
Efficient Neural Network Training via Forward and Backward Propagation Sparsification

Xiao Zhou^{*1}, Weizhong Zhang^{*1}, Zonghao Chen², Shizhe Diao¹, Tong Zhang^{†1}

¹ Hong Kong University of Science and Technology, ² Tsinghua University

xzhoubi@connect.ust.hk, zhangweizhongzju@gmail.com

czh17@mails.tsinghua.edu.cn, sdiaoaa@ust.hk, tongzhang@tongzhang-ml.org

Abstract

Sparse training is a natural idea to accelerate the training speed of deep neural networks and save the memory usage, especially since large modern neural networks are significantly over-parameterized. However, most of the existing methods cannot achieve this goal in practice because the chain rule based gradient (w.r.t. structure parameters) estimators adopted by previous methods require dense computation at least in the backward propagation step. This paper solves this problem by proposing an efficient sparse training method with completely sparse forward and backward passes. We first formulate the training process as a continuous minimization problem under global sparsity constraint. We then separate the optimization process into two steps, corresponding to weight update and structure parameter update. For the former step, we use the conventional chain rule, which can be sparse via exploiting the sparse structure. For the latter step, instead of using the chain rule based gradient estimators as in existing methods, we propose a variance reduced policy gradient estimator, which only requires two forward passes without backward propagation, thus achieving completely sparse training. We prove that the variance of our gradient estimator is bounded. Extensive experimental results on real-world datasets demonstrate that compared to previous methods, our algorithm is much more effective in accelerating the training process, up to an order of magnitude faster.

1 Introduction

In the last decade, deep neural networks (DNNs) [35, 11, 38] have proved their outstanding performance in various fields such as computer vision and natural language processing. However, training such large-sized networks is still very challenging, requiring huge computational power and storage. This hinders us from exploring larger networks, which are likely to have better performance. Moreover, it is a widely-recognized property that modern neural networks are significantly over-parameterized, which means that a fully trained network can always be sparsified dramatically by network pruning techniques [9, 8, 25, 46, 20] into a small sub-network with negligible degradation in accuracy. After pruning, the inference efficiency can be greatly improved. Therefore, a natural question is *can we exploit this sparsity to improve the training efficiency?*

The emerging technique called sparse network training [10] is closely related with our question, which can obtain sparse networks by training from scratch. We can divide existing methods into two categories, i.e., *parametric* and *non-parametric*, based on whether they explicitly parameterize network structures with trainable variables (termed *structure parameters*). Empirical results [24, 34, 44, 23] demonstrate that the sparse networks they obtain have comparable accuracy with those

^{*}Equal contribution

[†]Jointly with Google Research

obtained from network pruning. However, most of them narrowly aim at finding a sparse subnetwork instead of simultaneously sparsifying the computation of training by exploiting the sparse structure. As a consequence, it is hard for them to effectively accelerate the training process in practice on general platforms, e.g., Tensorflow [1] and Pytorch [31]. Detailed reasons are discussed below:

- Non-parametric methods find the sparse network by repeating a two-stage procedure that alternates between weight optimization and pruning [10, 6], or by adding a proper sparsity-inducing regularizer on the weights to the objective [22, 41]. The two-stage methods prune the networks in weight space and usually require retraining the obtained subnetwork from scratch every time when new weights are pruned, which makes training process even more time-consuming. Moreover, the computation of regularized methods is dense since the gradients of a zero-valued weights/filters are still nonzero.
- All the parametric approaches estimate the gradients based on chain rule. The gradient w.r.t. the structure parameters can be nonzero even when the corresponding channel/weight is pruned. Thus, to calculate the gradient via backward propagation, the error has to be propagated through all the neurons/channels. This means that the computation of backward propagation has to be dense. Concrete analysis can be found in Section 3.

We notice that some existing methods [4, 28] can achieve training speedup by careful implementation. For example, the dense to sparse algorithm [28] removes some channels if the corresponding weights are quite small for a long time. However, these methods always need to work with a large model at the beginning epochs and consume huge memory and heavy computation in the early stage. Therefore, even with such careful implementations, the speedups they can achieve are still limited.

In this paper, we propose an efficient channel-level parametric sparse neural network training method, which is comprised of **completely sparse** (See Remark 1) forward and backward propagation. We adopt channel-level sparsity since such sparsity can be efficiently implemented on the current training platforms to save the computational cost. In our method, we first parameterize the network structure by associating each filter with a binary mask modeled as an independent Bernoulli random variable, which can be continuously parameterized by the probability. Next, inspired by the recent work [47], we globally control the network size during the whole training process by controlling the sum of the Bernoulli distribution parameters. Thus, we can formulate the sparse network training problem into a constrained minimization problem on both the weights and structure parameters (i.e., the probability). The main novelty and contribution of this paper lies in our efficient training method called *completely sparse neural network training* for solving the minimization problem. Specifically, to fully exploit the sparse structure, we separate training iteration into two parts, i.e., weight update and structure parameter update. For weight update, the conventional backward propagation is used to calculate the gradient, which can be sparsified completely because the gradients of the filters with zero valued masks are also zero. For structure parameter update, we develop a new **variance reduced policy gradient estimator** (VR-PGE). Unlike the conventional chain rule based gradient estimators (e.g., straight through[2]), VR-PGE estimates the gradient via two forward propagations, which is completely sparse because of the sparse subnetwork. Finally, extensive empirical results demonstrate that our method can significantly accelerate the training process of neural networks.

The main contributions of this paper can be summarized as follows:

- We develop an efficient sparse neural network training algorithm with the following three appealing features:
 - In our algorithm, the computation in both forward and backward propagations is completely sparse, i.e., they do not need to go through any pruned channels, making the computational complexity significantly lower than that in standard training.
 - During the whole training procedure, our algorithm works on small sub-networks with the *target sparsity* instead of follows a dense-to-sparse scheme.
 - Our algorithm can be implemented easily on widely-used platforms, e.g., Pytorch and Tensorflow, to achieve practical speedup.
- We develop a variance reduced policy gradient estimator VR-PGE specifically for sparse neural network training, and prove that its variance is bounded.
- Experimental results demonstrate that our methods can achieve significant speed-up in training sparse neural networks. This implies that our method can enable us to explore larger-sized neural networks in the future.

Remark 1. We call a sparse training algorithm *completely sparse* if both its forward and backward propagation do not need to go through any pruned channels. For such algorithms, the computational cost in forward and backward propagation can be roughly reduced to $\rho^2 * 100\%$, with ρ being the ratio of remaining unpruned channels.

2 Related Work

In this section, we briefly review the studies on neural network pruning, which refers to the algorithms that prune DNNs after fully trained, and the recent works on sparse neural network training.

2.1 Neural Network Pruning

Network Pruning [10] is a promising technique for reducing the model size and inference time of DNNs. The key idea of existing methods [10, 8, 46, 20, 27, 13, 48, 40, 43, 32, 16] is to develop effective criteria (e.g, weight magnitude) to identify and remove the massive unimportant weights contained in networks after training. To achieve practical speedup on general devices, some of them prune networks in a structured manner, i.e., remove the weights in a certain group (e.g., filter) together, while others prune the weights individually. It has been reported in the literature [8, 25, 46, 20] that they can improve inference efficiency and reduce memory usage of DNNs by orders of magnitudes with minor loss in accuracy, which enables the deployment of DNNs on low-power devices.

We notice that although some pruning methods can be easily extended to train sparse networks, they cannot accelerate or could even slow down the training process. One reason is they are developed in the scenario that a fully trained dense network is given, and cannot work well on the models learned in the early stage of training. Another reason is after each pruning iteration, one has to fine tune or even retrain the network for lots of epoch to compensate the caused accuracy degradation.

2.2 Sparse Neural Network Training

The research on sparse neural network training has emerged in the recent years. Different from the pruning methods, they can find sparse networks without pre-training a dense one. Existing works can be divided into four categories based on their granularity in pruning and whether the network structures are explicitly parameterized. To the best of our knowledge, no significant training speedups achieved in practice are reported in the literature. Table 1 summarizes some representative works.

Table 1: Some representative works in sparse neural network training.

granularity	non-parametric	parametric
weight-level	[5, 6, 48, 22, 18, 29, 39, 30, 4]	[42, 37, 26, 47, 18]
channel-level	[41, 12]	[19, 24, 44, 26, 16]

Weight-level non-parametric methods, e.g., [6, 10, 48, 29, 30], always adopt a two-stage training procedure that alternates between weight optimization and pruning. They differ in the schedules of tuning the prune ratio over training and layers. [10] prunes the weights with the magnitude below a certain threshold and [48, 6] gradually increase the pruning rate during training. [30, 5] automatically reallocate parameters across layers during training via controlling the global sparsity.

Channel-level non-parametric methods [12, 41] are proposed to achieve a practical acceleration in inference. [41] is a structured sparse learning method, which adds a group Lasso regularization into the objective function of DNNs with each group comprised of the weights in a filter. [12] proposes a soft filter pruning method. It zeroizes instead of hard pruning the filters with small ℓ_2 norm, after which these filters are treated the same with other filters in training. It is obvious that these methods cannot achieve significant speedup in training since they need to calculate the full gradient in backward propagation although the forward propagation could be sparsified if implemented carefully.

Parametric methods multiply each weight/channel with a binary [47, 44, 37, 42] or continuous [24, 26, 19, 18] mask, which can be either deterministic [24, 42] or stochastic [47, 44, 26, 37, 19, 18]. The mask is always parameterized via a continuous trainable variable, i.e., structure parameter. The

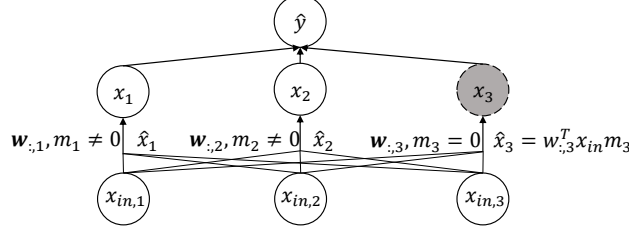


Figure 1: A fully connected network. w is the weight matrix of 1st layer, m_i is the mask of i -th neuron; $\hat{y}$, $\hat{x}_{in}$ and $\hat{x}_i$ are the output, input and preactivation. The 3rd neuron (in grey) is pruned.

sparsity is achieved by adding sparsity-inducing regularizers on the masks. The novelties of these methods lie in estimating the gradients w.r.t structure parameters in training. To be precise,

- *Deterministic Binary Mask.* [42] parameterizes its deterministic binary mask as a simple step function and estimates the gradients via sigmoid straight through estimator (STE) [2].
- *Deterministic Continuous Mask.* [24] uses the linear coefficients of batch normalization (BN) as a continuous mask and enforces most of them to 0 by penalizing the objective with ℓ_1 norm of the coefficients. [18] defines the mask as a soft threshold function with learnable threshold. These methods can estimate the gradients via standard backward propagation.
- *Stochastic Binary Mask.* [44, 37] model the mask as a bernoulli random variable and the gradients w.r.t. the parameters of bernoulli distributions are estimated via STE. [47] estimates the gradients via Gumbel-Softmax trick [15], which is more accurate than STE.
- *Stochastic Continuous Mask.* [26, 19] parameterize the mask as a continuous function $g(c, \epsilon)$, which is differentiable w.r.t. c , and ϵ is a parameter free noise, e.g., Gaussian noise $\mathcal{N}(0, 1)$. In this way, the gradients can be calculated via conventional backward propagation.

Therefore, we can see that all of these parametric methods estimate the gradients of the structure parameters based on the chain rule in backward propagation. This makes the training iteration cannot be sparsified by exploiting the sparse network structure. For the details, please refer to Section 3.

3 Why Existing Parametric Methods Cannot Achieve Practical Speedup?

In this section, we reformulate existing parametric channel-level methods into a unified framework to explain why they cannot accelerate the training process in practice.

Notice that convolutional layer can be viewed as a generalized fully connected layer, i.e., viewing the channels as neurons and convolution of two matrices as a generalized multiplication (see [7]). Hence, for simplicity, we consider the fully connected network in Figure 1. Moreover, since the channels in CNNs are corresponding to the neurons in fully connected networks, we consider *neuron-level instead of weight-level* sparse training in our example.

As discussed in Section 2, existing methods parameterize the 4 kinds of mask in the following ways:

- (i): $m_i = \phi(s_i)$; (ii): $m_i = \psi(s_i)$; (iii): $m_i = g(s_i, \epsilon), \epsilon \sim \mathcal{N}(0, 1)$; (iv): $m \sim \text{Bern}(p_i(s))$,

where the function $\phi(s_i)$ is binary, e.g., step function; $\psi(s_i)$ is a continuous function; $g(s_i, \epsilon)$ is differentiable w.r.t. s_i . All existing methods estimate the gradient of the loss $\ell(\hat{y}, y)$ w.r.t. s_i based on chain rule, which can be formulated into a unified form below.

Specifically, we take the pruned neuron x_3 in Figure 1 as an example, the gradient is calculated as

$$\nabla_{s_3} \ell(\hat{y}, y) = \underbrace{\frac{\partial \ell(\hat{y}, y)}{\partial \hat{x}_3}}_a \underbrace{(w_{:,3}^\top x_{in})}_{\text{forward}} \frac{\partial m_3}{\partial s_3}. \quad (1)$$

Existing parametric methods developed different ways to estimate $\frac{\partial m_3}{\partial s_3}$. Actually, for cases (ii) and (iii), the gradients are well-defined and thus can be calculated directly. STE is used to estimate the gradient in case (i) [42]. For cases (iv), [44, 37, 47] adopt STE and Gumbel-Softmax.

In Eqn.(1), the term (a) is always nonzero especially when $\hat{x}_3$ is followed by BN. Hence, we can see that even for the pruned neuron x_3 , the gradient $\frac{\partial m_3}{\partial s_3}$ can be nonzero in all four cases. This means the backward propagation has to go through all the neurons/channels, leading to dense computation.

At last, we can know from Eqn.(1) that forward propagation in existing methods cannot be completely sparse. Although $\mathbf{w}_{:,3}^\top \mathbf{x}_{in}$ can be computed sparsely as in general models $\mathbf{x}_{in}$ could be a sparse tensor of a layer with some channels being pruned, we need to calculate it for *each* neuron via forward propagation to calculate RHS of Eqn.(1). Thus, even if carefully implemented, the computational cost of forward propagation can only be reduced to $\rho * 100\%$ instead of $\rho^2 * 100\%$ as in inference.

That's why we argue that existing methods need dense computation at least in backward propagation. So they cannot speed up the training process effectively in practice.

Remark 2. *The authors of GrowEfficient [44] confirmed that actually they also calculated the gradient of q_c w.r.t. s_c in their Eqn.(6) via STE even if $q_c = 0$. Thus need dense backward propagation.*

4 Channel-level Completely Sparse Neural Network Training

Below, we present our sparse neural network training framework and the efficient training algorithm.

4.1 Framework of Channel-level Sparse Training

Given a convolutional network $f(x; \mathbf{w})$, let $\{\mathcal{F}_c : c \in \mathcal{C}\}$ be the set of filters with $\mathcal{C}$ being the set of indices of all the channels. To parameterize the network structure, we associate each $\mathcal{F}_c$ with a binary mask m_c , which is an independent Bernoulli random variable. Thus, each channel is computed as

$$\mathbf{x}_{out, c} = \mathbf{x}_{in} * (\mathcal{F}_c m_c),$$

with $*$ being the convolution operation. Inspired by [47], to avoid the problems, e.g., gradient vanishing, we parameterize m_c directly on the probability s_c , i.e., m_c equals to 1 and 0 with the probabilities s_c and $1 - s_c$, respectively. Thus, we can control the channel size by the sum of s_c . Following [47], we can formulate channel-level sparse network training into the following framework:

$$\begin{aligned} \min_{\mathbf{w}, \mathbf{s}} \mathbb{E}_{p(\mathbf{m}|\mathbf{s})} \mathcal{L}(\mathbf{w}, \mathbf{m}) &:= \frac{1}{N} \sum_{i=1}^N \ell(f(\mathbf{x}_i; \mathbf{w}, \mathbf{m}), \mathbf{y}_i) \\ s.t. \mathbf{w} &\in \mathbb{R}^n, \mathbf{s} \in \mathcal{S} := \{\mathbf{s} \in [0, 1]^{|\mathcal{C}|} : \mathbf{1}^\top \mathbf{s} \leq K\}, \end{aligned} \quad (2)$$

where $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$ is the training dataset, $\mathbf{w}$ is the weights of the original network, $f(\cdot; \cdot, \cdot)$ is the pruned network, and $\ell(\cdot, \cdot)$ is the loss function, e.g., cross entropy loss. $K = \rho|\mathcal{C}|$ controls the remaining channel size with ρ being the remaining ratio of the channels.

Discussion. We'd like to point out that although our framework is inspired by [47], our main contribution is the efficient solver comprised of completely sparse forward/backward propagation for Problem (2). Moreover, our framework can prune the weights in fully connected layers together, since we can associate each weight with an independent mask.

4.2 Completely Sparse Training with Variance Reduced Policy Gradient

Now we present our completely sparse training method, which can solve Problem (2) via *completely* sparse forward and backward propagation. The key idea is to separate the training iteration into filter update and structure parameter update so that the sparsity can be fully exploited.

4.2.1 Filter Update via Completely Sparse Computation

It is easy to see that the computation of the gradient w.r.t. the filters can be sparsified completely. To prove this point, we just need to clarify the following two things:

- *We do not need to update the filters corresponding to the pruned channels.* Consider a pruned channel c , i.e., $m_c = 0$, then due to the chain rule, we can have

$$\frac{\partial \ell(f(\mathbf{x}_i; \mathbf{w}, \mathbf{m}))}{\partial \mathcal{F}_c} = \frac{\partial \ell(f(\mathbf{x}_i; \mathbf{w}, \mathbf{m}))}{\partial \mathbf{x}_{out, c}} \frac{\partial \mathbf{x}_{out, c}}{\partial \mathcal{F}_c} \equiv 0,$$

the last equation holds since $\mathbf{x}_{out,c} \equiv 0$. This indicates that the gradient w.r.t the pruned filter $\mathcal{F}_c$ is always 0, and thus $\mathcal{F}_c$ does not need to be updated.

- *The error cannot pass the pruned channels via backward propagation.* Consider a pruned channel c , we denote its output before masking as $\hat{\mathbf{x}}_{out,c} = \mathbf{x}_{in} * \mathcal{F}_c$, then the error propagating through this channel can be computed as

$$\frac{\partial \ell(f(\mathbf{x}_i; \mathbf{w}, \mathbf{m}))}{\partial \hat{\mathbf{x}}_{out,c}} = \frac{\partial \ell(f(\mathbf{x}_i; \mathbf{w}, \mathbf{m}))}{\partial \mathbf{x}_{out,c}} \frac{\partial \mathbf{x}_{out,c}}{\hat{\mathbf{x}}_{out,c}} \equiv 0.$$

This demonstrates that to calculate the gradient w.r.t. the unpruned filters, the backward propagation does not need to go through any pruned channels.

Therefore, the filters can be updated via completely sparse backward propagation.

4.2.2 Structure Parameter Update via Variance Reduced Policy Gradient

We notice that **policy gradient estimator** (PGE) can estimate the gradient via forward propagation, avoiding the pathology of chain rule based estimators as discussed in Section 3. For abbreviation, we denote $\mathcal{L}(\mathbf{w}, \mathbf{m})$ as $\mathcal{L}(\mathbf{m})$ since $\mathbf{w}$ can be viewed as a constant here. The objective can be written as

$$\Phi(\mathbf{s}) = \mathbb{E}_{p(\mathbf{m}|\mathbf{s})} \mathcal{L}(\mathbf{m}),$$

which can be optimized using gradient descent:

$$\mathbf{s} \leftarrow \mathbf{s} - \eta \nabla \Phi(\mathbf{s}).$$

with learning rate η . One can obtain a stochastic unbiased estimate of the gradient $\nabla \Phi(\mathbf{s})$ using PGE:

$$\nabla \Phi(\mathbf{s}) = \mathbb{E}_{p(\mathbf{m}|\mathbf{s})} \mathcal{L}(\mathbf{m}) \nabla_{\mathbf{s}} \ln p(\mathbf{m}|\mathbf{s}), \quad (\text{PGE})$$

leading to Policy Gradient method, which may be regarded as a stochastic gradient descent algorithm:

$$\mathbf{s} \leftarrow \mathbf{s} - \eta \mathcal{L}(\mathbf{m}) \nabla_{\mathbf{s}} \ln p(\mathbf{m}|\mathbf{s}). \quad (3)$$

In Eqn.(3), $\mathcal{L}(\mathbf{m})$ can be computed via completely sparse forward propagation and the computational cost of $\nabla_{\mathbf{s}} \ln p(\mathbf{m}|\mathbf{s}) = \frac{\mathbf{m}-\mathbf{s}}{\mathbf{s}(1-\mathbf{s})}$ is negligible, therefore PGE is computationally efficient.

However, in accordance with the empirical results reported in [33, 15], we found that standard PGE suffers from high variance and does not work in practice. Below we will develop a **Variance Reduced Policy Gradient Estimator** (VR-PGE) starting from theoretically analyzing the variance of PGE.

Firstly, we know that this variance of PGE is

$$\mathbb{E}_{p(\mathbf{m}|\mathbf{s})} \mathcal{L}^2(\mathbf{m}) \|\nabla_{\mathbf{s}} \ln p(\mathbf{m}|\mathbf{s})\|_2^2 - \|\nabla \Phi(\mathbf{s})\|_2^2,$$

which can be large because $\mathcal{L}(\mathbf{m})$ is large.

Mean Field theory [36] indicates that, while $\mathcal{L}(\mathbf{m})$ can be large, the term $\mathcal{L}(\mathbf{m}) - \mathcal{L}(\mathbf{m}')$ is small when $\mathbf{m}$ and $\mathbf{m}'$ are two independent masks sampled from a same distribution $p(\mathbf{m}|\mathbf{s})$ (see the appendix for the details). This means that we may consider the following variance reduced preconditioned policy gradient estimator:

$$\mathbb{E}_{\mathbf{m}' \sim p(\mathbf{m}'|\mathbf{s})} \mathbb{E}_{\mathbf{m} \sim p(\mathbf{m}|\mathbf{s})} (\mathcal{L}(\mathbf{m}) - \mathcal{L}(\mathbf{m}')) H^\alpha(\mathbf{s}) \nabla_{\mathbf{s}} \ln p(\mathbf{m}|\mathbf{s}), \quad (\text{VR-PGE})$$

where $H^\alpha(\mathbf{s})$ is a specific diagonal preconditioning matrix

$$H^\alpha(\mathbf{s}) = \text{diag}(\mathbf{s} \circ (1 - \mathbf{s}))^\alpha, \quad (4)$$

with $\alpha \in (0, 1)$ and $\circ$ being the element-wise product. It plays a role as adaptive step size and it is shown that this term can reduce the variance of the stochastic PGE term $\nabla_{\mathbf{s}} \ln p(\mathbf{m}|\mathbf{s})$. The details can be found in the appendix. Thus $\Phi(\mathbf{s})$ can be optimized via:

$$\mathbf{s} \leftarrow \mathbf{s} - \eta (\mathcal{L}(\mathbf{m}) - \mathcal{L}(\mathbf{m}')) H^\alpha(\mathbf{s}) \nabla_{\mathbf{s}} \ln p(\mathbf{m}|\mathbf{s}). \quad (5)$$

In our experiments, we set α to be $\frac{1}{2}$ for our estimator VR-PGE. The theorem below demonstrates that VR-PGE can have bounded variance.

Algorithm 1 Completely Sparse Neural Network Training

Input: target remaining ratio ρ , a dense network $\mathbf{w}$, the step size η , and parameter α in (4) .

```
1: Initialize  $\mathbf{w}$ , let  $\mathbf{s} = \rho \mathbf{1}$ .  
2: for training epoch  $t = 1, 2, \dots, T$  do  
3:   for each training iteration do  
4:     Sample mini batch of data  $\mathcal{B} = \{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_B, \mathbf{y}_B)\}$ .  
5:     Sample  $\mathbf{m}^{(i)}$  from  $p(\mathbf{m}|\mathbf{s})$ ,  $i = 1, 2$ .  
6:     Update  $\mathbf{s}$  and  $\mathbf{w}$   
        $\mathbf{s} \leftarrow \text{proj}_{\mathcal{S}}(\mathbf{z})$  with  $\mathbf{z} = \mathbf{s} - \eta (\mathcal{L}_{\mathcal{B}}(\mathbf{w}, \mathbf{m}^{(1)}) - \mathcal{L}_{\mathcal{B}}(\mathbf{w}, \mathbf{m}^{(2)})) H^{\alpha}(\mathbf{s}) \frac{\mathbf{m}^{(1)} - \mathbf{s}}{\mathbf{s}(1 - \mathbf{s})}$ ,  
        $\mathbf{w} \leftarrow \mathbf{w} - \eta \nabla_{\mathbf{w}} \mathcal{L}_{\mathcal{B}}(\mathbf{w}, \mathbf{m}^{(1)})$   
7:   end for  
8: end for  
9: return A pruned network  $\mathbf{w} \circ \mathbf{m}$  by sampling a mask  $\mathbf{m}$  from the distribution  $p(\mathbf{m}|\mathbf{s})$ .
```

Theorem 1. Suppose $\mathbf{m}$ and $\mathbf{m}'$ are two independent masks sampled from the Bernoulli distribution $p(\mathbf{m}|\mathbf{s})$, then for any $\alpha \in [\frac{1}{2}, 1)$ and $\mathbf{s} \in (0, 1)^{|\mathcal{C}|}$, the variance is bounded for

$$(\mathcal{L}(\mathbf{m}) - \mathcal{L}(\mathbf{m}')) H^{\alpha}(\mathbf{s}) \nabla_{\mathbf{s}} \ln p(\mathbf{m}|\mathbf{s})$$

Finally, we provide a complete view of our sparse training algorithm in Algorithm 1, which is essentially a projected stochastic gradient descent equipped with our efficient gradient estimators above. The projection operator in Algorithm 1 can be computed efficiently using Theorem 1 of [47].

Discussion. In our algorithm, benefited from our constraint on $\mathbf{s}$, the channel size of the neural network during training can be strictly controlled. This is in contrast with GrowEfficient [44], which utilizes regularizer term to control the model size and has situations where model size largely drift away from desired. This will have larger demand for the GPU memory storage and have more risk that memory usage may explode, especially when we utilize sparse learning to explore larger models. Moreover, our forward and backward propagations are completely sparse, i.e., they do not need to go through any pruned channels. Therefore, the computational cost of each training iteration can be roughly reduced to $\rho^2 * 100\%$ of the dense network.

5 Experiments

In this section, we conduct a series of experiments to demonstrate the outstanding performance of our method. We divide the experiments into five parts. In part one, we compare our method with several state-of-the-art methods on CIFAR-10 [17] using VGG-16 [35], ResNet-20 [11] and WideResNet-28-10 [45] to directly showcase the superiority of our method. In part two, we directly compare with state-of-the-art method GrowEfficient [44] especially on extremely sparse regions, and on two high capacity networks VGG-19 [35] and ResNet-32 [11] on CIFAR-10/100 [17]. In part three, we conduct experiments on a large-scale dataset ImageNet [3] with ResNet-50 [11] and MobileNetV1 [14] and compare with GrowEfficient [44] across a wide sparsity region. In part four, we present the train-computational time as a supplementary to the conceptual train-cost savings to justify the applicability of sparse training method into practice. In part five, we present further analysis on epoch-wise train-cost dynamics and experimental justification of variance reduction of VR-PGE. Due to the space limitation, we postpone the experimental configurations, calculation schemes on train-cost savings and train-computational time and additional experiments into appendix.

5.1 VGG-16, ResNet-20 and WideResNet-28-10 on CIFAR-10

Table 2 presents Top-1 validation accuracy, parameters, FLOPs and train-cost savings comparisons with channel pruning methods L1-Pruning [20], SoftNet [12], ThiNet [27], Provable [21] and sparse training method GrowEfficient [44]. SoftNet can train from scratch but requires completely dense computation. Other pruning methods all require pretraining of dense model and multiple rounds of pruning and finetuning, which makes them slower than vanilla dense model training. Therefore the train-cost savings of these methods are below $1 \times$ and thus shown as ("") in Table 2.

Table 2: Comparison with the channel pruning methods L1-Pruning [20], SoftNet [12], ThiNet [27], Provable [21] and one channel sparse training method GrowEfficient [44] on CIFAR-10.

Model	Method	Val Acc(%)	Params(%)	FLOPs(%)	Train-Cost Savings($\times$)
VGG-16	Original	92.9	100	100	1 $\times$
	L1-Pruning	91.8	19.9	19.9	-
	SoftNet	92.1	36.0	36.1	-
	ThiNet	90.8	36.0	36.1	-
	Provable	92.4	5.7	15.0	-
	GrowEfficient	92.5	5.0	13.6	1.22 $\times$
	Ours	92.5	4.4	8.7	8.69$\times$
ResNet-20	Original	91.3	100	100	1 $\times$
	L1-Pruning	90.9	55.6	55.4	-
	SoftNet	90.8	53.6	50.6	-
	ThiNet	89.2	67.1	67.3	-
	Provable	90.8	37.3	54.5	-
	GrowEfficient	90.91	35.8	50.2	1.13 $\times$
	Ours	90.93	35.1	36.1	2.09$\times$
WRN-28-10	Original	96.2	100	100	1 $\times$
	L1-Pruning	95.2	20.8	49.5	-
	GrowEfficient	95.3	9.3	28.3	1.17 $\times$
	Ours	95.6	8.4	7.9	9.39$\times$

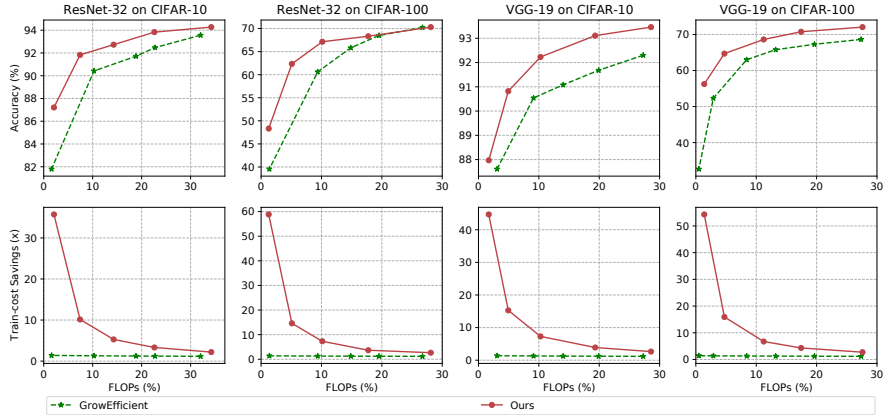


Figure 2: Comparison of Top-1 Validation Accuracy and Train-cost Savings on CIFAR-10/100.

GrowEfficient [44] is a recently proposed state-of-the-art channel-level sparse training method showing train-cost savings compared with dense training. As described in Section 3, GrowEfficient features completely dense backward and partially sparse forward pass, making its train-cost saving limited by $\frac{3}{2}$. By contrast, the train-cost savings of our method is not limited by any constraint. The details of how train-cost savings are computed can be found in appendix.

Table 2 shows that our method generally exhibits better performance in terms of validation accuracy, parameters and particularly FLOPs. In terms of train-cost savings, our method shows at least 1.85 $\times$ speed-up against GrowEfficient [44] and up to 9.39 $\times$ speed-up against dense training.

5.2 Wider Range of Sparsity on CIFAR-10/100 on VGG-19 and ResNet-32

In this section, we explore sparser regions of training efficiency to present a broader comparison with state-of-the-art channel sparse training method GrowEfficient [44].

We plot eight figures demonstrating the relationships between the Top-1 validation accuracy, FLOPs and train-cost savings. We find that our method generally achieves higher accuracy under same FLOPs

Table 3: Comparison with the channel pruning methods L1-Pruning [20], SoftNet [12], Provable [21] and one channel sparse training method GrowEfficient [44] on ImageNet-1K.

Model	Method	Val Acc(%)	Params(%)	FLOPs(%)	Train-Cost Savings($\times$)
ResNet-50	Original	77.0	100	100	1 $\times$
	L1-Pruning	74.7	85.2	77.5	-
	SoftNet	74.6	-	58.2	-
	Provable	75.2	65.9	70.0	-
	GrowEfficient	75.2	61.2	50.3	1.10 $\times$
	Ours	76.0	48.2	46.8	1.60 $\times$
	Ours	73.5	27.0	24.7	3.02 $\times$
	Ours	69.3	10.8	10.1	7.36$\times$

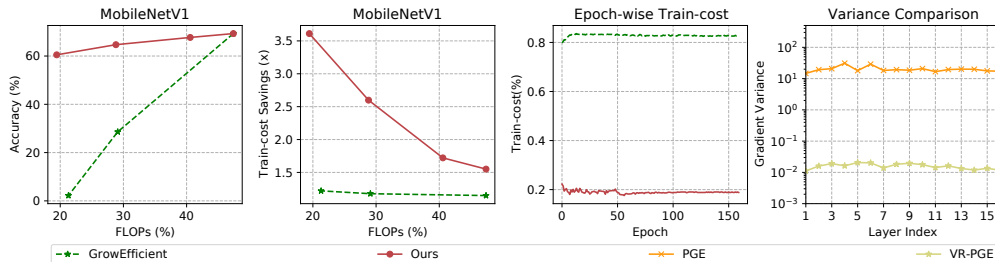


Figure 3: Top-1 Validation Accuracy and Train-cost Savings on MobileNetV1 on ImageNet. Epoch-wise Train-cost and Variance Comparison on VGG-19 on CIFAR-10.

settings. To be noted, the train-cost savings of our method is drastically higher than GrowEfficient [44], reaching up to 58.8 $\times$ when sparsity approaches 1.56% on ResNet-32 on CIFAR-100, while the speed-up of GrowEfficient is limited by $\frac{3}{2}$.

5.3 ResNet-50 and MobileNetV1 on ImageNet-1K

In this section, we present the performance boost obtained by our method on ResNet-50 and MobileNetV1 on ImageNet-1K [3]. Our method searches a model with 76.0% Top-1 accuracy, 48.2% parameters and 46.8% FLOPs beating all compared state-of-the-art methods. The train-cost saving comes up to 1.60 $\times$ and is not prominent due to the accuracy constraint to match up with compared methods. Therefore we give a harder limit to the channel size and present sparser results on the same Table 3, reaching up to 7.36 $\times$ speed-up while still preserving 69.3% Top-1 accuracy. For the already compact model MobileNetV1, we plot two figures in Figure 3 comparing with GrowEfficient [44]. We find that our method is much stabler in sparse regions and obtains much higher train-cost savings.

5.4 Actual Training Computational Time Testing

In this section, we provide actual training computational time on VGG-19 and CIFAR-10. The GPU in test is RTX 2080 Ti and the deep learning framework is Pytorch [31]. The intent of this section is to justify the feasibility of our method in reducing actual computational time cost, rather than staying in conceptual training FLOPs reduction. The computational time cost is measured by wall clock time, focusing on forward and backward propagation. We present training computational time in Table 4 with varying sparsity as in Figure 2. It shows that the computational time savings increases steadily with the sparsity. We also notice the gap between the savings in FLOPs and computational time. The gap comes from the difference between FLOPs and actual forward/backward time. More specifically, forward/backward time is slowed down by data-loading processes and generally affected by hardware latency and throughput, network architecture, etc. At extremely sparse regions, the pure computational time of sparse networks only occupies little of the forward/backward time and the cost of data management and hardware latency dominates the wall-clock time. Despite this gap, it can be expected that our train-cost savings can be better translated into real speed-up in exploring large

models where the pure computational time dominates the forward/backward time, which promises a bright future for making training infeasibly large models into practice.

Table 4: Train-computational Time on VGG-19 with CIFAR-10. The computational time saving is not as prominent as train-cost savings while still achieving nearly an order of reduction, preserving 87.97% accuracy.

Model	Val Acc(%)	Params(%)	FLOPs(%)	Train-Cost Savings($\times$)	Train-Computational Time(min)
VGG-19	93.84	100.00	100.00	1.00 $\times$	21.85 (1.00 $\times$)
	93.46	23.71	28.57	2.64 $\times$	14.04 (1.55 $\times$)
	93.11	12.75	19.33	3.89 $\times$	10.43 (2.09 $\times$)
	92.23	6.69	10.27	7.30 $\times$	6.83 (3.20 $\times$)
	90.82	3.06	4.94	15.28 $\times$	4.86 (4.50 $\times$)
	87.97	0.80	1.70	44.68 $\times$	2.95 (7.41 $\times$)

5.5 Further Analysis

[Epoch-wise Train-cost Dynamics of Sparse Training Process] We plot the train-cost dynamics in Figure 3. The vertical label is the ratio of train-cost to dense training, the inverse of train-cost savings. This demonstrates huge difference between our method and GrowEfficient [44]. The model searched by our method exhibits 92.73% Top-1 accuracy, 16.68% parameters, 14.28% FLOPs with 5.28 $\times$ train-cost savings, while the model searched by GrowEfficient exhibits 92.47% Top-1 accuracy, 18.08% parameters, 22.74% FLOPs with 1.21 $\times$ train-cost savings.

[Experimental Verification of Variance Reduction of VR-PGE against PGE] We plot the mean of variance of gradients of channels from different layers. The model checkpoint and input data are selected randomly. The gradients are calculated in two approaches, VR-PGE and PGE. From the rightmost graph of Figure 3, we find that the VR-PGE reduces variance significantly, up to 3 orders of magnitude.

6 Conclusion

This paper proposes an efficient sparse neural network training method with completely sparse forward and backward passes. A novel gradient estimator named VR-PGE is developed for updating structure parameters, which estimates the gradient via two sparse forward propagation. We theoretically proved that VR-PGE has bounded variance. In this way, we can separate the weight and structure update in training and making the whole training process completely sparse. Empirical results demonstrate that the proposed method can significantly accelerate the training process of DNNs in practice. This enables us to explore larger-sized neural networks in the future.

Acknowledgments and Disclosure of Funding

This work is supported by GRF 16201320.

References

- [1] M. Abadi, P. Barham, J. Chen, Z. Chen, A. Davis, J. Dean, M. Devin, S. Ghemawat, G. Irving, M. Isard, et al. Tensorflow: A system for large-scale machine learning. In *12th {USENIX} symposium on operating systems design and implementation ({OSDI} 16)*, pages 265–283, 2016.
- [2] Y. Bengio, N. Léonard, and A. Courville. Estimating or propagating gradients through stochastic neurons for conditional computation. *arXiv preprint arXiv:1308.3432*, 2013.
- [3] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei. Imagenet: A large-scale hierarchical image database. In *2009 IEEE conference on computer vision and pattern recognition*, pages 248–255. Ieee, 2009.
- [4] T. Dettmers and L. Zettlemoyer. Sparse networks from scratch: Faster training without losing performance. *arXiv preprint arXiv:1907.04840*, 2019.
- [5] U. Evci, T. Gale, J. Menick, P. S. Castro, and E. Elsen. Rigging the lottery: Making all tickets winners. In *International Conference on Machine Learning*, pages 2943–2952. PMLR, 2020.
- [6] J. Frankle, G. K. Dziugaite, D. M. Roy, and M. Carbin. Stabilizing the lottery ticket hypothesis. *arXiv preprint arXiv:1903.01611*, 2019.
- [7] Y. Gu, W. Zhang, C. Fang, J. D. Lee, and T. Zhang. How to characterize the landscape of overparameterized convolutional neural networks. In *Advances in Neural Information Processing Systems*, volume 33, pages 3797–3807, 2020.
- [8] Y. Guo, A. Yao, and Y. Chen. Dynamic network surgery for efficient dnns. In *Advances in neural information processing systems*, pages 1379–1387, 2016.
- [9] S. Han, H. Mao, and W. J. Dally. Deep compression: Compressing deep neural networks with pruning, trained quantization and huffman coding. *International Conference on Learning Representations*, 2016.
- [10] S. Han, J. Pool, J. Tran, and W. Dally. Learning both weights and connections for efficient neural network. In *Advances in neural information processing systems*, pages 1135–1143, 2015.
- [11] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [12] Y. He, G. Kang, X. Dong, Y. Fu, and Y. Yang. Soft filter pruning for accelerating deep convolutional neural networks. In *IJCAI International Joint Conference on Artificial Intelligence*, 2018.
- [13] Y. He, X. Zhang, and J. Sun. Channel pruning for accelerating very deep neural networks. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 1389–1397, 2017.
- [14] A. G. Howard, M. Zhu, B. Chen, D. Kalenichenko, W. Wang, T. Weyand, M. Andreetto, and H. Adam. Mobilenets: Efficient convolutional neural networks for mobile vision applications. *arXiv preprint arXiv:1704.04861*, 2017.
- [15] E. Jang, S. Gu, and B. Poole. Categorical reparameterization with gumbel-softmax. *International Conference on Learning Representations*, 2017.
- [16] M. Kang and B. Han. Operation-aware soft channel pruning using differentiable masks. In *International Conference on Machine Learning*, pages 5122–5131. PMLR, 2020.
- [17] A. Krizhevsky. Learning multiple layers of features from tiny images. *Master’s thesis, University of Tront*, 2009.
- [18] A. Kusupati, V. Ramanujan, R. Somani, M. Wortsman, P. Jain, S. Kakade, and A. Farhadi. Soft threshold weight reparameterization for learnable sparsity. In *Proceedings of the International Conference on Machine Learning*, July 2020.
- [19] C. Lemaire, A. Achkar, and P.-M. Jodoin. Structured pruning of neural networks with budget-aware regularization. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9108–9116, 2019.
- [20] H. Li, A. Kadav, I. Durdanovic, H. Samet, and H. P. Graf. Pruning filters for efficient convnets. *International Conference on Learning Representations*, 2017.
- [21] L. Liebenwein, C. Baykal, H. Lang, D. Feldman, and D. Rus. Provable filter pruning for efficient neural networks. In *International Conference on Learning Representations*, 2019.

- [22] B. Liu, M. Wang, H. Foroosh, M. Tappen, and M. Pensky. Sparse convolutional neural networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 806–814, 2015.
- [23] S. Liu, T. Chen, X. Chen, Z. Atashgahi, L. Yin, H. Kou, L. Shen, M. Pechenizkiy, Z. Wang, and D. C. Mocanu. Sparse training via boosting pruning plasticity with neuroregeneration. *arXiv preprint arXiv:2106.10404*, 2021.
- [24] Z. Liu, J. Li, Z. Shen, G. Huang, S. Yan, and C. Zhang. Learning efficient convolutional networks through network slimming. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 2736–2744, 2017.
- [25] Z. Liu, M. Sun, T. Zhou, G. Huang, and T. Darrell. Rethinking the value of network pruning. In *International Conference on Learning Representations*, 2018.
- [26] C. Louizos, M. Welling, and D. P. Kingma. Learning sparse neural networks through l₀ regularization. In *International Conference on Learning Representations*, 2018.
- [27] J.-H. Luo, J. Wu, and W. Lin. Thinet: A filter level pruning method for deep neural network compression. In *Proceedings of the IEEE international conference on computer vision*, pages 5058–5066, 2017.
- [28] S. Lym, E. Choukse, S. Zangeneh, W. Wen, S. Sanghavi, and M. Erez. Prunetrain: fast neural network training by dynamic sparse model reconfiguration. In *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*, pages 1–13, 2019.
- [29] D. C. Mocanu, E. Mocanu, P. Stone, P. H. Nguyen, M. Gibescu, and A. Liotta. Scalable training of artificial neural networks with adaptive sparse connectivity inspired by network science. *Nature communications*, 9(1):1–12, 2018.
- [30] H. Mostafa and X. Wang. Parameter efficient training of deep convolutional neural networks by dynamic sparse reparameterization. In *International Conference on Machine Learning*, pages 4646–4655. PMLR, 2019.
- [31] A. Paszke, S. Gross, F. Massa, A. Lerer, J. Bradbury, G. Chanan, T. Killeen, Z. Lin, N. Gimelshein, L. Antiga, et al. Pytorch: An imperative style, high-performance deep learning library. *Advances in neural information processing systems*, 32:8026–8037, 2019.
- [32] A. Renda, J. Frankle, and M. Carbin. Comparing rewinding and fine-tuning in neural network pruning. In *International Conference on Learning Representations*, 2019.
- [33] D. J. Rezende, S. Mohamed, and D. Wierstra. Stochastic backpropagation and approximate inference in deep generative models. In *International conference on machine learning*, pages 1278–1286. PMLR, 2014.
- [34] P. Savarese, H. Silva, and M. Maire. Winning the lottery with continuous sparsification. *Advances in Neural Information Processing Systems*, 33, 2020.
- [35] K. Simonyan and A. Zisserman. Very deep convolutional networks for large-scale image recognition. *International Conference on Learning Representations*, 2015.
- [36] M. Song, A. Montanari, and P. Nguyen. A mean field view of the landscape of two-layers neural networks. *Proceedings of the National Academy of Sciences*, 115:E7665–E7671, 2018.
- [37] S. Srinivas, A. Subramanya, and R. Venkatesh Babu. Training sparse neural networks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops*, pages 138–145, 2017.
- [38] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. u. Kaiser, and I. Polosukhin. Attention is all you need. In *Advances in Neural Information Processing Systems*, 2017.
- [39] C. Wang, G. Zhang, and R. Grosse. Picking winning tickets before training by preserving gradient flow. In *International Conference on Learning Representations*, 2019.
- [40] Z. Wang, J. Wohlwend, and T. Lei. Structured pruning of large language models. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 6151–6162, 2020.
- [41] W. Wen, C. Wu, Y. Wang, Y. Chen, and H. Li. Learning structured sparsity in deep neural networks. In *Advances in neural information processing systems*, pages 2074–2082, 2016.
- [42] X. Xiao, Z. Wang, and S. Rajasekaran. Autoprune: Automatic network pruning by regularizing auxiliary parameters. In *Advances in Neural Information Processing Systems*, pages 13681–13691, 2019.

- [43] M. Ye, C. Gong, L. Nie, D. Zhou, A. Klivans, and Q. Liu. Good subnetworks provably exist: Pruning via greedy forward selection. In *International Conference on Machine Learning*, pages 10820–10830. PMLR, 2020.
- [44] X. Yuan, P. H. P. Savarese, and M. Maire. Growing efficient deep networks by structured continuous sparsification. In *International Conference on Learning Representations*, 2020.
- [45] S. Zagoruyko and N. Komodakis. Wide residual networks. In *British Machine Vision Conference 2016*. British Machine Vision Association, 2016.
- [46] W. Zeng and R. Urtasun. MLPrune: Multi-layer pruning for automated neural network compression, 2019.
- [47] X. Zhou, W. Zhang, H. Xu, and T. Zhang. Effective sparsification of neural networks with global sparsity constraint. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 3599–3608, 2021.
- [48] M. Zhu and S. Gupta. To prune, or not to prune: exploring the efficacy of pruning for model compression. *arXiv preprint arXiv:1710.01878*, 2017.

Checklist

1. For all authors...
 - (a) Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope? [Yes]
 - (b) Did you describe the limitations of your work? [Yes]
 - (c) Did you discuss any potential negative societal impacts of your work? [No] This is a sparse training algorithm aiming to speed up deep neural network training process, with no potential negative societal impacts.
 - (d) Have you read the ethics review guidelines and ensured that your paper conforms to them? [Yes]
2. If you are including theoretical results...
 - (a) Did you state the full set of assumptions of all theoretical results? [Yes]
 - (b) Did you include complete proofs of all theoretical results? [Yes]
3. If you ran experiments...
 - (a) Did you include the code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL)? [Yes]
 - (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they were chosen)? [Yes]
 - (c) Did you report error bars (e.g., with respect to the random seed after running experiments multiple times)? [No]
 - (d) Did you include the total amount of compute and the type of resources used (e.g., type of GPUs, internal cluster, or cloud provider)? [Yes]
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...
 - (a) If your work uses existing assets, did you cite the creators? [Yes]
 - (b) Did you mention the license of the assets? [Yes]
 - (c) Did you include any new assets either in the supplemental material or as a URL? [No]
 - (d) Did you discuss whether and how consent was obtained from people whose data you're using/curating? [N/A]
 - (e) Did you discuss whether the data you are using/curating contains personally identifiable information or offensive content? [N/A]
5. If you used crowdsourcing or conducted research with human subjects...
 - (a) Did you include the full text of instructions given to participants and screenshots, if applicable? [N/A]
 - (b) Did you describe any potential participant risks, with links to Institutional Review Board (IRB) approvals, if applicable? [N/A]
 - (c) Did you include the estimated hourly wage paid to participants and the total amount spent on participant compensation? [N/A]