

# Optimally Deep Networks – Adapting Model Depth to Datasets for Superior Efficiency

Shaharyar Ahmed Khan Tareen  
University of Houston  
Houston, TX, USA  
stareen@cougarnet.uh.edu

Filza Khan Tareen  
National University of Sciences and Technology  
Islamabad, Pakistan  
ftareen.bse2021mcs@student.nust.edu.pk

**Abstract**—Deep neural networks (DNNs) have provided brilliant performance across various tasks. However, this success often comes at the cost of unnecessarily large model sizes, high computational demands, and substantial memory footprints. Typically, powerful architectures are trained at full depths but not all datasets or tasks require such high model capacity. Training big and deep architectures on relatively low-complexity datasets frequently leads to wasted computation, unnecessary energy consumption, and excessive memory usage, which in turn makes deployment of models on resource-constrained devices impractical. To address this problem, we introduce the concept of Optimally Deep Networks (ODNs), which provides a balance between model depth and task complexity. Specifically, we propose a NAS like training strategy called “progressive depth expansion”, which begins by training neural networks at shallower depths and incrementally increases their depth as the earlier blocks converge, continuing this process until the target accuracy is reached. ODNs use only the “optimal” depth for the tasks at hand, removing redundant layers. This cuts down future training and inference costs, lowers the model’s memory footprint, enhances computational efficiency, and facilitates deployment on edge devices. Empirical results show that the optimal depths of ResNet-18 and ResNet-34 for MNIST and SVHN, achieve up to 98.64 % and 96.44 % reduction in memory footprint, while maintaining a competitive accuracy of 99.31 % and 96.08 %, respectively.

**Index Terms**—*optimally deep networks, optimal depth, adaptive depth, progressive depth expansion, neural architecture search, model efficiency, efficient deep learning.*

## I. INTRODUCTION

Deep Neural Networks (DNNs) have become the cornerstone of modern machine learning, due to their brilliant performance across various domains [1]–[3]. However, these feats come at the expense of large model sizes due to deeper architectures and correspondingly higher computational and memory demands. Although large models have proven essential for complex datasets (ImageNet-1K, ImageNet-21K, large-scale NLP datasets etc.) [2], not all tasks have such high level of complexity. Many domains involve tasks that are significantly simple or have lower complexity such as some datasets in medical imaging [4], industrial inspection [5], digit or character classification [6], [7], and image matching [8]. Using deep architectures (ResNet-18, ResNet-34, ResNet-50, ResNet-101 etc.) in simple scenarios leads to inefficiencies, including wasted computation/time, increased latency during inference, unnecessary energy consumption, and excessive memory footprint, hindering their deployability on resource-constrained devices e.g. mobile phones and edge devices.

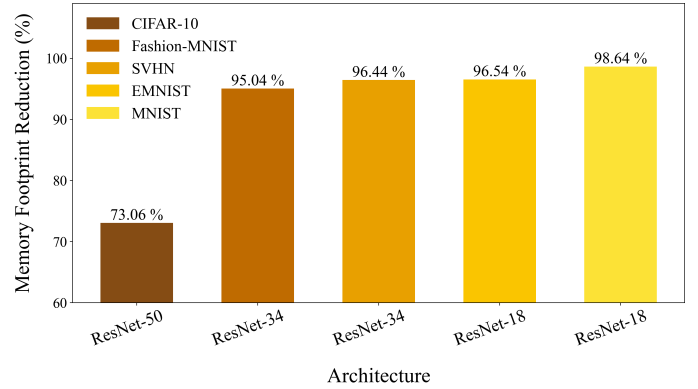


Fig. 1: Memory footprint reduction of networks at their optimal depths across five datasets (The accuracies of the ODNs remain within 1 % of the full-depth networks).

Many strategies for efficient deep learning [9], [10] have been proposed including pruning [11], quantization [12], knowledge distillation [13]–[17], dynamic inference [18], and Neural Architecture Search (NAS) [19]–[21]. Pruning and quantization compress the networks but cannot substantially reduce their size on disk. Knowledge distillation transfers knowledge from a large teacher model to a smaller student model, however, it does not provide the optimized architecture for the task at hand. Dynamic inference techniques adaptively skip layers in a model during inference to reduce the latency but fail to address the problem of large memory footprint. Neural Architecture Search (NAS) provides an architecture search process, however, it has huge computation cost leading to large carbon footprint [22], extreme time consumption (as it typically trains thousands of candidates to find the optimum architecture), and the complexity of defining the search space, which may sometimes miss powerful architectures.

We introduce the concept of Optimally Deep Networks (ODNs), which use a simple and smart NAS like training strategy that enables adapting the depth of a model to the complexity of the task at hand. Instead of directly training the model with full depth, “progressive depth expansion” is applied that begins by activating only the shallow layers and incrementally deepens the network once the earlier layers have converged, continuing this process until the desired performance is achieved. When the target performance is achieved

with a depth shallower than the original architecture, the Optimally Deep Network is extracted and fine-tuned. This depth-search strategy enables networks to allocate only the necessary depth capacity for the tasks, reducing substantial memory footprint, future training costs, computation overheads, and inference run-time, all without a significant compromise on accuracy. Once the optimal depth of an architecture is obtained, it can be used to efficiently deploy the model at a large scale, reducing resource wastage. The same depth can be reused for future trainings on datasets of similar complexity, avoiding the cost of performing another search. Moreover, the ODNs can be further compressed using pruning and quantization techniques.

Unlike Neural Architecture Search (NAS), progressive depth expansion is a simple depth search strategy which does not require defining a complex search space. It can be straightforwardly applied on many powerful architectures such as Residual Networks (ResNet-18, ResNet-24, ResNet-50, ResNet-101 etc.) [2], Wide Residual Networks (WRN-40-2, WRN-16-8, WRN-28-10 etc.) [23], Mobile Networks (MobileNetV1, MobileNetV2, MobileNetV3 etc.) [24], [25] and many more, by dividing them into different depth levels (blocks) and progressively training them from shallower to deeper blocks until the desired performance is achieved with the optimal depth. ODNs offer a new perspective on efficient deep learning by searching optimal depth of well-known networks according to the complexity of tasks, bridging the gap between performance and efficiency.

**Key Contributions:** Our contributions are highlighted below:

- We introduce Optimally Deep Networks (ODNs), in which “model depth” is adapted according to the complexity of task at hand, with competitive performance.
- We propose a novel training paradigm called “progressive depth expansion” that gradually deepens the model-depth during training to achieve the target accuracy.
- We show that ODNs significantly reduce memory footprint and inference costs, enabling efficient deployment on edge and resource-limited devices.
- We empirically demonstrate the effectiveness of ODNs on five datasets: MNIST, EMNIST, SVHN, Fashion-MNIST, and CIFAR-10, by optimizing the depths of three architectures: ResNet-18, ResNet-34, and ResNet-50.
- ResNet-18 trained through progressive depth expansion on MNIST provides **98.64 % reduction** in model size (from 44.78 MB to 0.61 MB) while preserving competitive performance with 99.39 % accuracy. Similarly, ResNet-34 trained on SVHN provides **96.44 % reduction** in model size (from 85.29 MB to 3.04 MB) with a competitive accuracy of 96.08 %.
- The code of ODNs is released at the following link: [https://github.com/saktx/Optimally\\_Deep\\_Networks](https://github.com/saktx/Optimally_Deep_Networks).

## II. RELATED WORK

The pursuit of greater efficiency in deep learning inspired researchers to introduce various approaches aimed at cutting

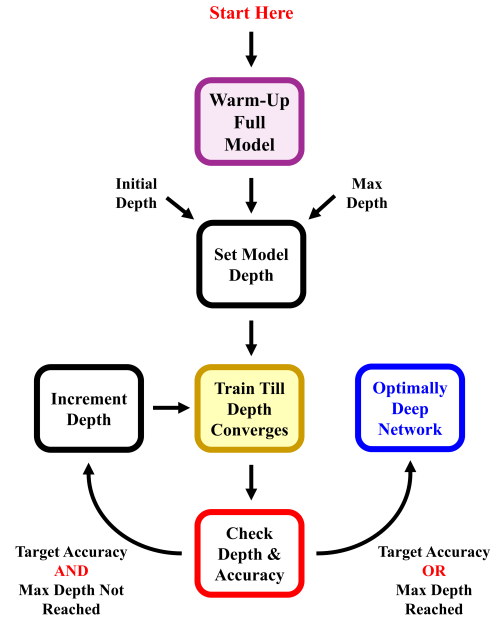


Fig. 2: Procedure for obtaining Optimally Deep Networks.

down the computational costs, memory footprint, inference latency and other deployment barriers of the models, while retaining their performance [9], [10]. Most popular research directions related to efficient deep learning include Neural Architecture Search (NAS), model compression techniques (sparsification, pruning, and quantization), knowledge distillation, dynamic inference, and Once-For-All (OFA) training.

### A. Neural Architecture Search (NAS)

NAS automated the process of finding a neural network that is optimized with respect to accuracy, robustness, and/or efficiency by exploring a vast search space of possible architectures without requiring extensive evaluation of trial-and-error manually [19]–[21]. NAS has been able to provide high performing architectures by leveraging reinforcement learning such as RENAS [26], gradient based optimization like DARTS [27], and evolutionary search algorithms e.g. NPENAS [28]. However, the search process of NAS can be computationally intensive, requiring up to thousands of GPU hours and since NAS generated architectures are not universally optimal across all the datasets, a new architecture search is required for each task, undermining the scalability [19]. This limitation makes NAS extremely costly for such scenarios.

### B. Model Compression Techniques

Sparsification, pruning, and quantization are widely used techniques for compressing large neural networks [11], [12]. Sparsification encourages the model to use as few parameters as possible by pushing their magnitudes toward zero during training. Subsequently, pruning is applied to remove these zeroed-out weights or neurons. Quantization on the other hand reduces the numerical precision of model parameters to decrease its memory footprint [12]. Although these compression methods reduce the number of parameters or their numerical

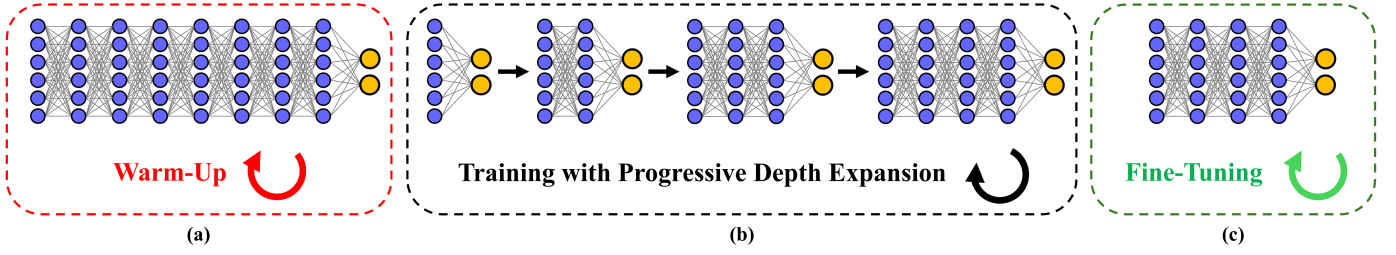


Fig. 3: Training and depth search process of Optimally Deep Networks (ODNs): (a) The selected network is first warmed up by training for a few epochs with a small learning rate for all the depth levels. (b) Progressive depth expansion is applied during training, starting with a shallow depth and incrementally adding deeper blocks once the shallower ones have fully converged. (c) The network searched with the optimal depth is then fine-tuned to maximize its performance.

precision, they do not provide substantial reduction in memory footprint, particularly in the case of unstructured sparsification. Moreover, pruning and quantization are typically applied post training, meaning the model is first trained at full depth before compression, still leading to larger memory footprints because the model’s structure remains unchanged.

### C. Once-for-All Training Frameworks

Once-for-All (OFA) training frameworks aim to train a large overparametrized super-net only once, that can provide multiple optimized sub-networks of varying sizes later, facilitating flexible deployment across diverse hardware platforms [22], [29], [30]. Following this approach, SNNs [31], [32] train a super-net for different widths, enabling accuracy and efficiency trade-off. Although OFA training provides flexibility and avoids retraining for each device, it introduces significant training complexity and does not perform well across different datasets. Moreover, the super-net parameters are shared among all the sub-nets that leads to challenges in optimization, high training costs, and reduced specialization of the sub-nets compared to individually trained models [22], [31].

While all these methods have advanced the domain of efficient deep learning, each one suffers from its limitations which highlight the need for an alternative and simple paradigm that directly adapts model size to the complexity of the task at hand. Optimally Deep Networks (ODNs) tackle this gap by progressively expanding model depth in a data-driven manner, cutting down the unnecessary memory footprint while maintaining competitive accuracy.

## III. OPTIMALLY DEEP NETWORKS (ODNs)

We propose **Optimally Deep Networks (ODNs)**, which adapt model-depth to the complexity of the dataset via “progressive depth expansion”. Instead of training the model with full-depth, we first divide its depth in partitions and assign an output layer for each depth. Next, we warm up the model by training it for few epochs at all the depth levels, with a small learning rate. This stabilizes the optimization process for all depth partitions of the model. The warmed-up model parameters are saved for assisting in the search of optimal depth later. Afterwards, the user specifies a “target

accuracy” according to the desired performance requirements. The progressive depth expansion based training starts by first activating a shallow portion of the network that only contains its initial block(s). Once this shallow portion is converged, additional blocks are appended, and training is restarted using the warmed-up parameters (till this new depth of the network) that were saved after the warm-up. This process continues until the model achieves the target accuracy or maximum depth is reached. If the target accuracy is obtained with a depth smaller than the model’s original depth, the current depth is returned as the “**optimal depth**” for which the ODN is extracted and fine-tuned. This dynamic depth search approach ensures that depth capacity of the network is allocated for the dataset as needed, preventing over-parameterization and reducing the model’s memory footprint without compromising accuracy.

Algorithm 1 outlines the procedure of training Optimally Deep Networks and the workflow to obtain ODNs is shown in Figure 2. In the following subsections, we describe the key phases involved in obtaining an Optimally Deep Network.

### A. Depth Partitioning & Multiple Output Layers Setup

To train a model at different depths, its depth is partitioned into different levels. When a depth is activated, the data passes only through the activated depth to provide the output. We divide the depth of ResNet-18, ResNet-34, and ResNet-50 into 8, 16, and 16 blocks, respectively. The partitions are based on number of residual blocks. An activated depth level of 4 means that only the first 4 blocks of the model are active. Each block in ResNet-18 and ResNet-34 contains two convolutional layers each followed by a BatchNorm layer, whereas, ResNet-50 blocks contain three convolutional layers (each followed by a BatchNorm layer). A separate output layer is employed in the network architecture for each depth level (to facilitate data flow), meaning for 16 depth levels there will be 16 output layers, one layer for each depth. Formally, let a neural network  $\mathcal{N}$  be parameterized by depth  $\mathcal{D}$  as in Eq. (1).

$$\mathcal{D} \in \{1, 2, 3, \dots, \mathcal{D}_{max}\} \quad (1)$$

where  $\mathcal{D}_{max}$  denotes the maximum depth (total number of blocks) in the network.

TABLE I: Comparison of Model Depth, Parameter Count, Memory Footprint (Float32 Precision), FLOPs, and Performance of Standard Networks with Full Depths versus Corresponding Optimally Deep Networks (ODNs). Reduced Parameter Counts Show Significant Reduction in Memory Footprints whereas Reduced FLOPs Indicate Improvements in Inference Speed.

| DATASET       | Architecture               | Depth | Parameters ↓ | Size on Disk ↓  | % Size Reduction ↑ | FLOPs ↓   | Accuracy ↑ |
|---------------|----------------------------|-------|--------------|-----------------|--------------------|-----------|------------|
| MNIST         | ResNet-18                  | 8/8   | 11.17 M      | 44.78 MB        | 98.64 %            | 458.21 M  | 99.61 %    |
|               | ResNet-18 <sub>2/8</sub>   | 2/8   | 0.15 M       | <b>0.61 MB</b>  |                    | 117.36 M  | 99.31 %    |
| EMNIST        | ResNet-18                  | 8/8   | 11.17 M      | 44.78 MB        | 96.54 %            | 458.21 M  | 95.17 %    |
|               | ResNet-18 <sub>3/8</sub>   | 3/8   | 0.38 M       | <b>1.55 MB</b>  |                    | 162.37 M  | 95.04 %    |
| SVHN          | ResNet-34                  | 16/16 | 21.28 M      | 85.29 MB        | 96.44 %            | 939.91 M  | 96.51 %    |
|               | ResNet-34 <sub>5/16</sub>  | 5/16  | 0.75 M       | <b>3.04 MB</b>  |                    | 365.63 M  | 96.08 %    |
| Fashion-MNIST | ResNet-34                  | 16/16 | 21.28 M      | 85.29 MB        | 95.04 %            | 939.91 M  | 94.52 %    |
|               | ResNet-34 <sub>6/16</sub>  | 6/16  | 1.04 M       | <b>4.23 MB</b>  |                    | 337.08 M  | 93.54 %    |
| CIFAR-10      | ResNet-50                  | 16/16 | 23.52 M      | 94.43 MB        | 73.06 %            | 1315.02 M | 95.01 %    |
|               | ResNet-50 <sub>11/16</sub> | 11/16 | 6.31 M       | <b>25.44 MB</b> |                    | 906.19 M  | 94.12 %    |

### B. Warm-Up with Full Depth

After depth partitioning and assignment of separate output layers for each depth, the model is warmed up at all the depths for a few epochs with a small learning rate, see Figure 3(a). This gives all depth levels good initialization points, stabilizing the progressive depth expansion based training, preventing the parameters of newly activated deeper blocks from random initialization that leads to jumps in the loss curve, gradient instability, and slower convergence as shown in Figure 4.

### C. Training with Progressive Depth Expansion

Progressive depth expansion based training starts from the shallow portion of the network that comprises initial block(s) up to a depth level. Once the shallow block(s) have converged, next block is appended to increase the model depth and training is resumed. At each depth increment, training is started from the warmed-up state of the model instead of starting from scratch. This progressive depth expansion continues until either the target accuracy is achieved or the model reaches its maximum depth. If the target accuracy is achieved using a depth smaller than the full model depth, the current depth is declared as the optimal depth, and the ODN is extracted with the corresponding output layer for fine-tuning as shown in Figure 3(b). The training objective with model depth “ $\mathcal{D}$ ” is to minimize a standard supervised learning loss function:

$$\mathcal{L}(\Theta_{\mathcal{D}}) = \frac{1}{N} \sum_{i=1}^N \mathcal{L}(\mathcal{N}(x_i; \Theta_{\mathcal{D}}), y_i) \quad (2)$$

where  $\mathcal{L}$  is cross entropy loss,  $\Theta_{\mathcal{D}}$  represents the parameters of the network  $\mathcal{N}$  activated till depth  $\mathcal{D}$ ,  $N$  is the number of training samples in a mini-batch,  $x_i$  represents input data and  $y_i$  denotes corresponding ground-truth labels. Traditionally, a model is always trained using depth  $\mathcal{D} = \mathcal{D}_{max}$ , regardless of the dataset’s complexity. The goal of ODNs is to optimally adapt the depth  $\mathcal{D}_{opt} \leq \mathcal{D}_{max}$  of the model for the task at hand, bridging the gap between performance and efficiency.

**Stopping Criterion:** Each depth is trained until convergence and the convergence is determined using a stopping

criterion. During training or fine-tuning, the learning rate is reduced by 60 % after every 5 consecutive epochs with “no-improvement” on the validation set. Whenever a new best validation accuracy is achieved, the counter is reset to zero. The model is considered converged when the “no-improvement” epochs reach the value of 23, as shown in Algorithm 1.

### D. Fine-tuning the Optimally Deep Network

Once the Optimally Deep Network is obtained, it is fine-tuned till convergence to maximize its performance as shown in Figure 3(c). The fine-tuning stage is the final step in training Optimally Deep Networks (ODNs) to ensure the selected optimal depth achieves its best possible performance. Fine-tuning allows the model parameters with optimal depth to adjust more precisely to the target dataset, further improving accuracy (without altering the depth). The ODN can be extracted from the backbone at the end (by copying weights) for deployment.

## IV. EXPERIMENTS & RESULTS

We evaluate Optimally Deep Networks (ODNs) across five benchmark datasets of varying complexity: MNIST, EMNIST, SVHN, Fashion-MNIST, and CIFAR-10. All the datasets have 10 classes except EMNIST, which has 26 classes (for “letters”). To demonstrate the simplicity and applicability of our approach, we use three widely used architectures: ResNet-18, ResNet-34, and ResNet-50. We train the models using SGD optimizer with momentum 0.9, weight decay 5e-4, batch size 128, and initial learning rate of 0.1.

Table I presents the performance of ODNs as compared to networks with full depths. It can be observed that the optimally deep ResNet-18 trained on MNIST and EMNIST achieves a comparable accuracy of 99.31 % and 95.04 % while reducing the model size (memory footprint) by 98.64 % and 96.54 %, respectively. Only 2 out of the 8 blocks of ResNet-18 are sufficient to match the complexity of the MNIST dataset whereas only 3 blocks are sufficient for EMNIST. Similarly, the optimally deep ResNet-34 achieves accuracies of 96.08 % and 93.54 %, with model size reductions of 96.44 % and 95.04 % on SVHN and Fashion-MNIST, respectively. For the

CIFAR-10 dataset, the optimally deep ResNet-50 with a depth of 11 blocks instead of 16, provides the accuracy of 94.12 % with a reduction of 73.06 % in model size. Compared to full-depth models, the accuracies of ODNs can be achieved within **1% tolerance**, while the memory footprint reductions are substantial, **exceeding 95 %** for low-complexity datasets: MNIST, EMNIST, SVHN, and Fashion-MNIST.

Table I also shows that the substantial reduction in parameter count and model size of ODNs leads to a decrease in FLOPs. Lower FLOPs result in faster inference times and reduced latency in edge devices. For MNIST and EMNIST, the memory footprint is reduced from 44.78 MB to 0.61 MB

---

**Algorithm 1** : Training Optimally Deep Networks (ODNs)

---

**Require:** training dataset  $\mathcal{S}$ , neural network  $\mathcal{N}$ , depth levels  $\mathcal{D} = \{1, 2, 3, \dots, \mathcal{D}_{\max}\}$ , warm-up epochs  $E_w$ , stopping epochs  $E_{stop}$ , warm-up learning rate  $\alpha_w$ , learning rate  $\alpha$ , initial depth  $\mathcal{D}_i$ , final depth  $\mathcal{D}_f$ , target accuracy  $\mathcal{A}_{tar}$

**Ensure:** trained network  $\mathcal{N}_{\Theta}$  with optimal depth  $\mathcal{D}_{opt}$

```

1: Activate  $\mathcal{N}$  with depth  $\mathcal{D}_{\max}$ 
2: Initialize network parameters  $\Theta$ 
3: Initialize optimizer with  $\Theta$  and  $\alpha_w$ 
4: for  $e = 1$  to  $E_w$  do
5:   Clear gradients: optimizer.zero_grad()
6:   Sample batch  $(x, y)$  from  $\mathcal{S}$ 
7:   Compute loss:  $\mathcal{L}(y, \mathcal{N}(x))$ 
8:   Backpropagate:  $\mathcal{L.backward}()$ 
9:   Update parameters: optimizer.step()
10: Save: Save parameters of  $\mathcal{N}$  as  $\Theta_w$ 
11: Activate  $\mathcal{N}$  with depth  $\mathcal{D}_i$ 
12: Set learning rate  $\alpha$ 
13: Initialize  $A_{best} \leftarrow 0$ ,  $A_{current} \leftarrow 0$ ,  $E_{current} \leftarrow 0$ ,
    $E_{stop} \leftarrow 23$ ,  $\mathcal{D}_{current} \leftarrow \mathcal{D}_i$ 
14: while  $A_{current} \leq A_{target}$  and  $\mathcal{D}_{current} \leq \mathcal{D}_f$  do
15:    $E_{current} += 1$ 
16:   for each batch  $(x, y)$  in  $\mathcal{S}$  do
17:     Clear gradients: optimizer.zero_grad()
18:     Sample batch  $(x, y)$  from  $\mathcal{S}$ 
19:     Compute loss:  $\mathcal{L}(y, \mathcal{N}(x))$ 
20:     Backpropagate:  $\mathcal{L.backward}()$ 
21:     Update parameters: optimizer.step()
22:     Evaluate  $\mathcal{N}$  till  $\mathcal{D}_{current}$  to get  $A_{current}$ 
23:   if  $A_{current} > A_{best}$  then
24:     Save: Save best model  $\mathcal{N}_{\Theta}$ 
25:      $A_{best} = A_{current}$ 
26:      $E_{current} = 0$ 
27:   if  $E_{current} \bmod 5 = 0$  then
28:      $\alpha \leftarrow 0.6 \times \alpha$ 
29:   if  $E_{current} = E_{stop}$  then
30:      $\mathcal{D}_{current} += 1$ 
31:   Load: Load parameters of  $\mathcal{N}$  from  $\Theta_w$ 
32:  $\mathcal{D}_{opt} \leftarrow \mathcal{D}_{current}$ 
33: return  $\mathcal{N}_{\Theta}$ ,  $\mathcal{D}_{opt}$ 

```

---

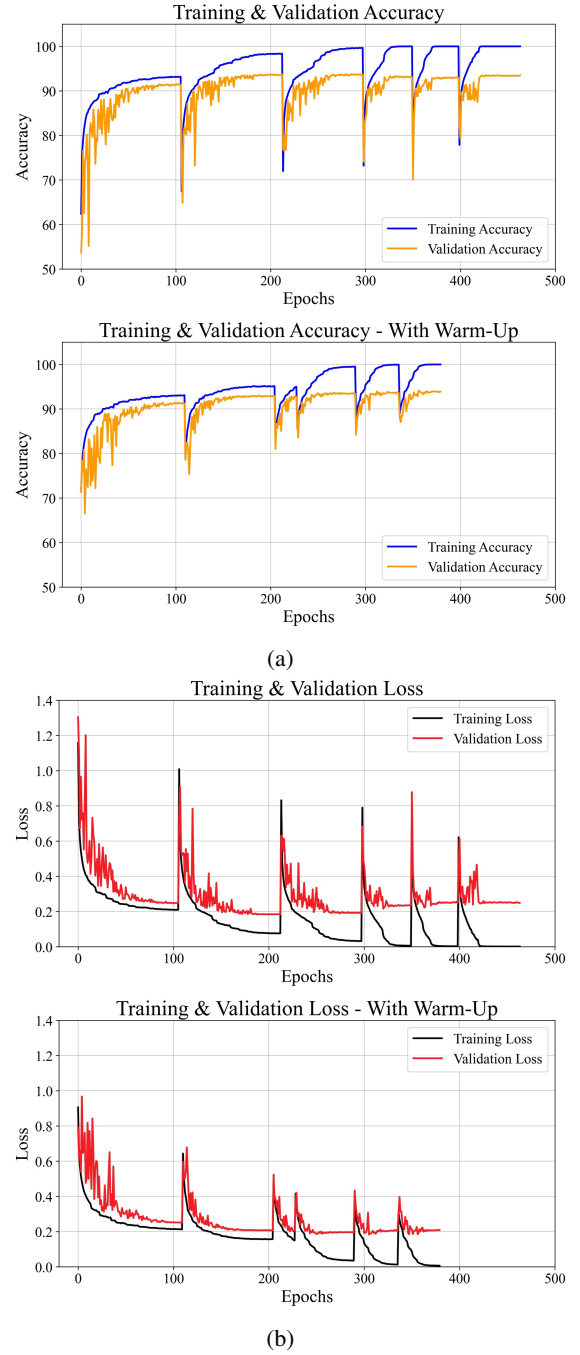


Fig. 4: The accuracy and loss curves of ResNet-34 for Fashion-MNIST show that use of warm-up mitigates large fluctuations and accelerates convergence during training.

and 1.55 MB, respectively. Similarly, for SVHN and Fashion-MNIST, the memory footprint decreases from 85.29 MB to 3.04 MB and 4.23 MB, respectively. In the case of CIFAR-10, the memory footprint is reduced by 73.06 %, from 94.43 MB to 25.44 MB, reflecting the higher complexity of this dataset.

ODNs eliminate the need for an expensive and complex search space as compared to complex NAS methods. Unlike sparsification and pruning, ODNs achieve a substantial

reduction in the memory footprints of the models. In contrast to Once-for-All (OFA) training, ODNs are simpler, requiring only a single progressively expanding model without the complexity of maintaining a super-net [22]. Another advantage of ODNs is that in order to find the optimal depth, the model is trained till convergence using shallower depths during progressive depth expansion stage. These fully converged intermediate models can be saved as by-products at each depth and later used for deployment on devices (like OFA) when stricter accuracy-efficiency trade-offs are required.

## CONCLUSION

We propose Optimally Deep Networks (ODNs), models which adapt their depth according to the complexity of the task at hand. Typically, powerful architectures are trained at full depths but we show that not all datasets require high model capacity. ODNs are trained using progressive depth expansion: training begins with shallow depths and the network is incrementally deepened only as needed, until the desired performance is achieved. This enables the extraction of Optimally Deep Networks that significantly reduce memory footprint, computational overhead, and inference cost, while maintaining competitive accuracy. Empirical evaluations on MNIST, EMNIST, SVHN, Fashion-MNIST, and CIFAR-10 using ResNet-18, ResNet-34, and ResNet-50 show that ODNs can achieve up to 98.64 % reduction in model size while maintaining good performance. ODNs present a practical framework for searching optimally deep networks that balance resource efficiency with accuracy, paving the way for scalable and deployable neural networks in real-world applications.

## REFERENCES

- [1] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *nature*, vol. 521, no. 7553, pp. 436–444, 2015.
- [2] 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*, 2016, pp. 770–778.
- [3] R. Mikkilainen, J. Liang, E. Meyerson, A. Rawal, D. Fink, O. Francon, B. Raju, H. Shahrzad, A. Navruzian, N. Duffy *et al.*, "Evolving deep neural networks," in *Artificial intelligence in the age of neural networks and brain computing*. Elsevier, 2024, pp. 269–287.
- [4] C. Affonso, A. L. D. Rossi, F. H. A. Vieira, A. C. P. de Leon Ferreira *et al.*, "Deep learning for biological image classification," *Expert systems with applications*, vol. 85, pp. 114–122, 2017.
- [5] M. S. Hossain, M. Al-Hammadi, and G. Muhammad, "Automatic fruit classification using deep learning for industrial applications," *IEEE transactions on industrial informatics*, vol. 15, no. 2, pp. 1027–1034, 2018.
- [6] S. Minaee, N. Kalchbrenner, E. Cambria, N. Nikzad, M. Chenaghlu, and J. Gao, "Deep learning-based text classification: a comprehensive review," *ACM computing surveys (CSUR)*, vol. 54, no. 3, pp. 1–40, 2021.
- [7] S. S. Kadam, A. C. Adamuthe, and A. B. Patil, "Cnn model for image classification on mnist and fashion-mnist dataset," *Journal of scientific research*, vol. 64, no. 2, pp. 374–384, 2020.
- [8] G. Koch, R. Zemel, R. Salakhutdinov *et al.*, "Siamese neural networks for one-shot image recognition," in *ICML deep learning workshop*, vol. 2, no. 1. Lille, 2015, pp. 1–30.
- [9] G. Huang, D. Chen, T. Li, F. Wu, L. Van Der Maaten, and K. Q. Weinberger, "Multi-scale dense networks for resource efficient image classification," *arXiv preprint arXiv:1703.09844*, 2017.
- [10] J. Pan, S. Yang, L. G. Foo, Q. Ke, H. Rahmani, Z. Fan, and J. Liu, "Progressive channel-shrinking network," *IEEE Transactions on Multimedia*, vol. 26, pp. 2016–2026, 2023.
- [11] S. Han, J. Pool, J. Tran, and W. Dally, "Learning both weights and connections for efficient neural network," *Advances in neural information processing systems*, vol. 28, 2015.
- [12] I. Hubara, M. Courbariaux, D. Soudry, R. El-Yaniv, and Y. Bengio, "Quantized neural networks: Training neural networks with low precision weights and activations," *journal of machine learning research*, vol. 18, no. 187, pp. 1–30, 2018.
- [13] G. Hinton, O. Vinyals, and J. Dean, "Distilling the knowledge in a neural network," *arXiv preprint arXiv:1503.02531*, 2015.
- [14] M. Goldblum, L. Fowl, S. Feizi, and T. Goldstein, "Adversarially robust distillation," in *Proceedings of the AAAI conference on artificial intelligence*, vol. 34, no. 04, 2020, pp. 3996–4003.
- [15] J. Gou, B. Yu, S. J. Maybank, and D. Tao, "Knowledge distillation: A survey," *International Journal of Computer Vision*, vol. 129, no. 6, pp. 1789–1819, 2021.
- [16] N. Papernot, P. McDaniel, X. Wu, S. Jha, and A. Swami, "Distillation as a Defense to Adversarial Perturbations against Deep Neural Networks," Mar. 2016, arXiv:1511.04508 [cs, stat]. [Online]. Available: <http://arxiv.org/abs/1511.04508>
- [17] H. Zhang, D. Chen, and C. Wang, "Confidence-aware multi-teacher knowledge distillation," in *ICASSP 2022-2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2022, pp. 4498–4502.
- [18] X. Wang, F. Yu, Z.-Y. Dou, T. Darrell, and J. E. Gonzalez, "Skipnet: Learning dynamic routing in convolutional networks," in *Proceedings of the European conference on computer vision (ECCV)*, 2018, pp. 409–424.
- [19] T. Elsken, J. H. Metzen, and F. Hutter, "Neural architecture search: A survey," *Journal of Machine Learning Research*, vol. 20, no. 55, pp. 1–21, 2019.
- [20] C. Liu, B. Zoph, M. Neumann, J. Shlens, W. Hua, L.-J. Li, L. Fei-Fei, A. Yuille, J. Huang, and K. Murphy, "Progressive neural architecture search," in *Proceedings of the European conference on computer vision (ECCV)*, 2018, pp. 19–34.
- [21] H. Pham, M. Guan, B. Zoph, Q. Le, and J. Dean, "Efficient neural architecture search via parameters sharing," in *International conference on machine learning*. PMLR, 2018, pp. 4095–4104.
- [22] H. Cai, C. Gan, T. Wang, Z. Zhang, and S. Han, "Once-for-All: Train One Network and Specialize it for Efficient Deployment," Apr. 2020, arXiv:1908.09791 [cs, stat]. [Online]. Available: <http://arxiv.org/abs/1908.09791>
- [23] S. Zagoruyko and N. Komodakis, "Wide residual networks," *arXiv preprint arXiv:1605.07146*, 2016.
- [24] 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.
- [25] M. Sandler, A. Howard, M. Zhu, A. Zhmoginov, and L.-C. Chen, "Mobilenetv2: Inverted residuals and linear bottlenecks," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2018, pp. 4510–4520.
- [26] Y. Chen, G. Meng, Q. Zhang, S. Xiang, C. Huang, L. Mu, and X. Wang, "Renas: Reinforced evolutionary neural architecture search," in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2019, pp. 4787–4796.
- [27] H. Liu, K. Simonyan, and Y. Yang, "Darts: Differentiable architecture search," *arXiv preprint arXiv:1806.09055*, 2018.
- [28] C. Wei, C. Niu, Y. Tang, Y. Wang, H. Hu, and J. Liang, "Npenas: Neural predictor guided evolution for neural architecture search," *IEEE transactions on neural networks and learning systems*, vol. 34, no. 11, pp. 8441–8455, 2022.
- [29] H. Wang, T. Chen, S. Gui, T.-K. Hu, J. Liu, and Z. Wang, "Once-for-All Adversarial Training: In-Situ Tradeoff between Robustness and Accuracy for Free," Nov. 2020, arXiv:2010.11828 [cs]. [Online]. Available: <http://arxiv.org/abs/2010.11828>
- [30] T. Chen, B. Ji, T. Ding, B. Fang, G. Wang, Z. Zhu, L. Liang, Y. Shi, S. Yi, and X. Tu, "Only train once: A one-shot neural network training and pruning framework," *Advances in Neural Information Processing Systems*, vol. 34, pp. 19637–19651, 2021.
- [31] J. Yu, L. Yang, N. Xu, J. Yang, and T. Huang, "Slimmable neural networks," *arXiv preprint arXiv:1812.08928*, 2018.
- [32] J. Yu and T. S. Huang, "Universally slimmable networks and improved training techniques," in *Proceedings of the IEEE/CVF international conference on computer vision*, 2019, pp. 1803–1811.