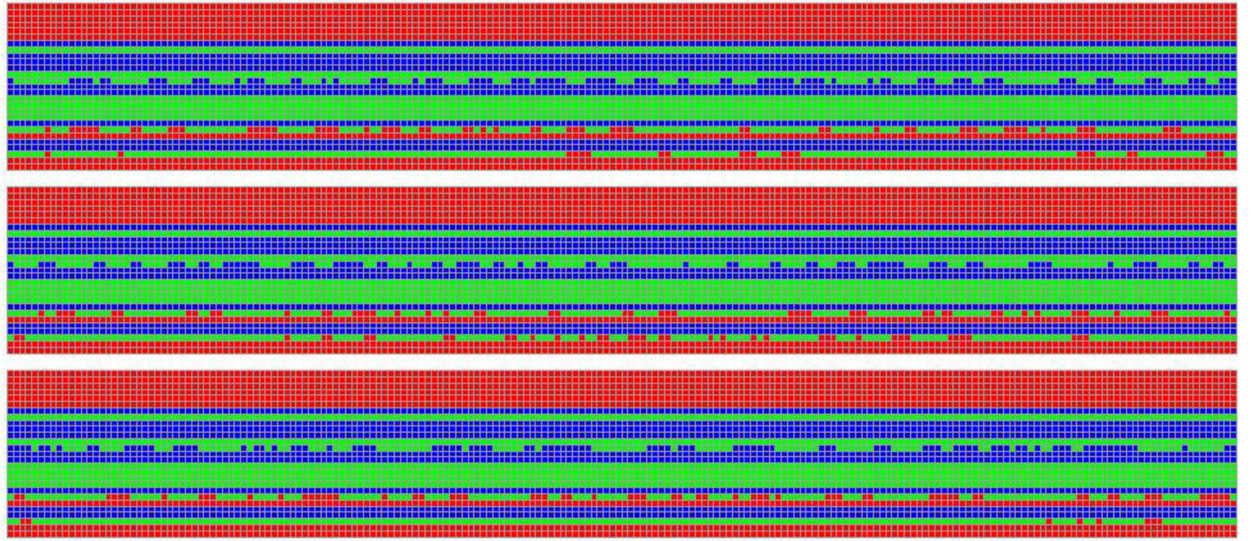


Fig. 8. Evolution of signs of weight filters in the last training epoch (or 600 iterations) of ResNet-20. Each of the three 27×200 blocks corresponds to evolution of the $3 \times 3 \times 3$ convolutional filter over 200 iterations. Ternary weights over the last 600 iterations of training, red/green/blue for sign values $1/0/-1$.

It follows that

$$\sum_{j=2}^k \mu_j^t = \text{Constant} + \left(\sum_{s=0}^{t-1} \eta_s \right) \frac{\|\mathbf{v}\|^2}{2\sqrt{2\pi}} \left[\left(\sum_{j=2}^k \lambda_j \right) - 1 \right] < 0,$$

for large t 's. Now we see that when t is large enough, $\tilde{\mathbf{y}}^t$ is bounded away from $\text{Cone}(\mathbf{w}^*)$ which contradicts Lemma 7 and our desired result follows (see Figs. 6–10). $\square$

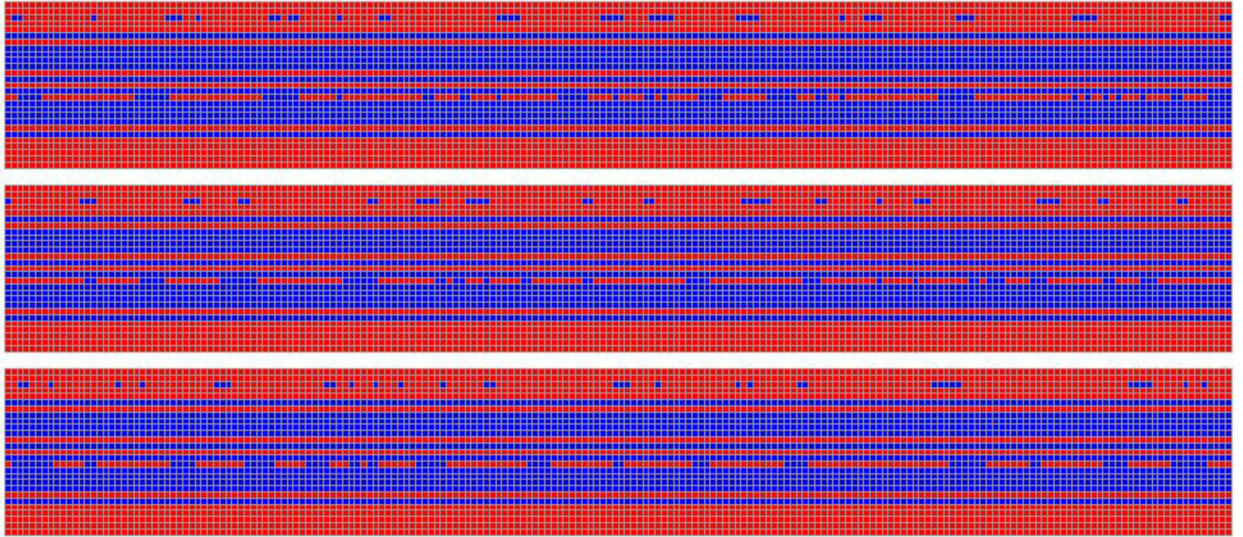


Fig. 9. Evolution of signs of weight filters in the last training epoch (or 600 iterations) of VGG-11. Each of the three 27×200 blocks corresponds to evolution of the $3 \times 3 \times 3$ convolutional filter over 200 iterations. Binary weights over the last 600 iterations of training, red/blue for sign values $1/-1$.

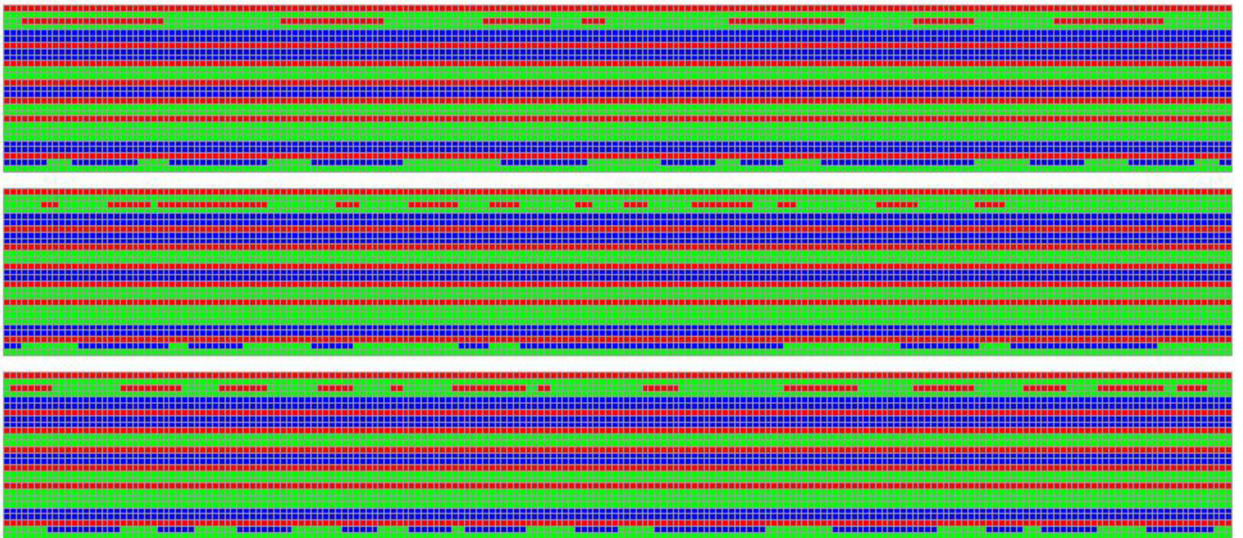


Fig. 10. Evolution of signs of weight filters in the last training epoch (or 600 iterations) of VGG-11. Each of the three 27×200 blocks corresponds to evolution of the $3 \times 3 \times 3$ convolutional filter over 200 iterations. Ternary weights over the last 600 iterations of training, red/green/blue for sign values $1/0/-1$.

References

- [1] Yoshua Bengio, Nicholas Léonard, Aaron Courville, Estimating or propagating gradients through stochastic neurons for conditional computation, arXiv preprint, arXiv:1308.3432, 2013.
- [2] Yaniv Blumenfeld, Dar Gilboa, Daniel Soudry, A mean field theory of quantized deep networks: the quantization-depth trade-off, in: Advances in Neural Information Processing Systems, 2019, pp. 7036–7046.
- [3] Zhaowei Cai, Xiaodong He, Jian Sun, Nuno Vasconcelos, Deep learning with low precision by half-wave gaussian quantization, in: IEEE Conference on Computer Vision and Pattern Recognition, 2017.
- [4] Jungwook Choi, Zhuo Wang, Swagath Venkataramani, Pierce I-Jen Chuang, Vijayalakshmi Srinivasan, Kailash Gopalakrishnan Pact, Parameterized clipping activation for quantized neural networks, arXiv preprint, arXiv:1805.06085, 2018.
- [5] Matthieu Courbariaux, Yoshua Bengio, Jean-Pierre David Binaryconnect, Training deep neural networks with binary weights during propagations, in: Advances in Neural Information Processing Systems, 2015, pp. 3123–3131.
- [6] J. Deng, W. Dong, R. Socher, L. Li, K. Li, F. Li, Imagenet: a large-scale hierarchical image database, in: IEEE Conference on Computer Vision and Pattern Recognition, 2009, pp. 248–255.

- [7] Song Han, Huizi Mao, William J. Dally, Deep compression: compressing deep neural networks with pruning, trained quantization and huffman coding, arXiv preprint, arXiv:1510.00149, 2015.
- [8] Kaiming He, Xiangyu Zhang, Shaoqing Ren, Jian Sun, Deep residual learning for image recognition, in: IEEE Conference on Computer Vision and Pattern Recognition, 2016, pp. 770–778.
- [9] Geoffrey Hinton, Neural networks for machine learning, coursera. Coursera, video lectures, 2012.
- [10] Itay Hubara, Matthieu Courbariaux, Daniel Soudry, Ran El-Yaniv, Yoshua Bengio, Binarized neural networks: training neural networks with weights and activations constrained to +1 or -1, arXiv preprint, arXiv:1602.02830, 2016.
- [11] Itay Hubara, Matthieu Courbariaux, Daniel Soudry, Ran El-Yaniv, Yoshua Bengio, Quantized neural networks: training neural networks with low precision weights and activations, J. Mach. Learn. Res. 18 (2018) 1–30.
- [12] Eric Jang, Shixiang Gu, Ben Poole, Categorical reparameterization with gumbel-softmax, in: International Conference on Learning Representations (ICLR), 2017.
- [13] Alex Krizhevsky, Learning multiple layers of features from tiny images, Tech Report, 2009.
- [14] Fengfu Li, Bo Zhang, Bin Liu, Ternary weight networks, arXiv preprint, arXiv:1605.04711, 2016.
- [15] Hao Li, Soham De, Zheng Xu, Christoph Studer, Hanan Samet, Tom Goldstein, Training quantized nets: a deeper understanding, in: Advances in Neural Information Processing Systems, 2017, pp. 5811–5821.
- [16] Tao Lin, Sebastian U. Stich, Luis Barba, Daniil Dmitriev, Martin Jaggi, Dynamic model pruning with feedback, in: International Conference on Learning Representations, 2020.
- [17] Christos Louizos, Matthias Reisser, Tijmen Blankevoort, Efstratios Gavves, Max Welling, Relaxed quantization for discretized neural networks, in: International Conference on Learning Representations, 2019.
- [18] Mohammad Rastegari, Vicente Ordonez, Joseph Redmon, Ali Farhadi Xnor-net, Imagenet classification using binary convolutional neural networks, in: European Conference on Computer Vision, Springer, 2016, pp. 525–542.
- [19] Karen Simonyan, Andrew Zisserman, Very deep convolutional networks for large-scale image recognition, arXiv preprint, arXiv:1409.1556, 2014.
- [20] Dimitrios Stamoulis, Ruizhou Ding, Di Wang, Dimitrios Lymberopoulos, Nissanka Bodhi Priyantha, Jie Liu, Diana Marculescu, Single-path mobile automl: efficient convnet design and nas hyperparameter optimization, IEEE J. Sel. Top. Signal Process. (2020).
- [21] Stefan Uhlich, Lukas Mauch, Fabien Cardinaux, Kazuki Yoshiyama, Javier Alonso García, Stephen Tiedemann, Thomas Kemp, Akira Nakamura, Mixed precision dnns: all you need is a good parametrization, in: International Conference on Learning Representations (ICLR), 2020.
- [22] Xia Xiao, Zigeng Wang, Sanguthevar Rajasekaran Autoprune, Automatic network pruning by regularizing auxiliary parameters, in: Advances in Neural Information Processing Systems, 2019, pp. 13681–13691.
- [23] Canran Xu, Ruijiang Li, Relation embedding with dihedral group in knowledge graph, in: Annual Conference of the Association for Computational Linguistics, 2019.
- [24] Penghang Yin, Jiancheng Lyu, Shuai Zhang, Stanley J. Osher, Yingyong Qi, Jack Xin, Understanding straight-through estimator in training activation quantized neural nets, in: International Conference on Learning Representations, 2019.
- [25] Penghang Yin, Shuai Zhang, Jack Xin, Yingyong Qi, Training ternary neural networks with exact proximal operator, arXiv:1612.06052 [abs], 2016.
- [26] Aojun Zhou, Anbang Yao, Yiwen Guo, Lin Xu, Yurong Chen, Incremental network quantization: towards lossless CNNs with low-precision weights, arXiv preprint, arXiv:1702.03044, 2017.
- [27] Shuchang Zhou, Yuxin Wu, Zekun Ni, Xinyu Zhou, He Wen, Yuheng Zou Dorefa-net, Training low bitwidth convolutional neural networks with low bitwidth gradients, arXiv preprint, arXiv:1606.06160, 2016.
- [28] Chenzhuo Zhu, Song Han, Huizi Mao, William J. Dally, Trained ternary quantization, arXiv preprint, arXiv:1612.01064, 2016.