

Efficient & Scalable Deep Learning and Beyond

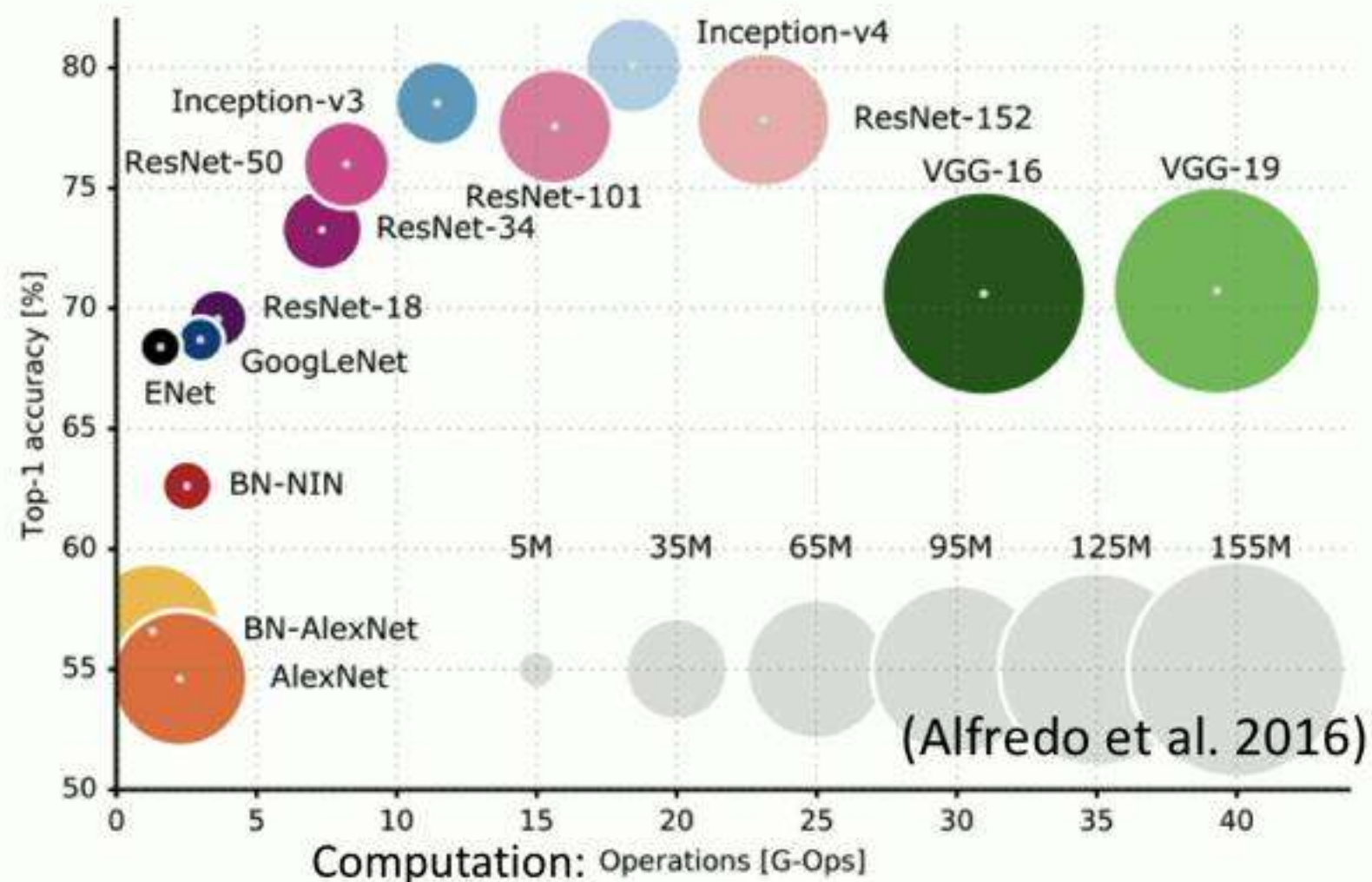
Wei Wen

Duke University

<http://www.pittnuts.com/>

Deep Learning Model Scaling

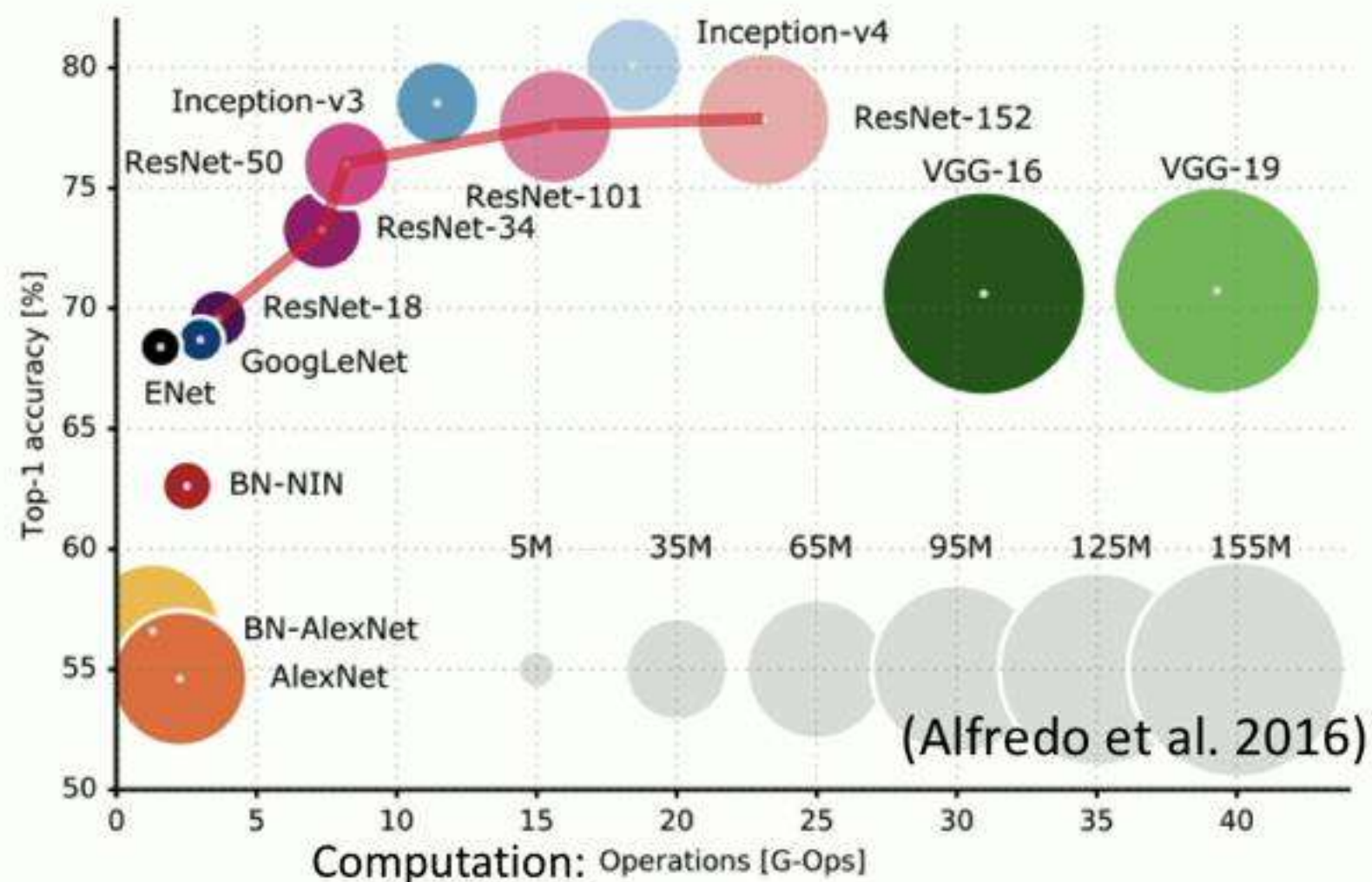
Image Classification on ImageNet



Trend: we can obtain better prediction using larger models

Deep Learning Model Scaling

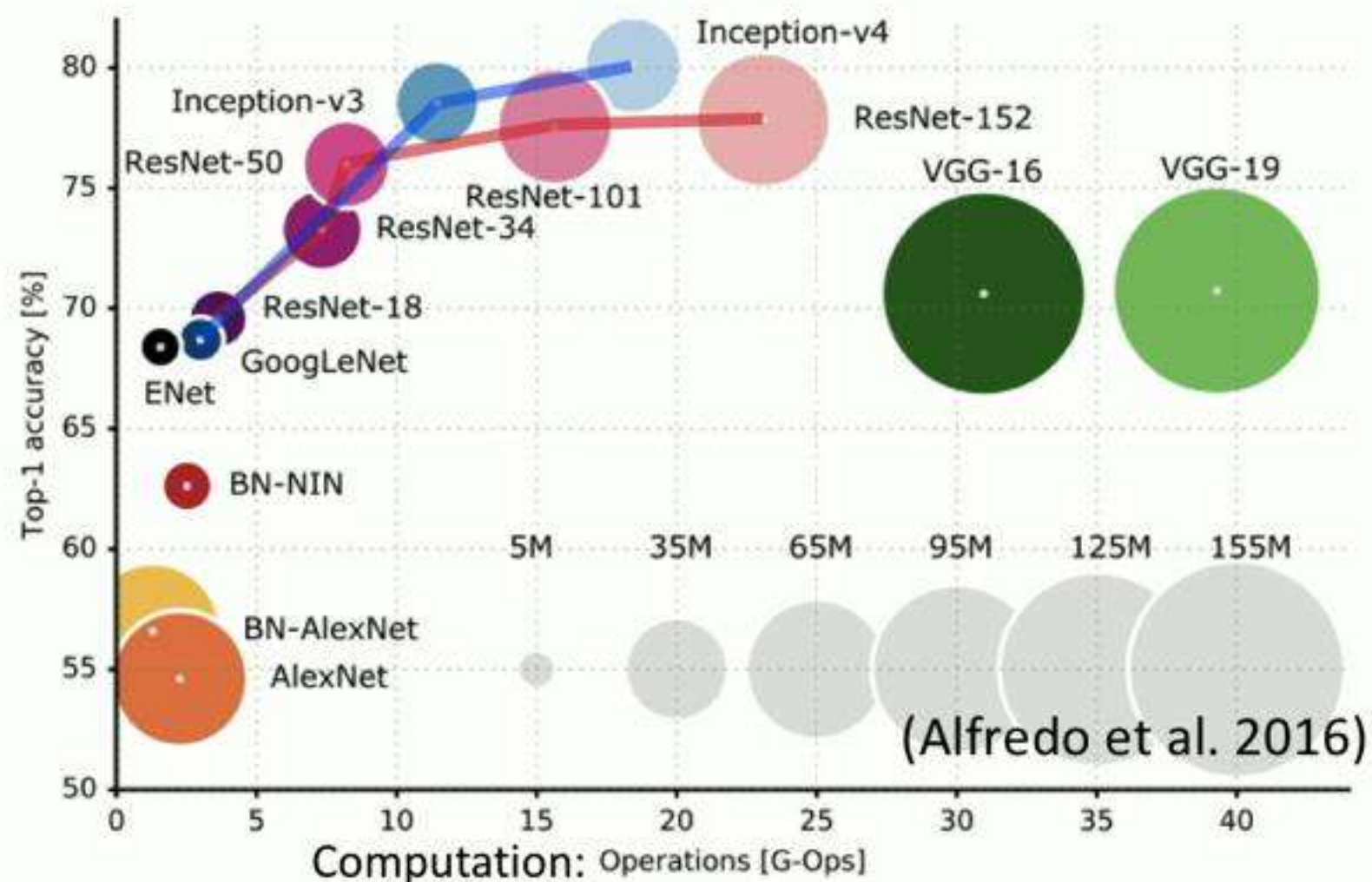
Image Classification on ImageNet



Trend: we can obtain better prediction using larger models

Deep Learning Model Scaling

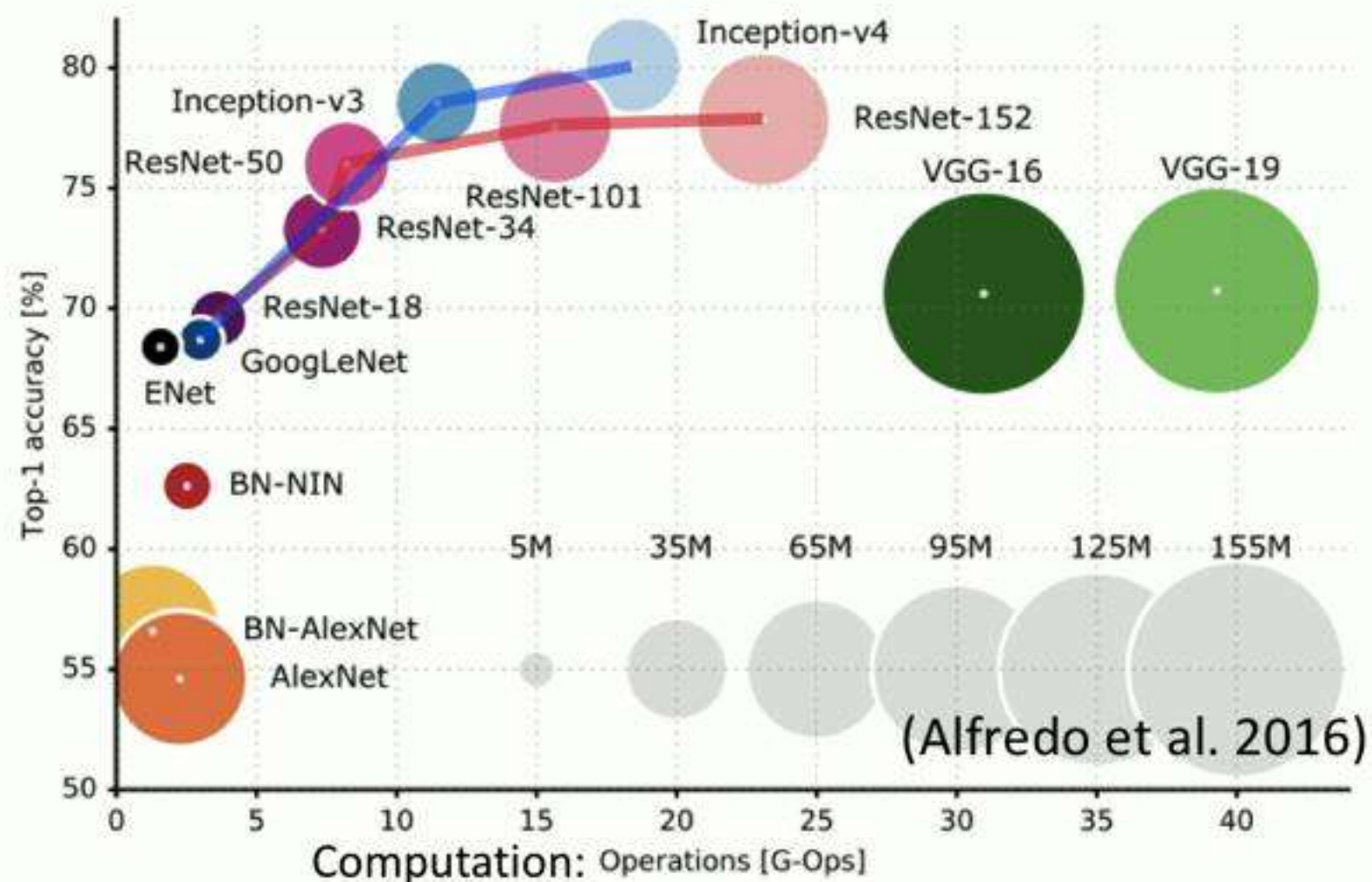
Image Classification on ImageNet



Trend: we can obtain better prediction using larger models

Deep Learning Model Scaling

Image Classification on ImageNet

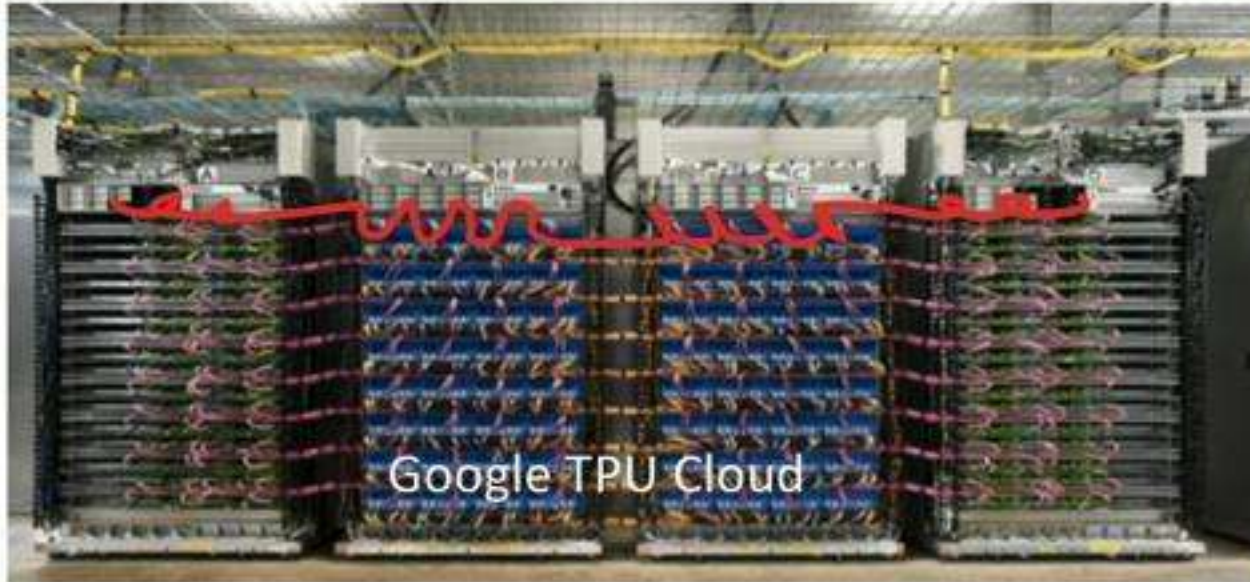


Language Modeling on Wikitext-103



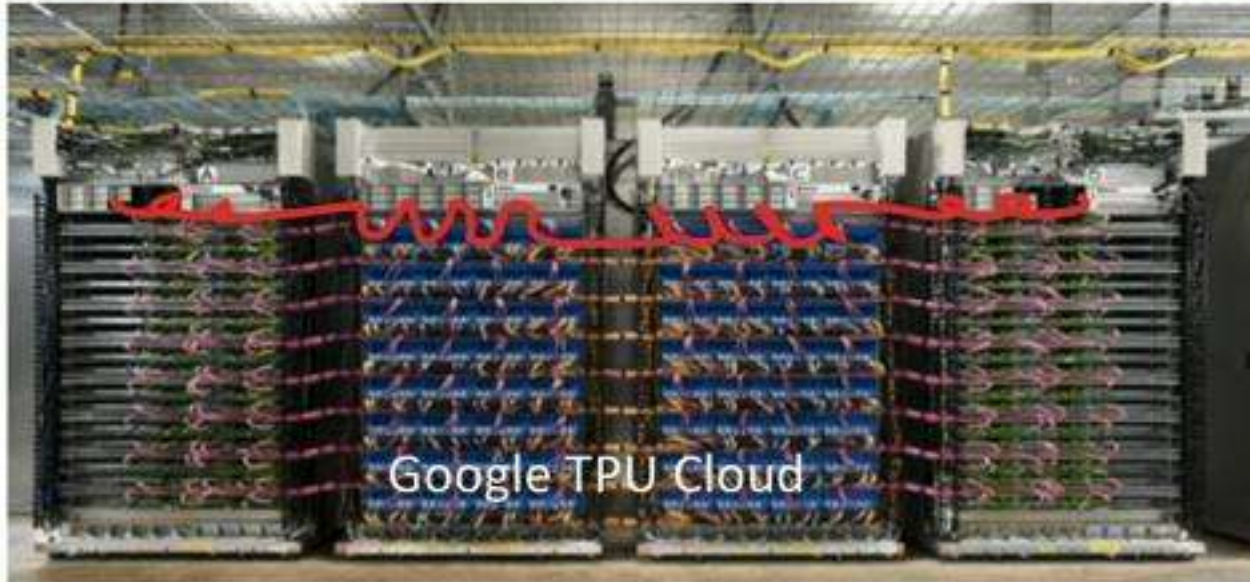
Trend: we can obtain better prediction using larger models

Obstacles of Building Larger Models



- Training larger models is slower
- Longer research & production cycles

Obstacles of Building Larger Models

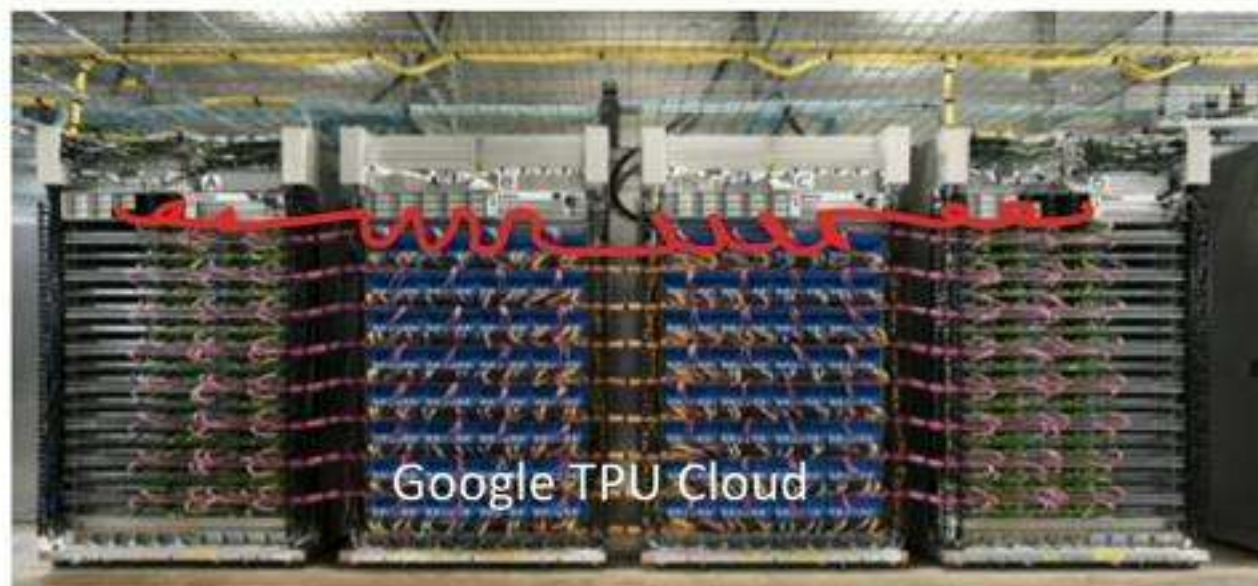


- Training larger models is slower
- Longer research & production cycles



- Inference speed is slow
- Challenging to deploy large models to applications with compute and memory constraints

Obstacles of Building Larger Models



- Training larger models is slower
- Longer research & production cycles
- Inference speed is slow
- Challenging to deploy large models to applications with compute and memory constraints

My research: make training and inference faster.

Outline

- Previous Research
 - Distributed deep learning acceleration (15 min)
 - Sparse deep neural networks (20 min)
- Future Research (6 min)

***TernGrad*: Ternary Gradients to Reduce Communication in Distributed Deep Learning**

NeurIPS 2017 (Oral)

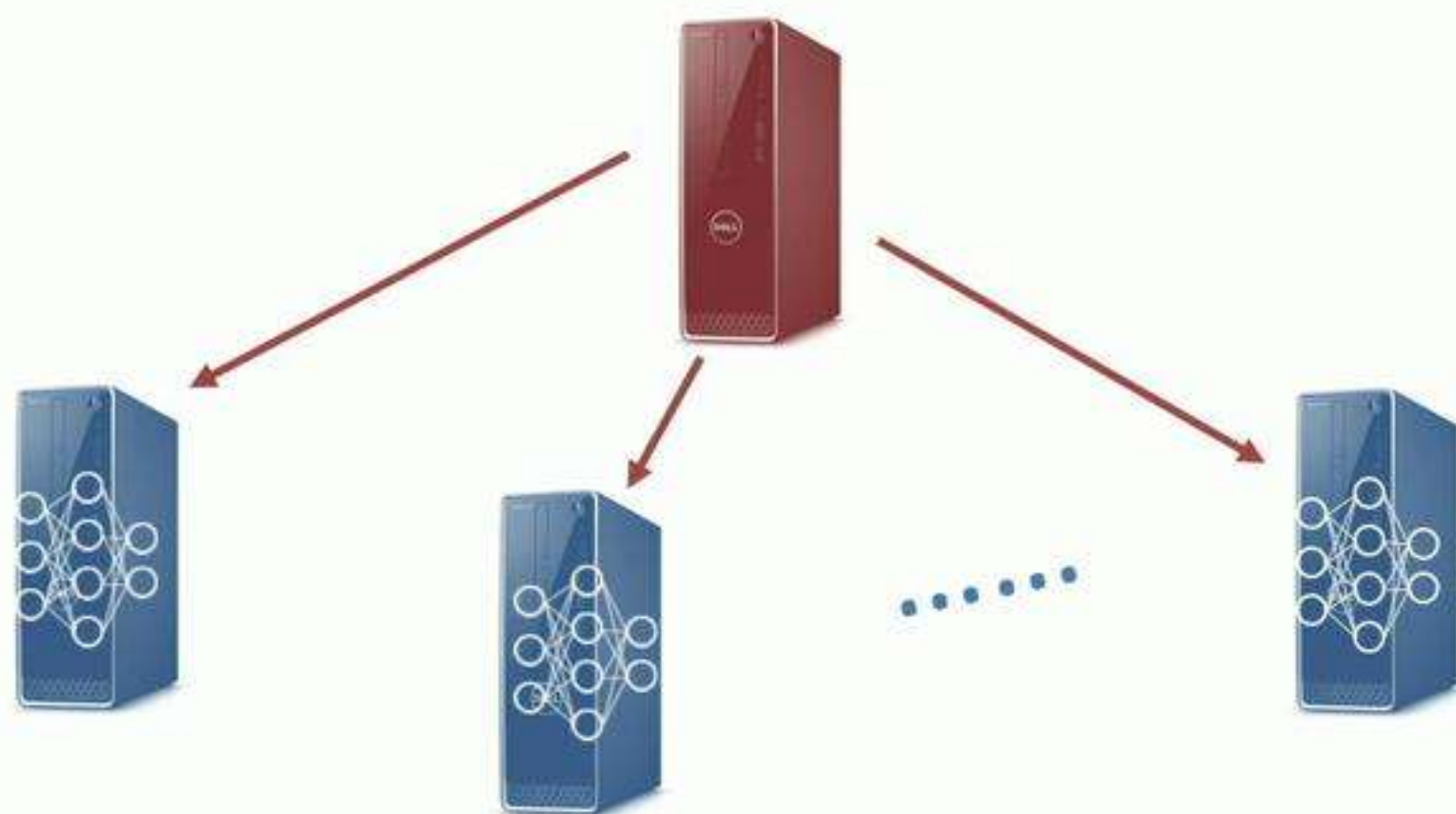
Wen, Wei, Cong Xu, Feng Yan, Chunpeng Wu, Yandan Wang, Yiran Chen, and Hai Li. "Terngrad: Ternary gradients to reduce communication in distributed deep learning." In *Advances in neural information processing systems*, pp. 1509-1519. 2017.

Communication Bottleneck in Distributed Deep Learning



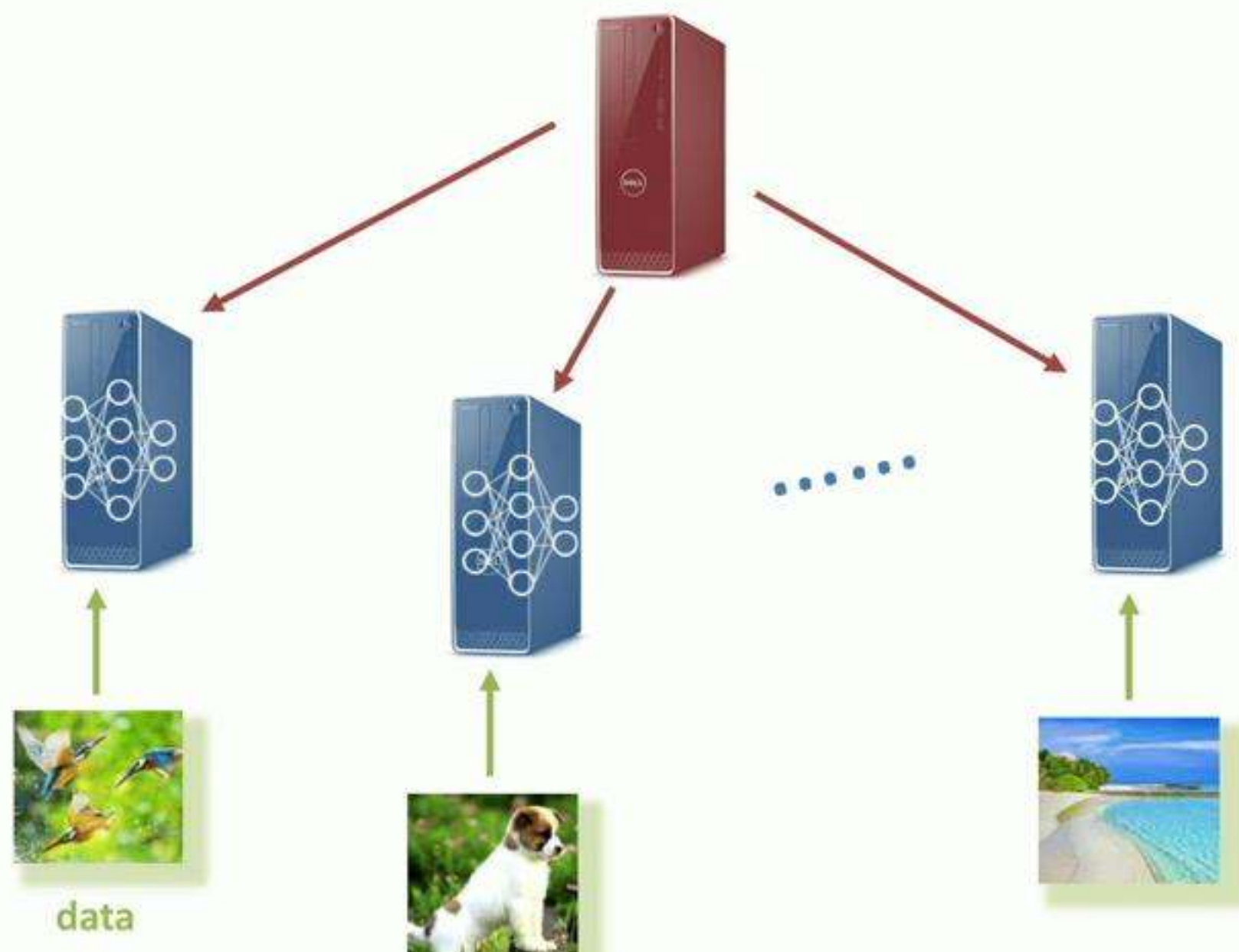
Synchronous SGD

Communication Bottleneck in Distributed Deep Learning



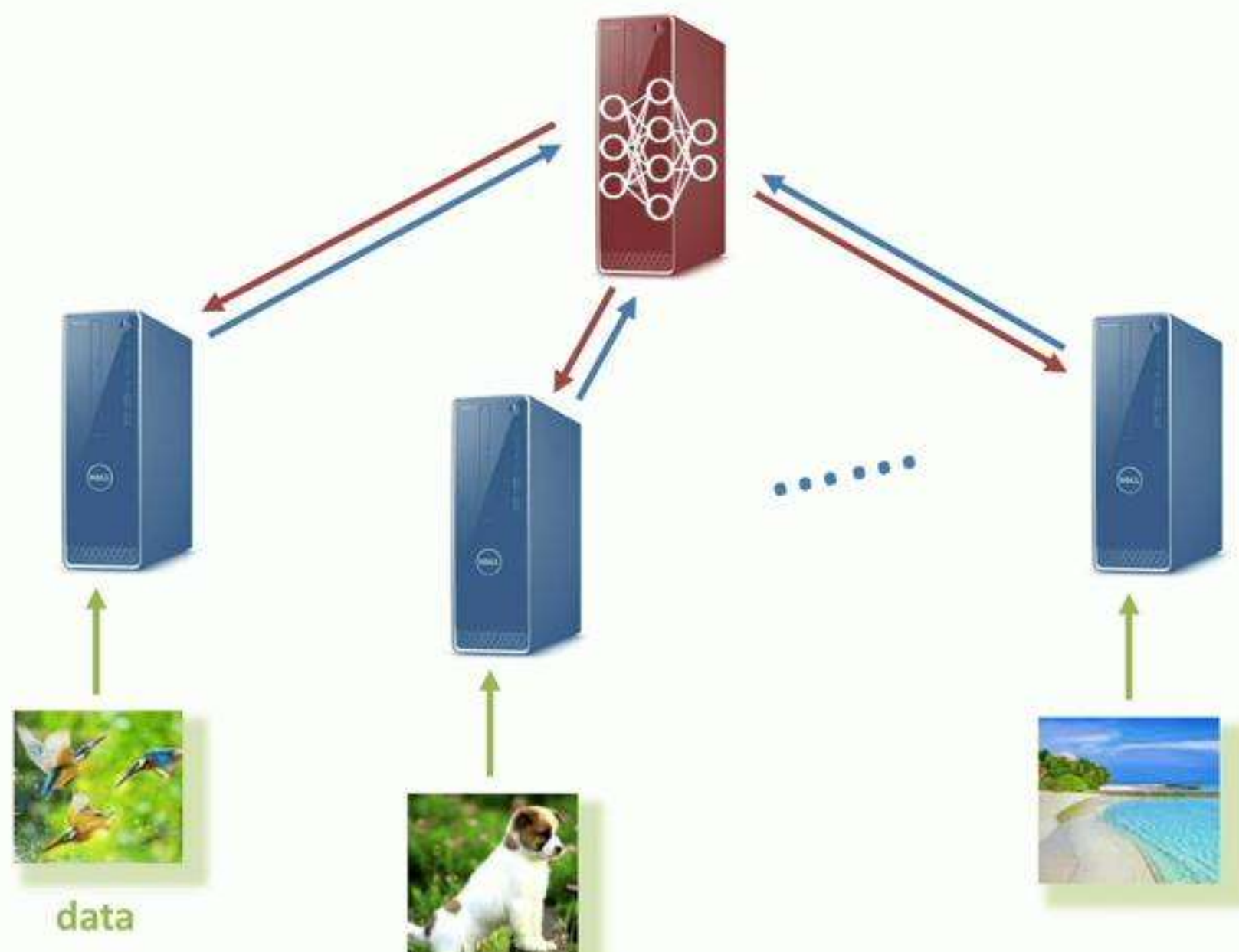
Synchronous SGD

Communication Bottleneck in Distributed Deep Learning



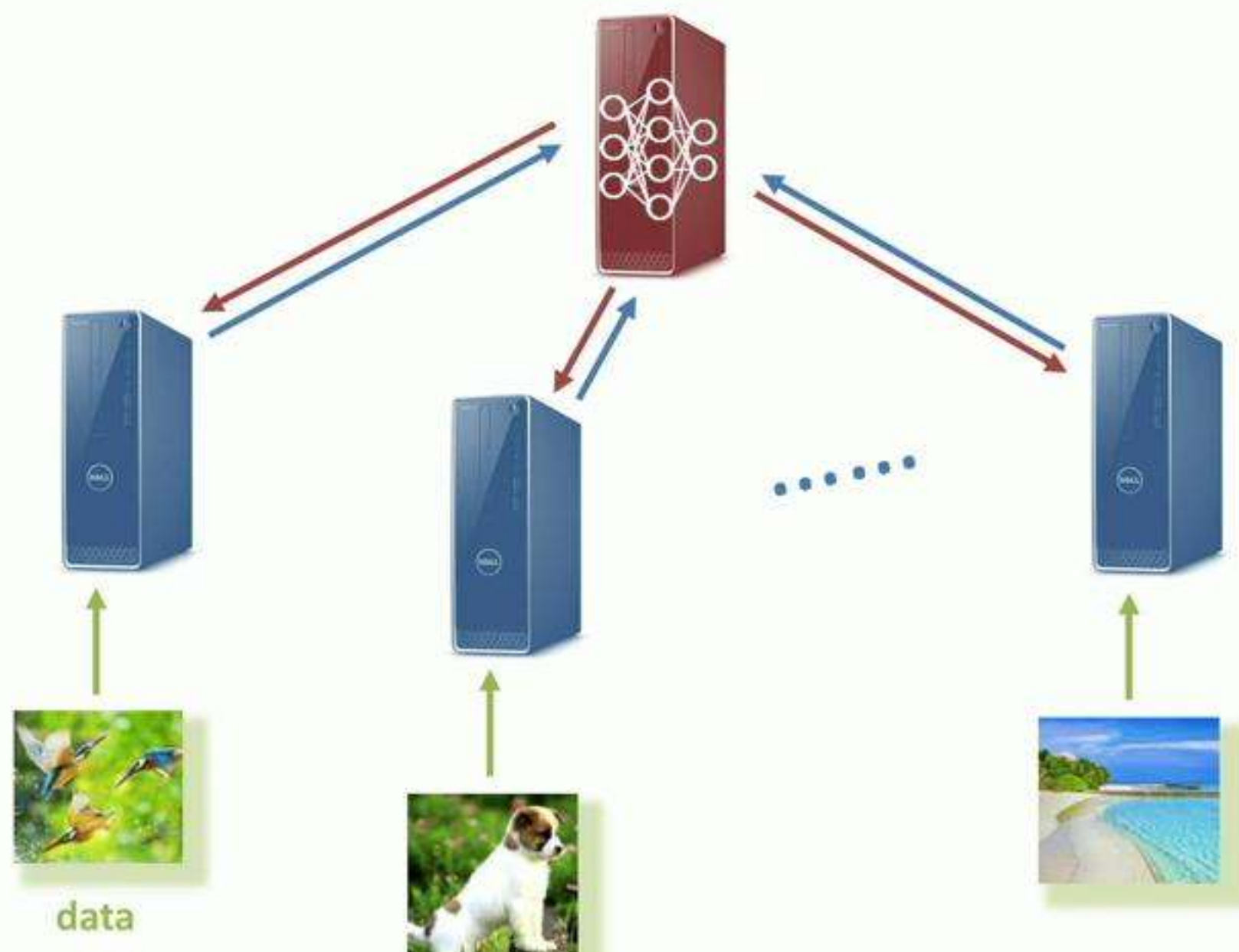
Synchronous SGD

Communication Bottleneck in Distributed Deep Learning



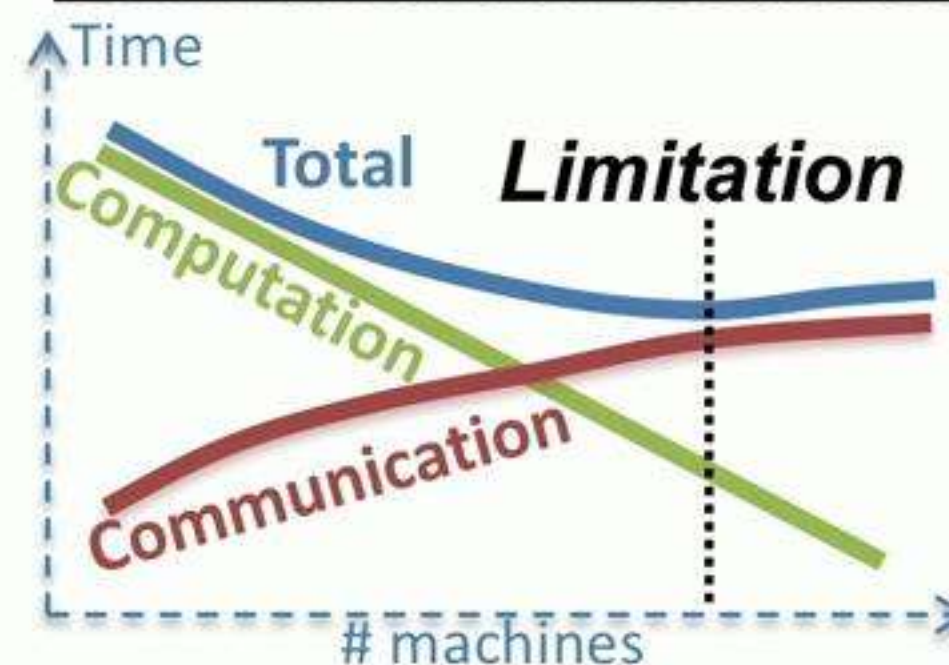
Synchronous SGD

Communication Bottleneck in Distributed Deep Learning

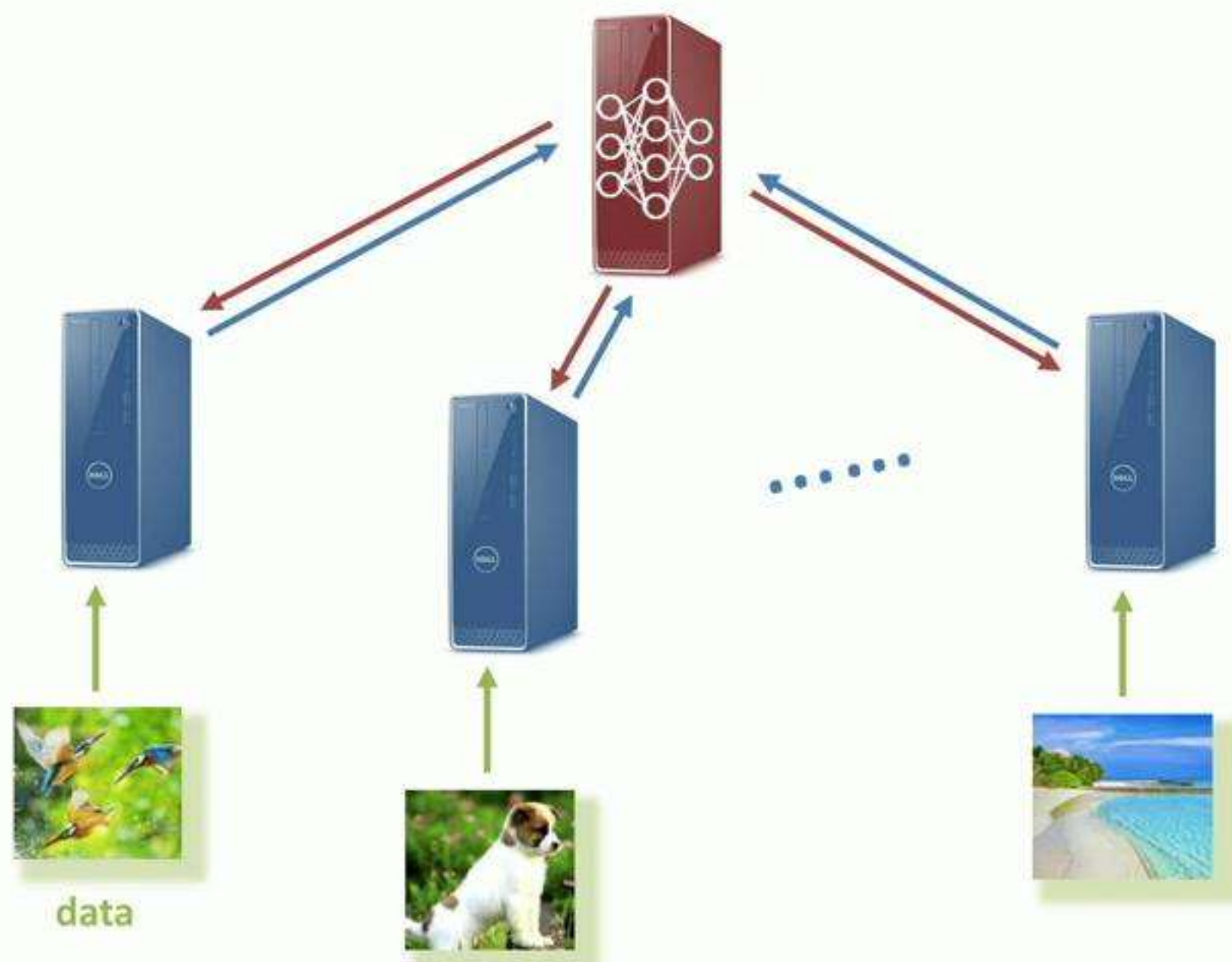


Synchronous SGD

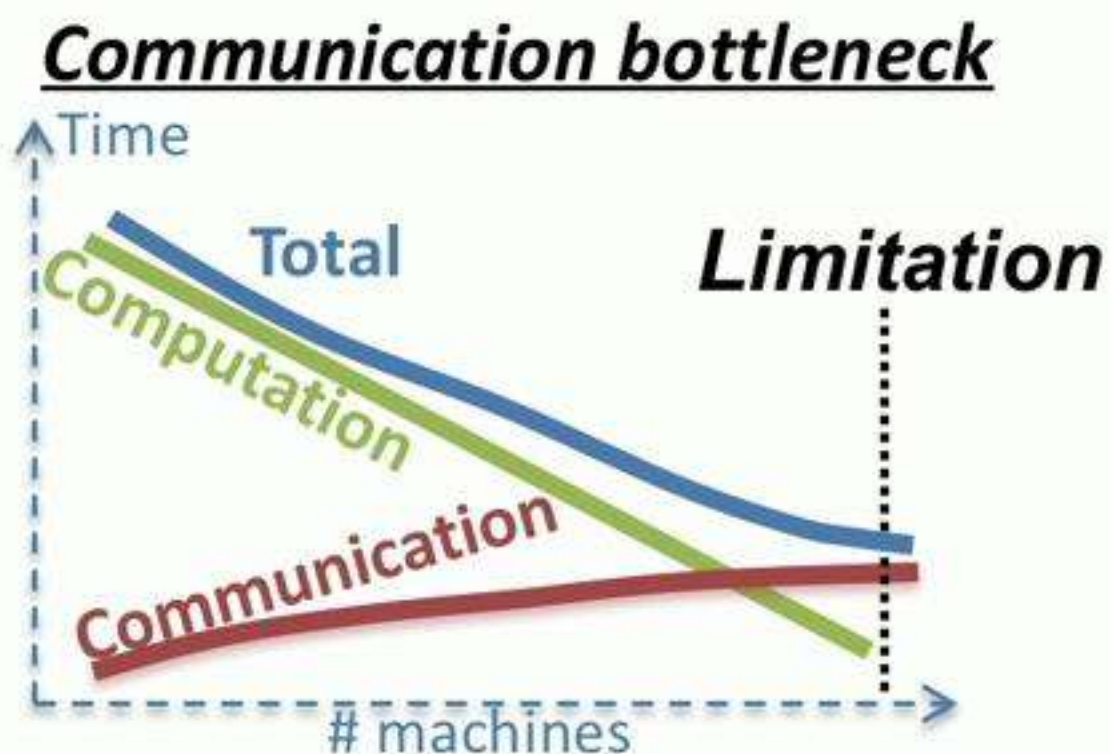
Communication bottleneck



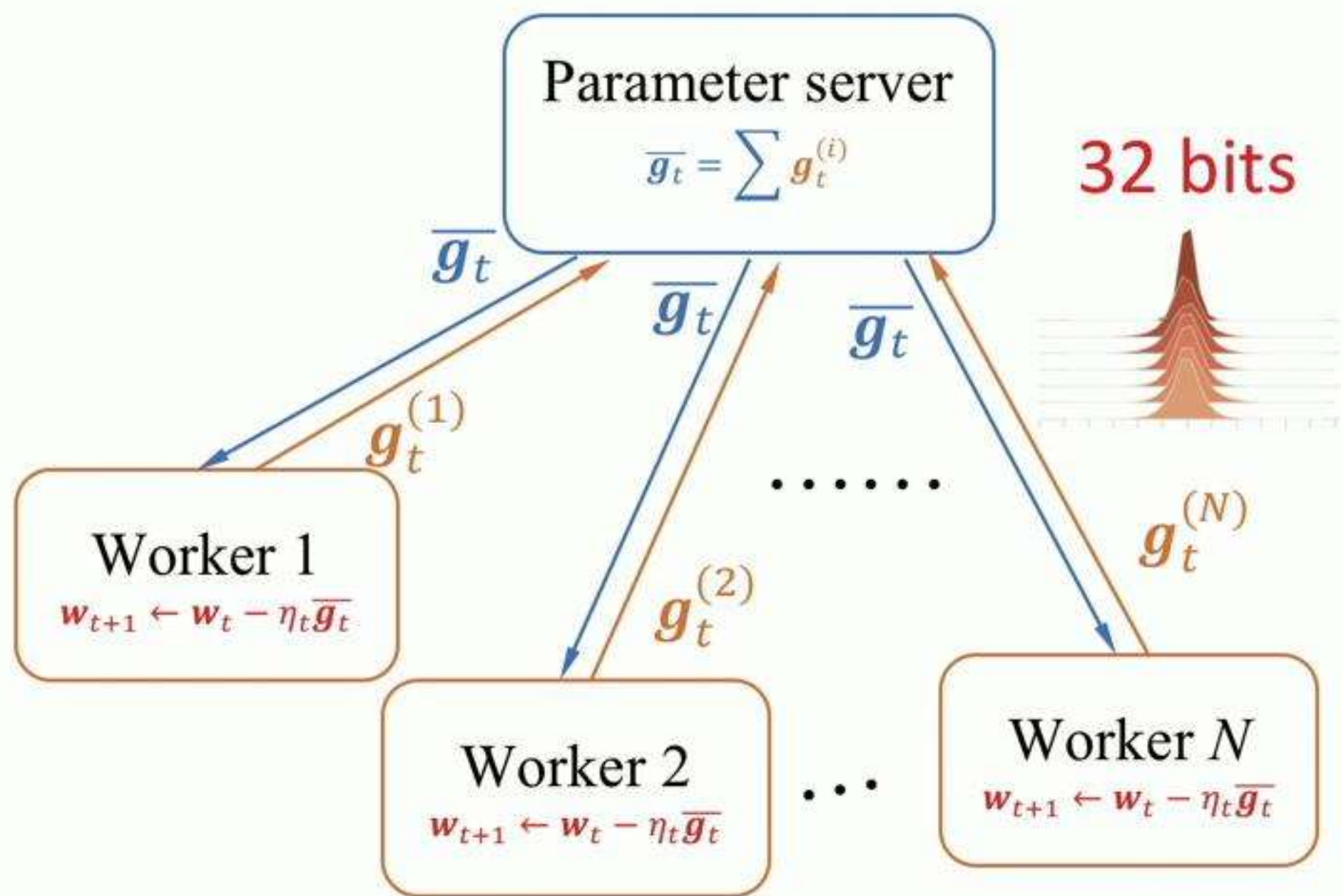
Communication Bottleneck in Distributed Deep Learning



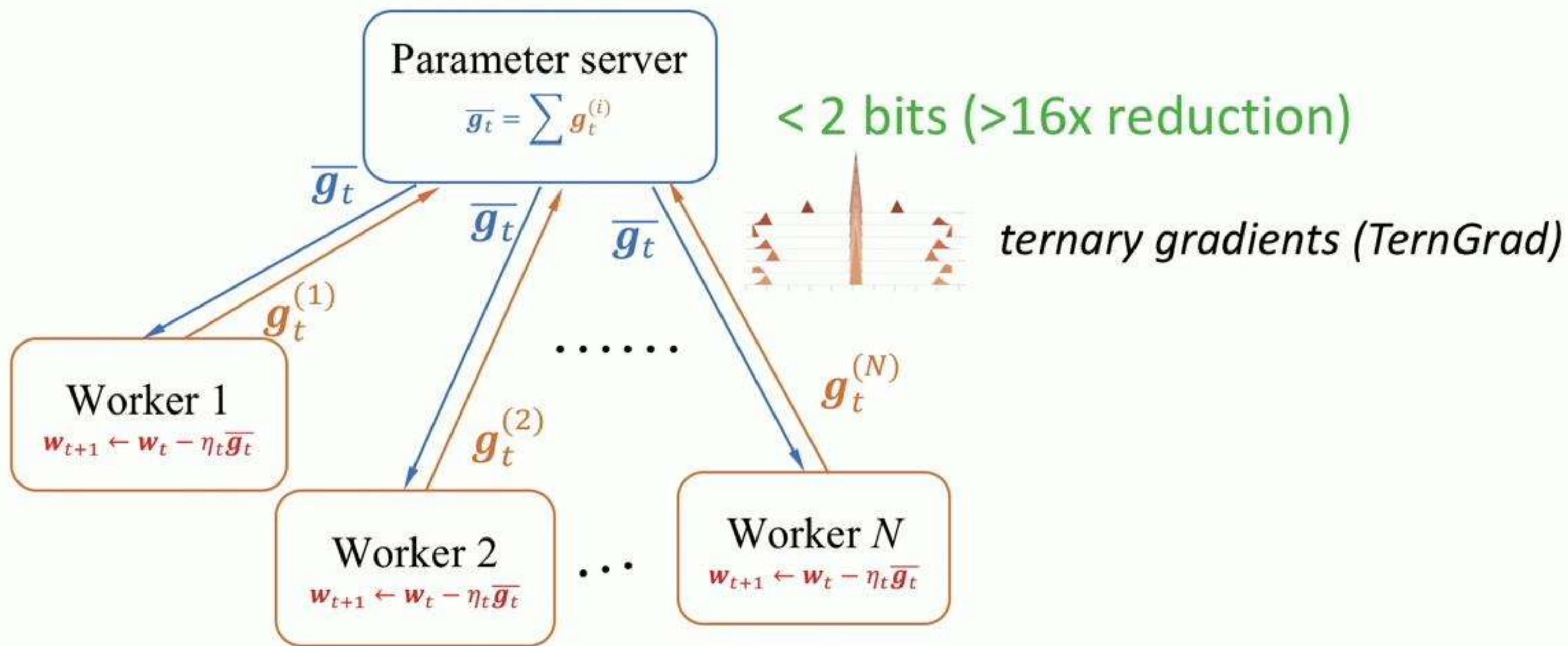
Synchronous SGD



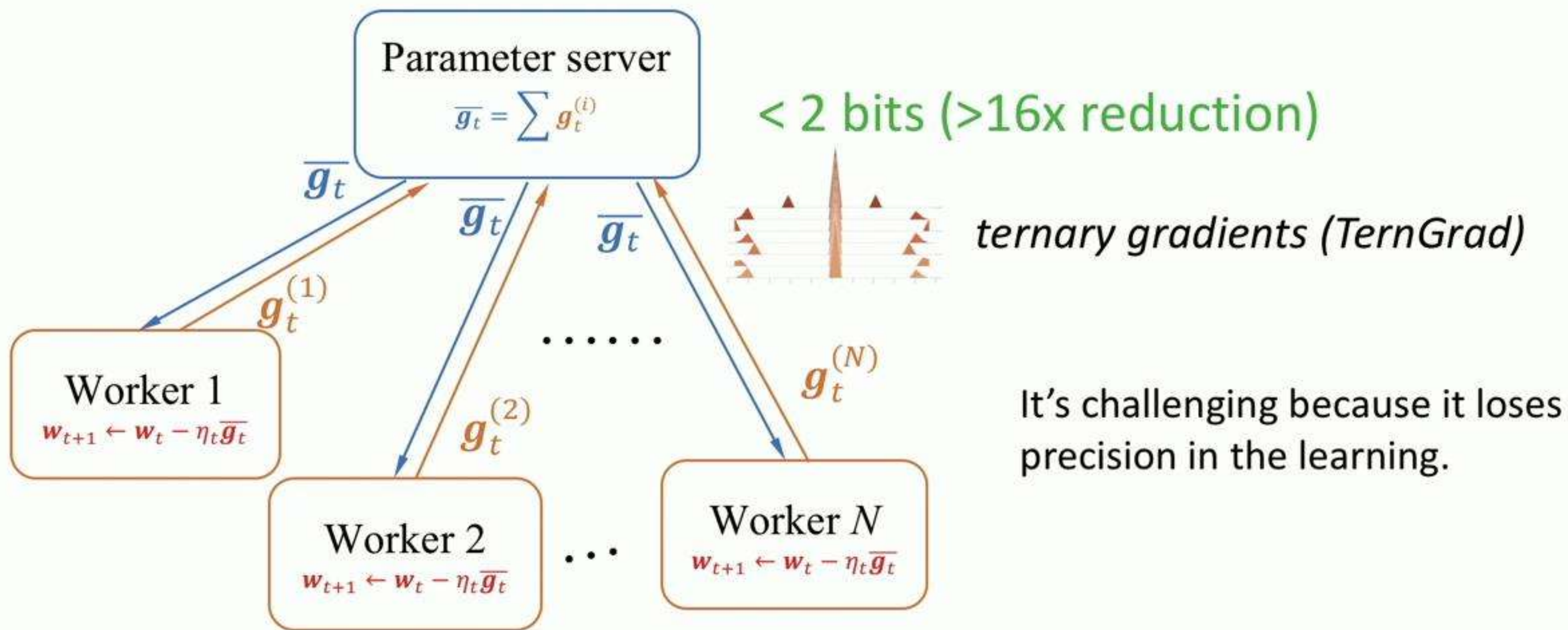
Gradient Quantization for Communication Reduction



Gradient Quantization for Communication Reduction



Gradient Quantization for Communication Reduction



Basic Idea: Stochastic Gradients without Bias

Batch Gradient Descent

$$C(\mathbf{w}) \triangleq \frac{1}{n} \sum_{i=1}^n Q(\mathbf{z}_i, \mathbf{w})$$

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \frac{\eta_t}{n} \sum_{i=1}^n g_t^{(i)}$$

SGD

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \cdot g_t^{(I)}$$

I is randomly drawn from $[1, n]$

$$\mathbb{E}\{g_t^{(I)}\} = \nabla C(\mathbf{w})$$

No bias

TernGrad

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \cdot \textit{ternarize}\left(g_t^{(I)}\right)$$
$$\mathbb{E}\{\textit{ternarize}\left(g_t^{(I)}\right)\} = \nabla C(\mathbf{w})$$

No bias

TernGrad is Simple

$$\tilde{\mathbf{g}}_t = \text{ternarize}(\mathbf{g}_t) = s_t \cdot \text{sign}(\mathbf{g}_t) \circ \mathbf{b}_t$$

$$s_t \triangleq \|\mathbf{g}_t\|_\infty \triangleq \max(\text{abs}(\mathbf{g}_t))$$

$$\begin{cases} P(b_{tk} = 1 \mid \mathbf{g}_t) = |g_{tk}|/s_t \\ P(b_{tk} = 0 \mid \mathbf{g}_t) = 1 - |g_{tk}|/s_t \end{cases}$$

Example:

$$\mathbf{g}_t^{(i)}: [0.30, -1.20, \dots, 0.9]$$

$$s_t: 1.20$$

$$\text{Signs: } [1, -1, \dots, 1]$$

$$P(b_{tk} = 1 \mid \mathbf{g}_t): [\frac{0.3}{1.2}, \frac{1.2}{1.2}, \dots, \frac{0.9}{1.2}]$$

$$\mathbf{b}_t: [0, 1, \dots, 1]$$

$$\tilde{\mathbf{g}}_t^{(i)}: [0, -1, \dots, 1] * 1.20$$

TernGrad is Simple

$$\tilde{\mathbf{g}}_t = \text{ternarize}(\mathbf{g}_t) = s_t \cdot \text{sign}(\mathbf{g}_t) \circ \mathbf{b}_t$$

$$s_t \triangleq \|\mathbf{g}_t\|_\infty \triangleq \max(\text{abs}(\mathbf{g}_t))$$

$$\begin{cases} P(b_{tk} = 1 \mid \mathbf{g}_t) = |g_{tk}|/s_t \\ P(b_{tk} = 0 \mid \mathbf{g}_t) = 1 - |g_{tk}|/s_t \end{cases}$$

$$\begin{aligned} \mathbf{E}_{\mathbf{z}, \mathbf{b}} \{\tilde{\mathbf{g}}_t\} &= \mathbf{E}_{\mathbf{z}, \mathbf{b}} \{s_t \cdot \text{sign}(\mathbf{g}_t) \circ \mathbf{b}_t\} \\ &= \mathbf{E}_{\mathbf{z}} \{s_t \cdot \text{sign}(\mathbf{g}_t) \circ \mathbf{E}_{\mathbf{b}} \{\mathbf{b}_t \mid \mathbf{z}_t\}\} = \mathbf{E}_{\mathbf{z}} \{\mathbf{g}_t\} = \nabla_{\mathbf{w}} C(\mathbf{w}_t) \end{aligned}$$

No bias

Example:

$$\mathbf{g}_t^{(i)}: [0.30, -1.20, \dots, 0.9]$$

$$s_t: 1.20$$

$$\text{Signs: } [1, -1, \dots, 1]$$

$$P(b_{tk} = 1 \mid \mathbf{g}_t): [\frac{0.3}{1.2}, \frac{1.2}{1.2}, \dots, \frac{0.9}{1.2}]$$

$$\mathbf{b}_t: [0, 1, \dots, 1]$$

$$\tilde{\mathbf{g}}_t^{(i)}: [0, -1, \dots, 1] * 1.20$$

Convergence

Standard SGD almost truly converges under assumptions (Fisk 1965, Metivier 1981&1983, Bottou 1998)

Assumption 1:

$C(\mathbf{w})$ has a single minimum $\mathbf{w}^*$ and $\forall \epsilon > 0, \inf_{\|\mathbf{w} - \mathbf{w}^*\|^2 > \epsilon} (\mathbf{w} - \mathbf{w}^*)^T \nabla_{\mathbf{w}} C(\mathbf{w}) > 0$

Assumption 2:

Learning rate γ_t decreases neither very fast nor very slow $\begin{cases} \sum_{t=0}^{+\infty} \gamma_t^2 < +\infty \\ \sum_{t=0}^{+\infty} \gamma_t = +\infty \end{cases}$

Assumption 3 (gradient bound):

$$\mathbf{E} \{ \|\mathbf{g}\|^2 \} \leq A + B \|\mathbf{w} - \mathbf{w}^*\|^2$$

Standard SGD *almost-truly* converges

Convergence

Standard SGD almost truly converges under assumptions (Fisk 1965, Metivier 1981&1983, Bottou 1998)

Assumption 1:

$C(\mathbf{w})$ has a single minimum $\mathbf{w}^*$ and $\forall \epsilon > 0, \inf_{\|\mathbf{w} - \mathbf{w}^*\|^2 > \epsilon} (\mathbf{w} - \mathbf{w}^*)^T \nabla_{\mathbf{w}} C(\mathbf{w}) > 0$

Assumption 2:

Learning rate γ_t decreases neither very fast nor very slow $\begin{cases} \sum_{t=0}^{+\infty} \gamma_t^2 < +\infty \\ \sum_{t=0}^{+\infty} \gamma_t = +\infty \end{cases}$

Assumption 3 (gradient bound):

$$\mathbf{E} \{ \|\mathbf{g}\|^2 \} \leq A + B \|\mathbf{w} - \mathbf{w}^*\|^2$$

Standard SGD *almost-truly* converges

Assumption 3 (gradient bound):

$$\mathbf{E} \{ \|\mathbf{g}\|_{\infty} \cdot \|\mathbf{g}\|_1 \} \leq A + B \|\mathbf{w} - \mathbf{w}^*\|^2$$

TernGrad almost-truly converges

Convergence

Standard SGD almost truly converges under assumptions (Fisk 1965, Metivier 1981&1983, Bottou 1998)

Assumption 1:

$C(\mathbf{w})$ has a single minimum $\mathbf{w}^*$ and $\forall \epsilon > 0, \inf_{\|\mathbf{w} - \mathbf{w}^*\|^2 > \epsilon} (\mathbf{w} - \mathbf{w}^*)^T \nabla_{\mathbf{w}} C(\mathbf{w}) > 0$

Assumption 2:

Learning rate γ_t decreases neither very fast nor very slow $\begin{cases} \sum_{t=0}^{+\infty} \gamma_t^2 < +\infty \\ \sum_{t=0}^{+\infty} \gamma_t = +\infty \end{cases}$

Assumption 3 (gradient bound):

$$\mathbf{E} \{ \|\mathbf{g}\|^2 \} \leq A + B \|\mathbf{w} - \mathbf{w}^*\|^2$$

Standard SGD *almost-truly* converges

Assumption 3 (gradient bound):

$$\mathbf{E} \{ \|\mathbf{g}\|_{\infty} \cdot \|\mathbf{g}\|_1 \} \leq A + B \|\mathbf{w} - \mathbf{w}^*\|^2$$

TernGrad almost-truly converges

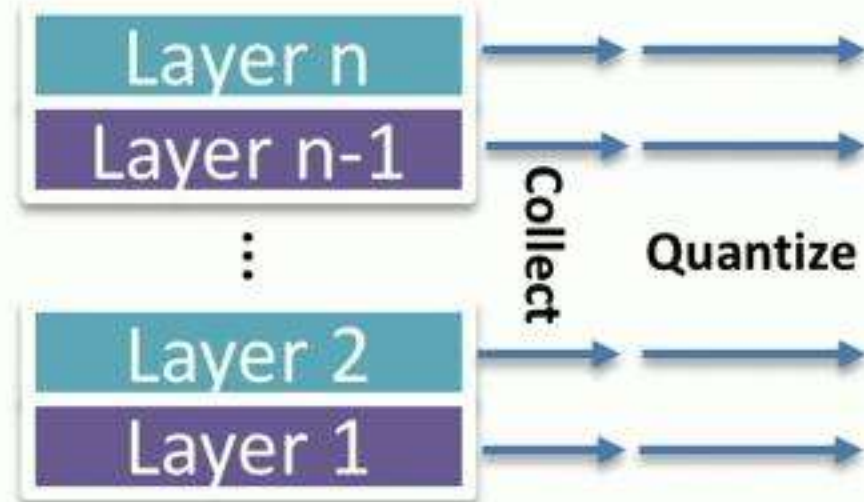
$$\mathbf{E} \{ \|\mathbf{g}\|^2 \} \leq \mathbf{E} \{ \|\mathbf{g}\|_{\infty} \cdot \|\mathbf{g}\|_1 \} \leq A + B \|\mathbf{w} - \mathbf{w}^*\|^2$$

Stronger gradient bound in TernGrad

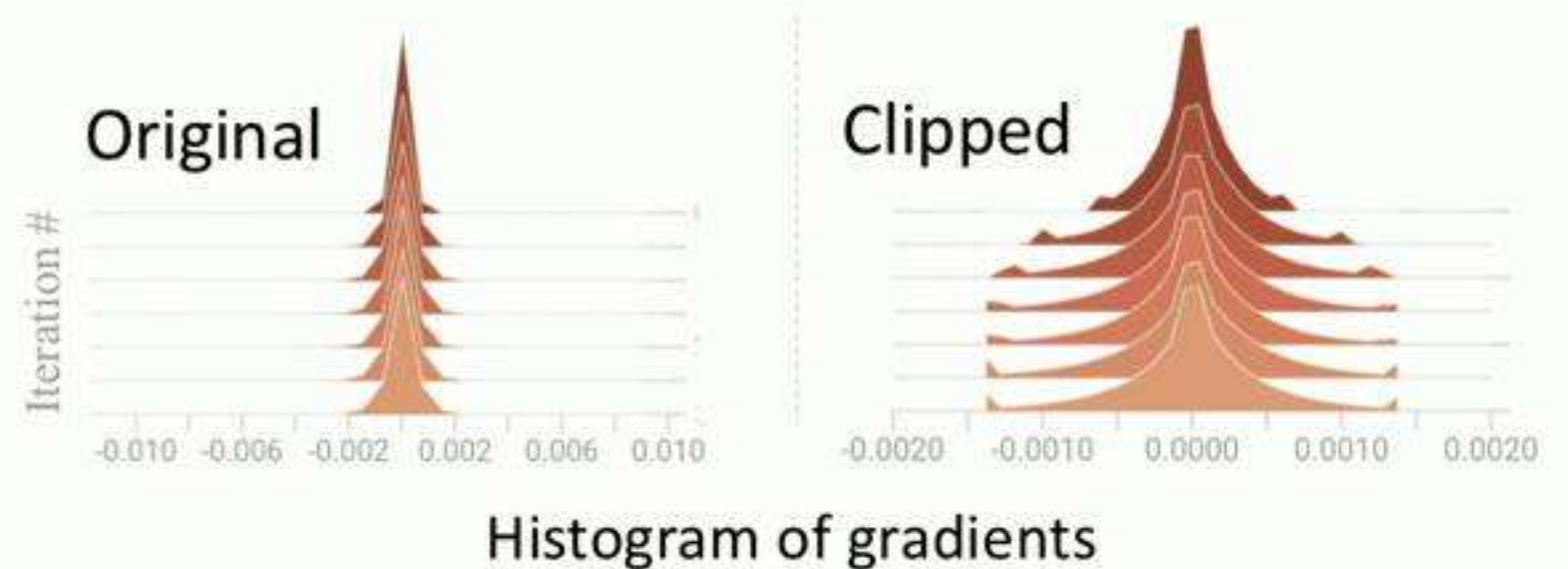
Closing Bound Gap

Two methods to push the gradient bound of *TernGrad* closer to the bound of standard SGD

Layer-wise ternarizing



Gradient clipping

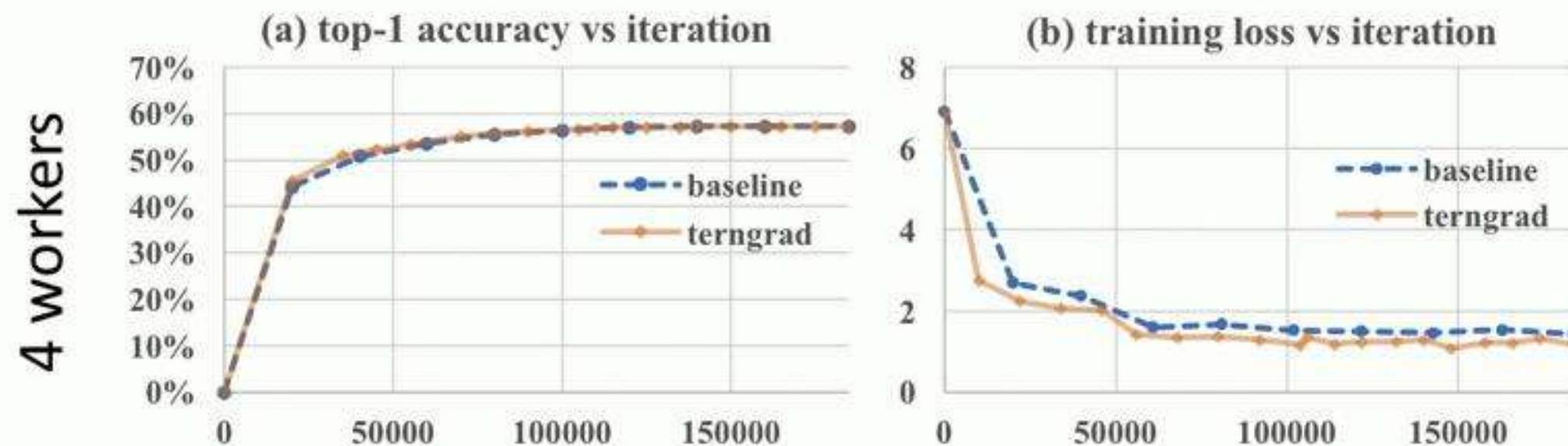


Training ImageNet with Ternary Gradients

AlexNet

base LR	mini-batch size	workers	gradients	top-1
0.01	256	2	floating	57.33%
			<i>TernGrad</i>	57.61%
			<i>TernGrad</i> -noclip	54.63%
0.02	512	4	floating	57.32%
			<i>TernGrad</i>	57.28%
0.04	1024	8	floating	56.62%
			<i>TernGrad</i>	57.54%

Keskar, et al., 2017



Training ImageNet with Ternary Gradients

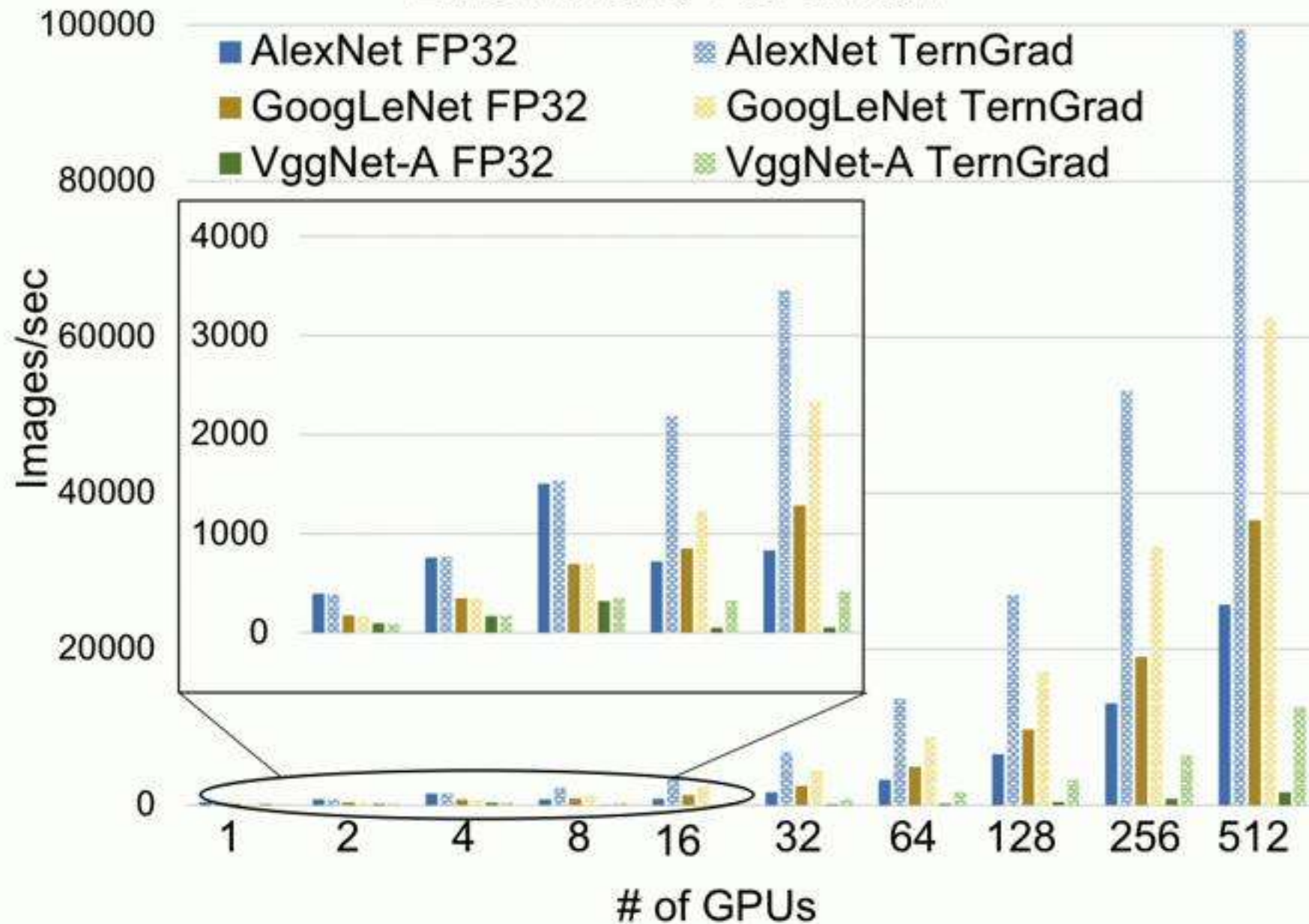
GoogLeNet accuracy loss is <2% on average.

base LR	mini-batch size	workers	gradients	top-5
0.04	128	2	floating	88.30%
			<i>TernGrad</i>	86.77%
0.08	256	4	floating	87.82%
			<i>TernGrad</i>	85.96%
0.10	512	8	floating	89.00%
			<i>TernGrad</i>	86.47%

All hyper-parameters in standard SGD are directed used in *TernGrad* (such as learning rate, batch size, total epochs, etc.)

Speedup

Training throughput on GPU cluster with Ethernet and PCI switch



TernGrad gives higher speedup when communication time occupies a higher ratio, such as:

1. using more machines;
2. having smaller communication bandwidth, e.g., Ethernet vs. InfiniBand;
3. training a DNN with a higher “parameter #: computation #”
4. using GPU-based distributed systems vs. CPU-based

TernGrad in Production

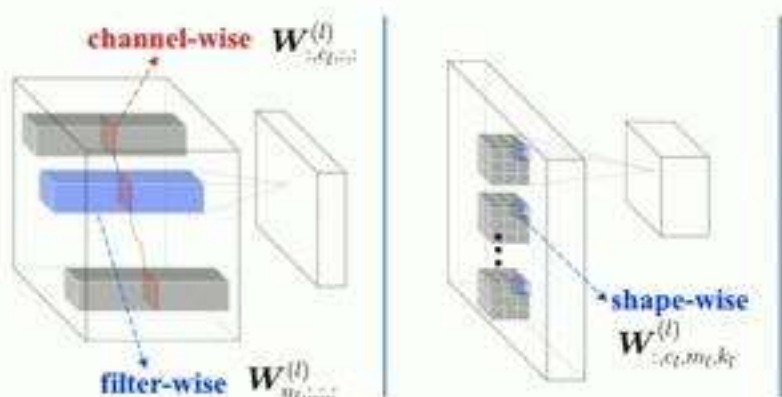
- Facebook is using *TernGrad* to overcome communication bottleneck in AI Infra
- Effectiveness is evaluated by ad ranking models, which have zero toleration on prediction loss
- Available in PyTorch/Caffe2



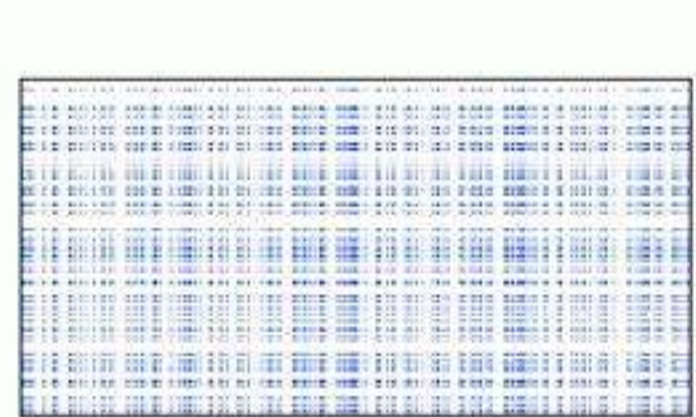
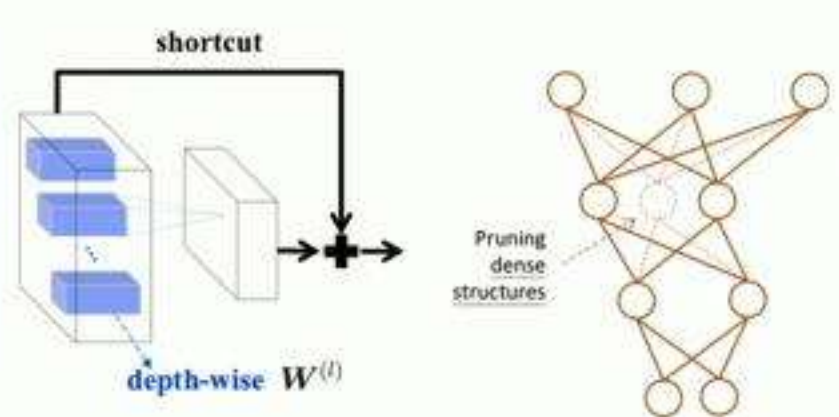
PyTorch



Inference Acceleration



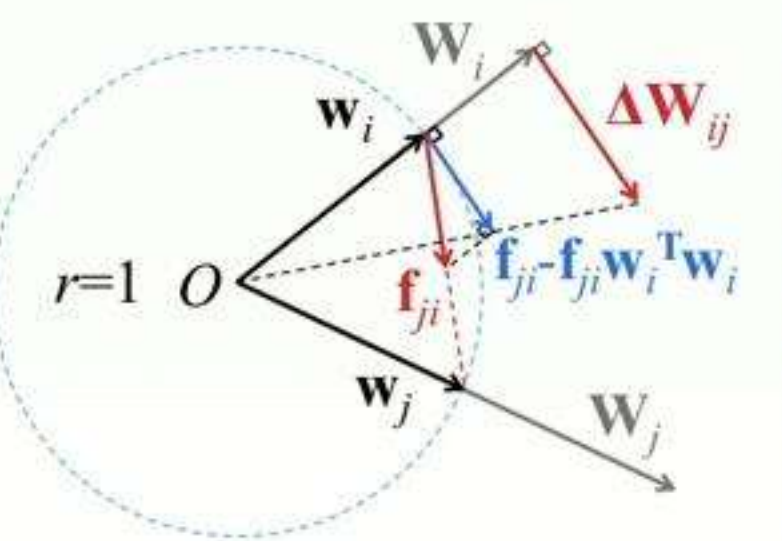
Wen et al. NeurIPS 2016



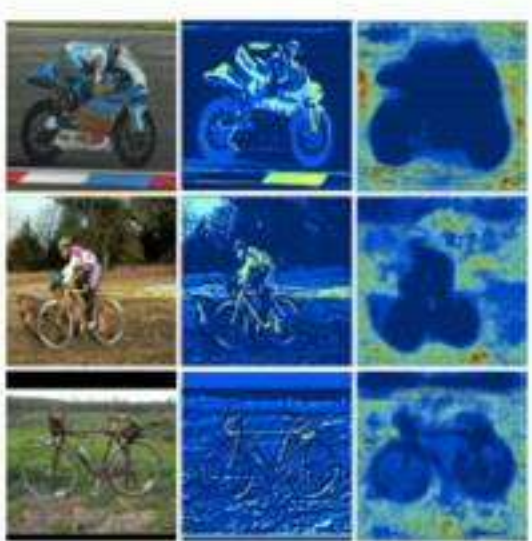
Wen et al. ICLR 2018



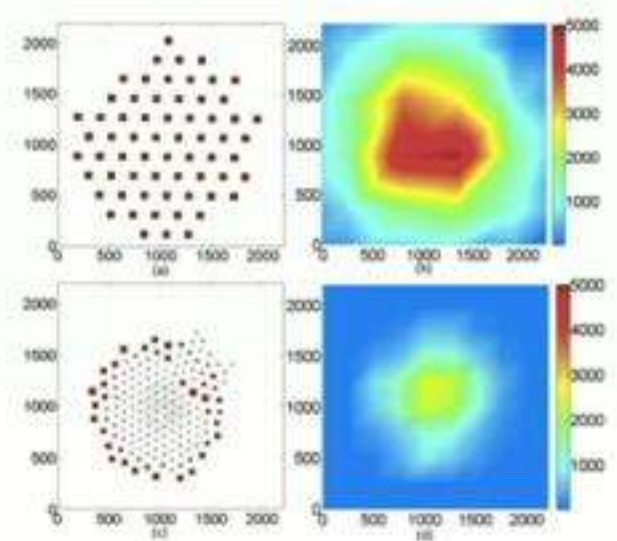
Wang et al. ASP-DAC 2017
(Best Paper Award)



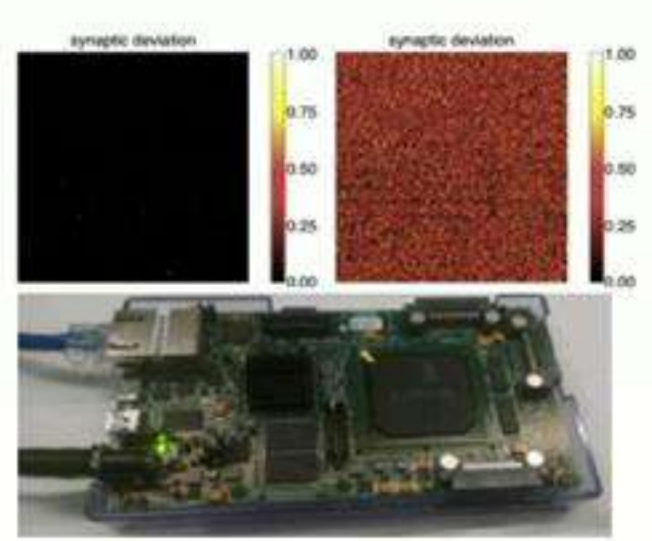
Wen et al. ICCV 2017



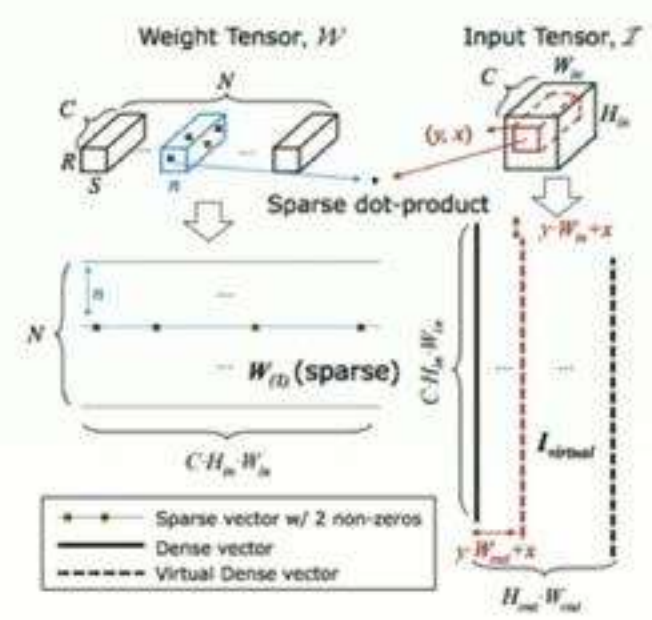
Wu et al. CVPR 2017



Wen et al. DAC 2015
(Best Paper Nomination)



Wen et al. DAC 2016
(Best Paper Nomination)



Park et al. ICLR 2017

TernGrad in Production

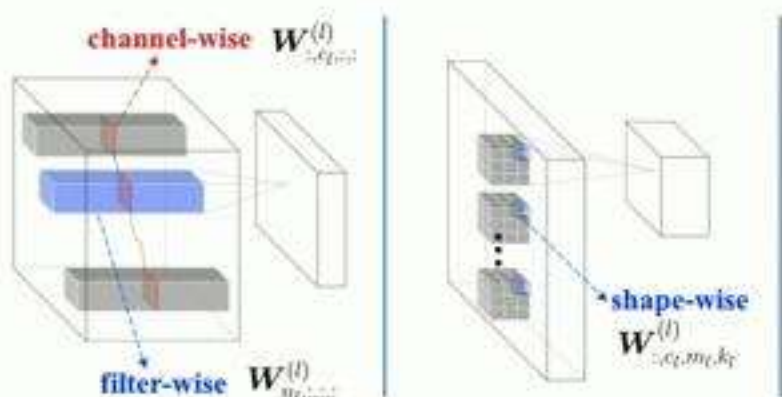
- Facebook is using *TernGrad* to overcome communication bottleneck in AI Infra
- Effectiveness is evaluated by ad ranking models, which have zero toleration on prediction loss
- Available in PyTorch/Caffe2



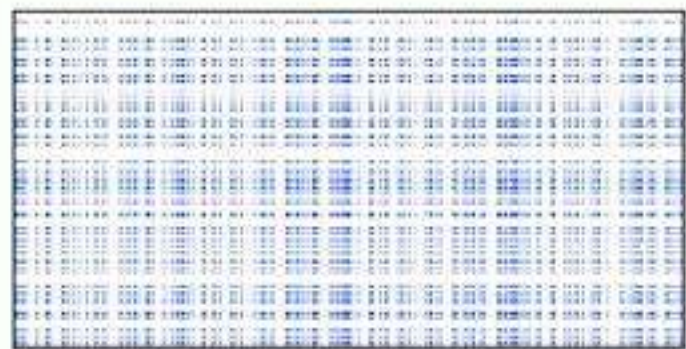
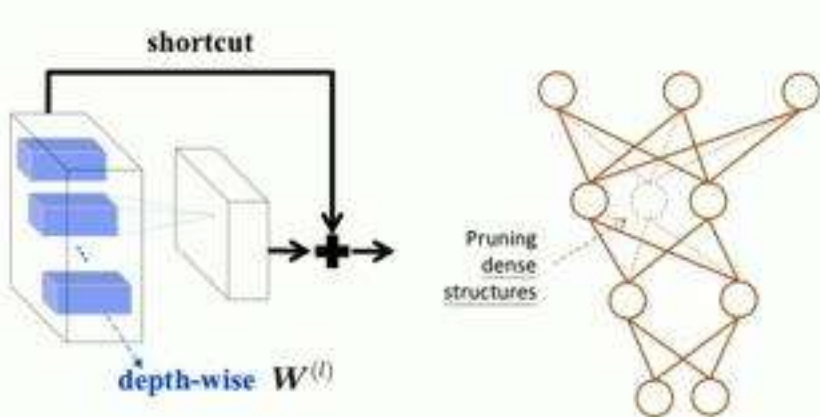
PyTorch



Inference Acceleration



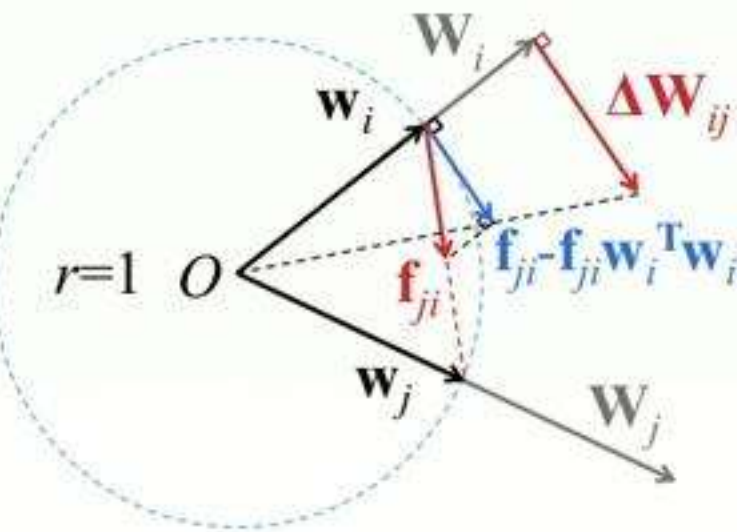
Wen et al. NeurIPS 2016



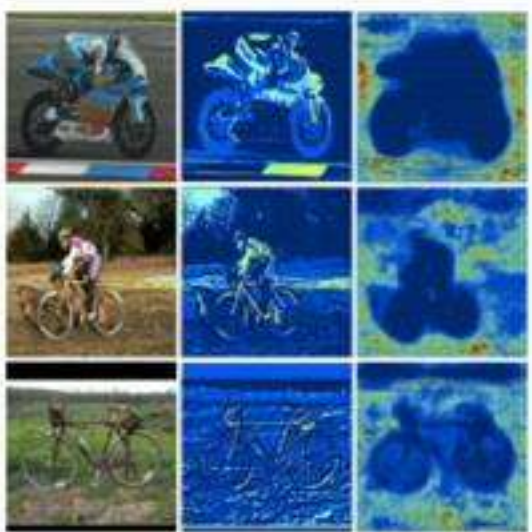
Wen et al. ICLR 2018



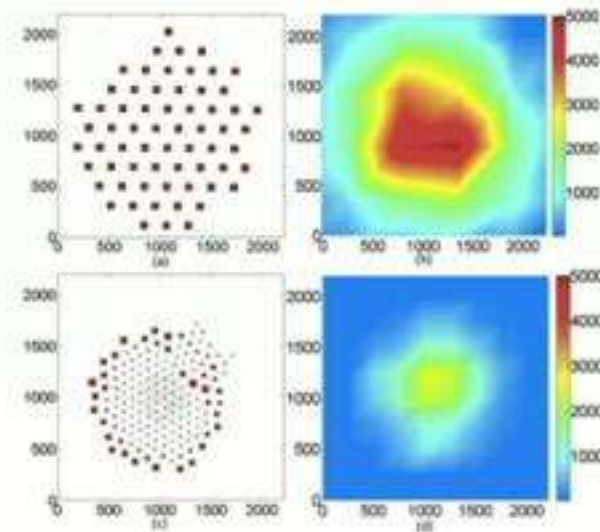
Wang et al. ASP-DAC 2017
(Best Paper Award)



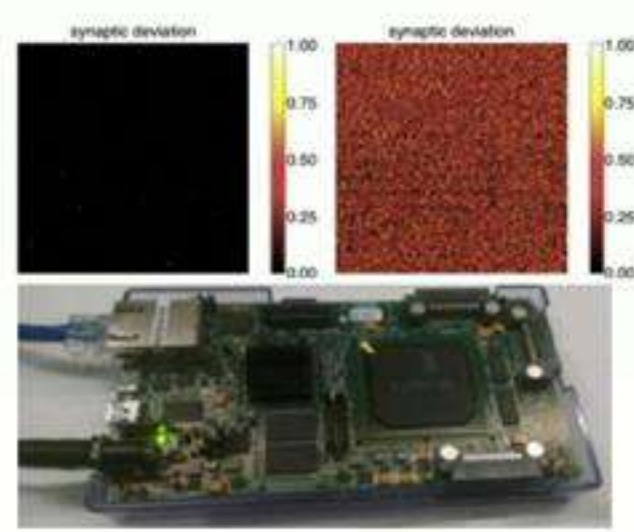
Wen et al. ICCV 2017



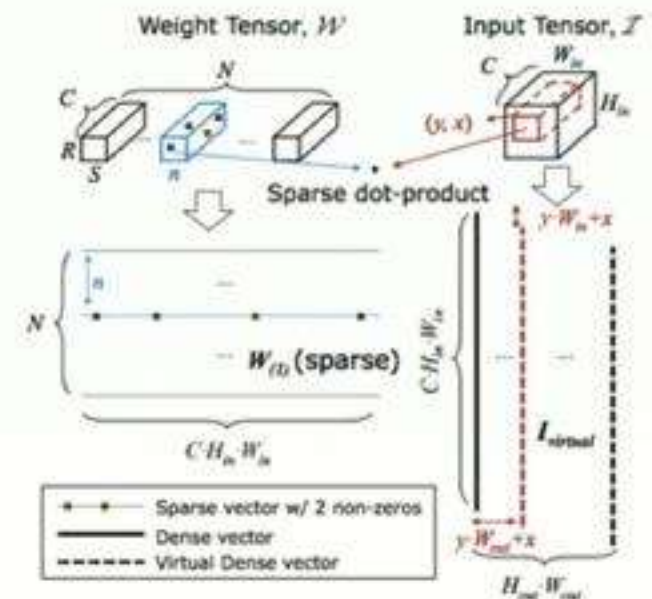
Wu et al. CVPR 2017



Wen et al. DAC 2015
(Best Paper Nomination)

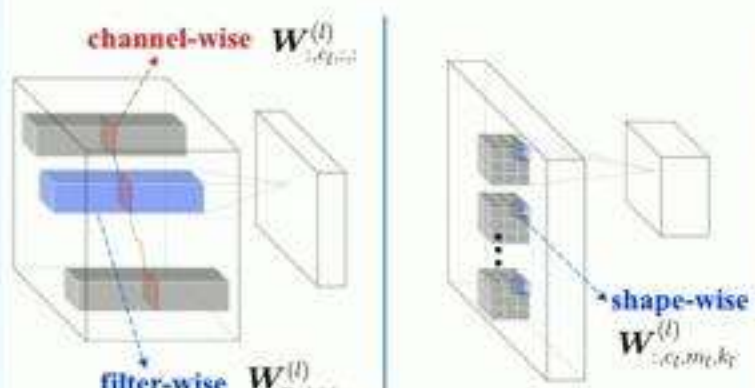


Wen et al. DAC 2016
(Best Paper Nomination)



Park et al. ICLR 2017

Inference Acceleration



channel-wise $W_{i,l,l+1}^{(l)}$

filter-wise $W_{n_l, \dots}^{(l)}$

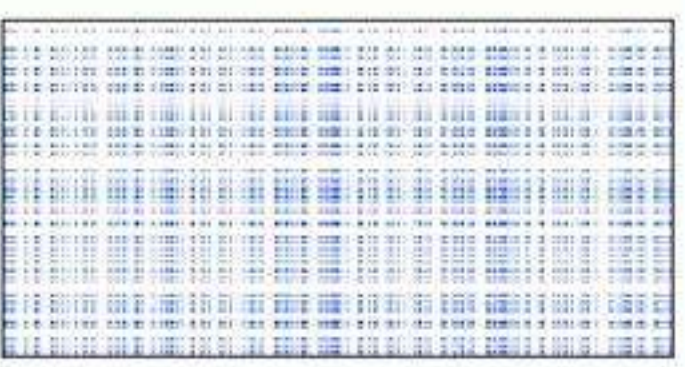
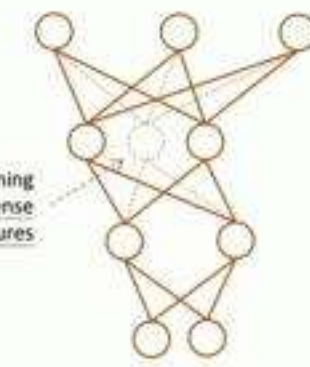
shape-wise $W_{:,c_l,m_l,k_l}^{(l)}$

depth-wise $W^{(l)}$

shortcut

Pruning dense structures

This presentation

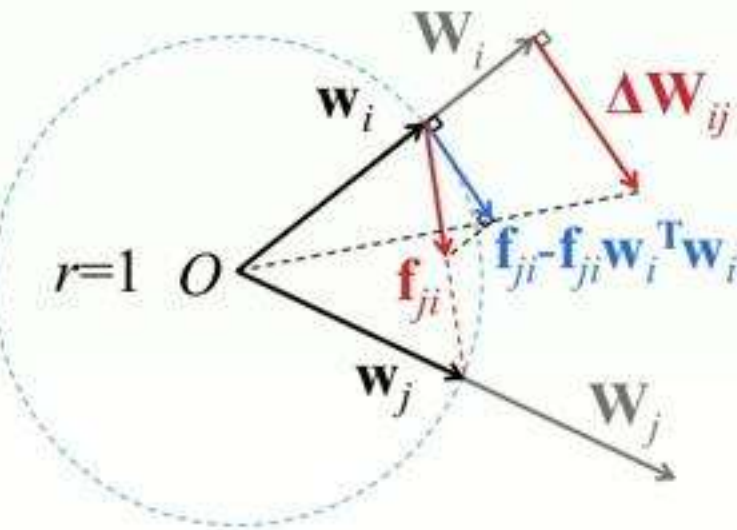


Wen et al. NeurIPS 2016

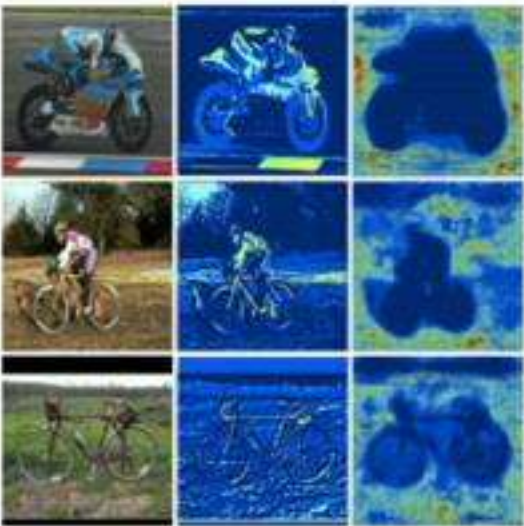
Wen et al. ICLR 2018



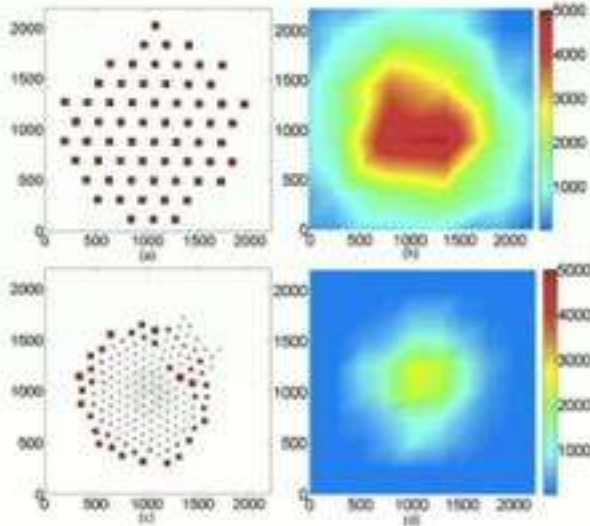
Wang et al. ASP-DAC 2017
(Best Paper Award)



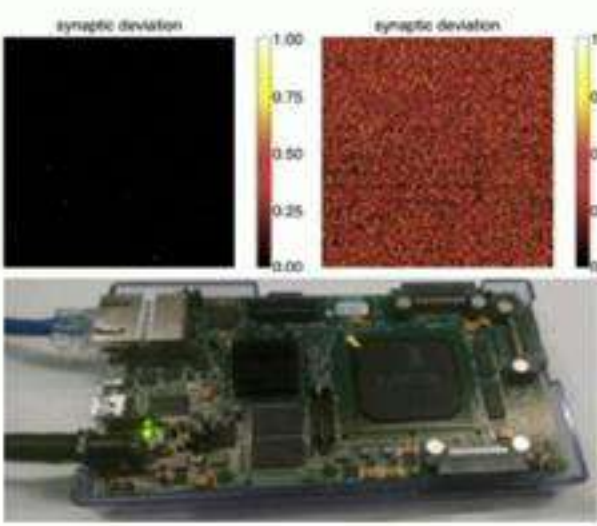
Wen et al. ICCV 2017



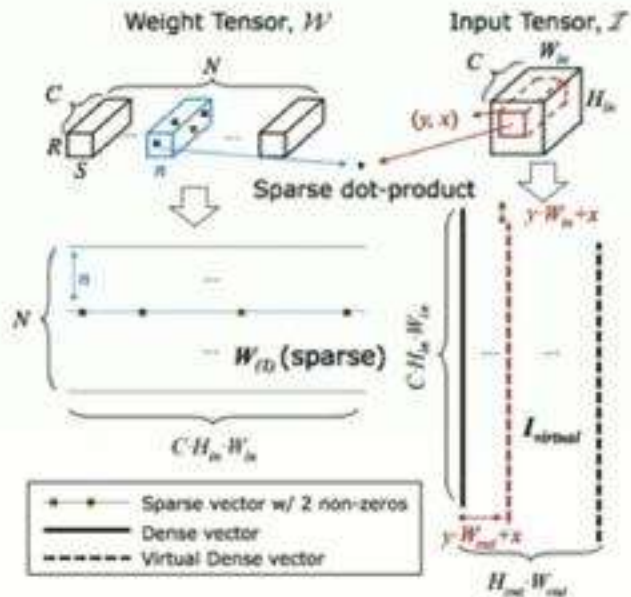
Wu et al. CVPR 2017



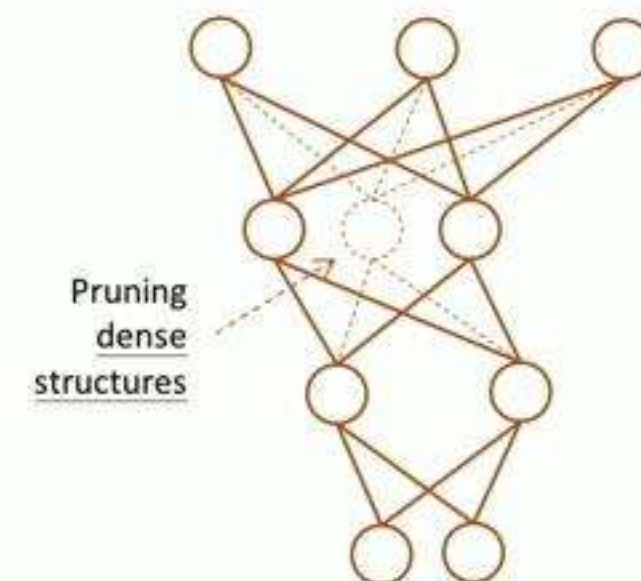
Wen et al. DAC 2015
(Best Paper Nomination)



Wen et al. DAC 2016
(Best Paper Nomination)



Park et al. ICLR 2017

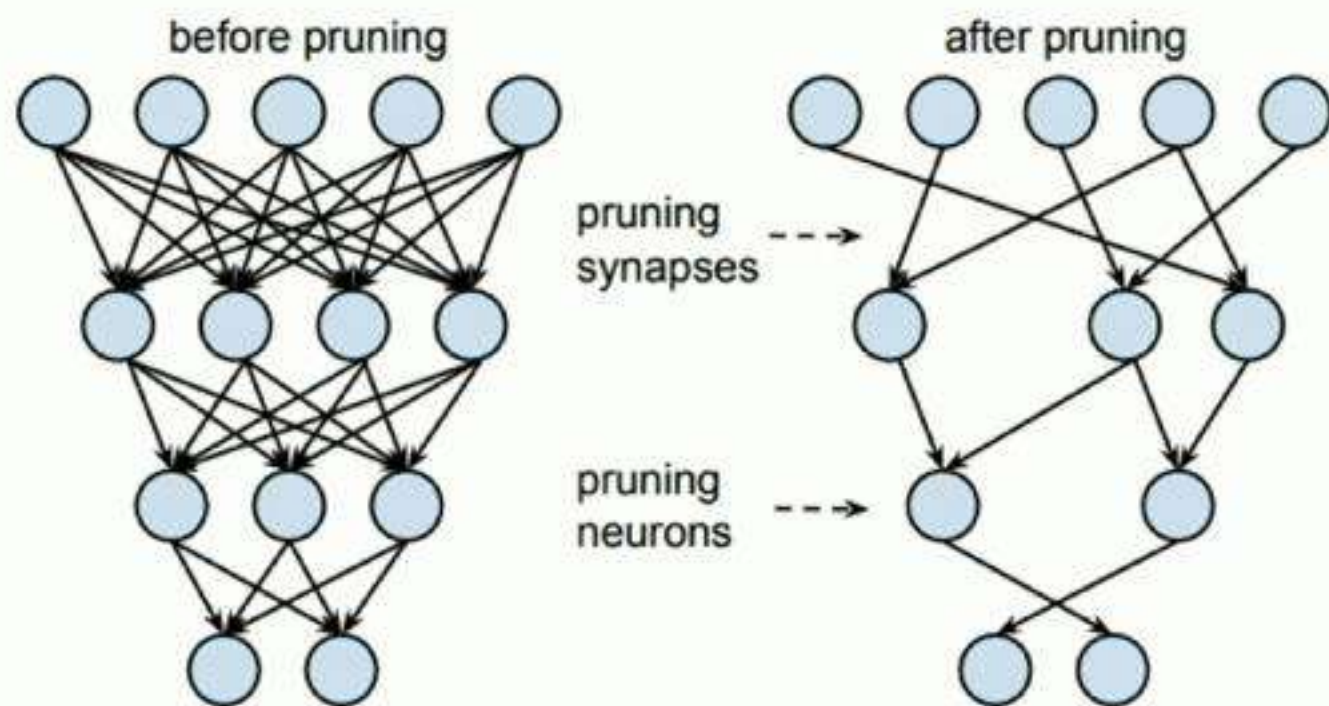


Sparse Deep Neural Networks

Wen, Wei, Chunpeng Wu, Yandan Wang, Yiran Chen, and Hai Li. "Learning structured sparsity in deep neural networks." In *Advances in neural information processing systems*, pp. 2074-2082. 2016.

Wen, Wei, Yuxiong He, Samyam Rajbhandari, Minjia Zhang, Wenhan Wang, Fang Liu, Bin Hu, Yiran Chen, and Hai Li. "Learning intrinsic sparse structures within long short-term memory." *ICLR* 2018.

Inefficiency of Sparse Convolutional Neural Networks



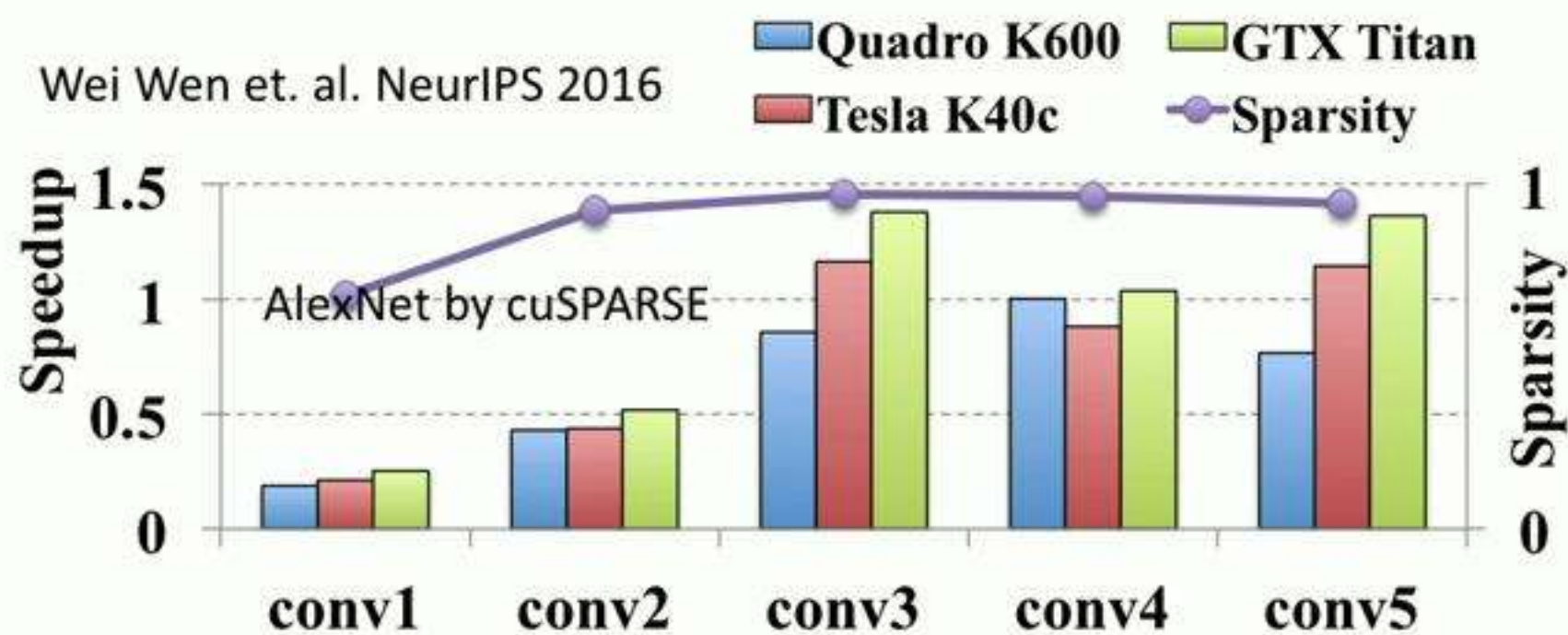
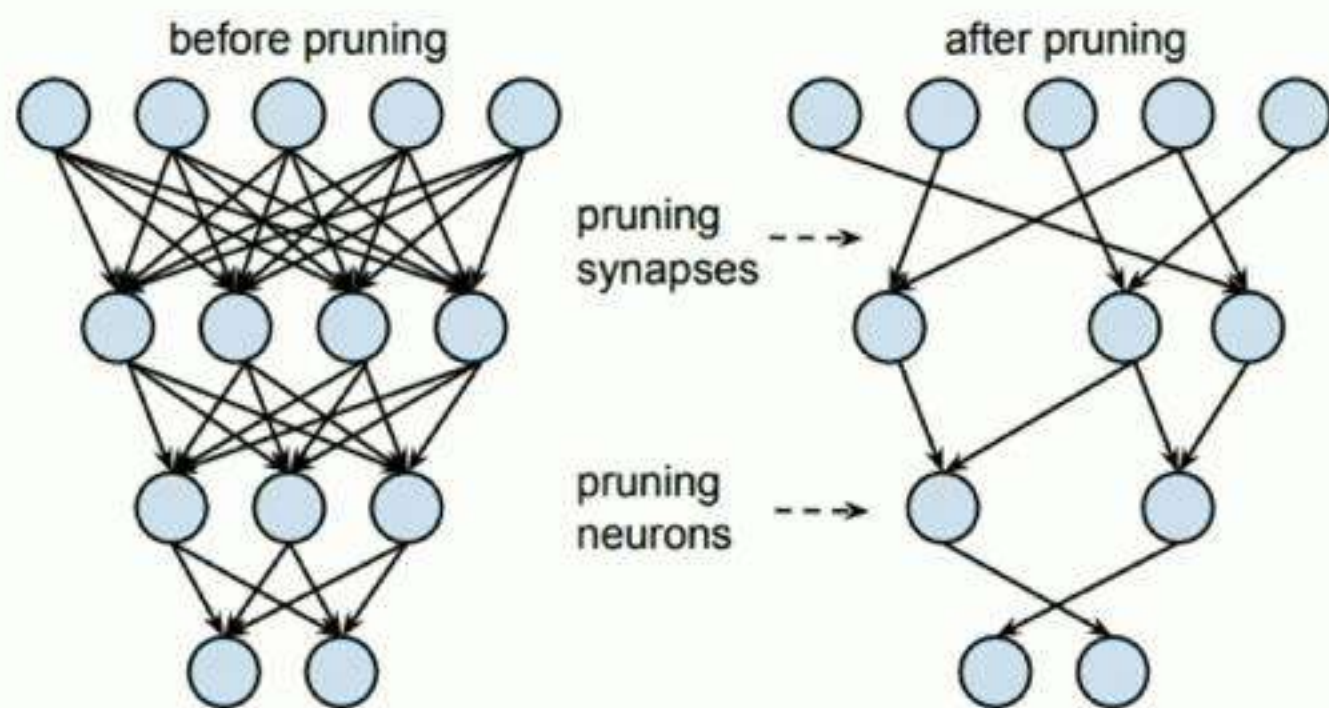
Significantly reduces the storage size of DNNs [1]
Good speedup with customized hardware [2]

[1] S. Han et. al. NeurIPS 2015

[2] S. Han et. al. ISCA 2016.



Inefficiency of Sparse Convolutional Neural Networks

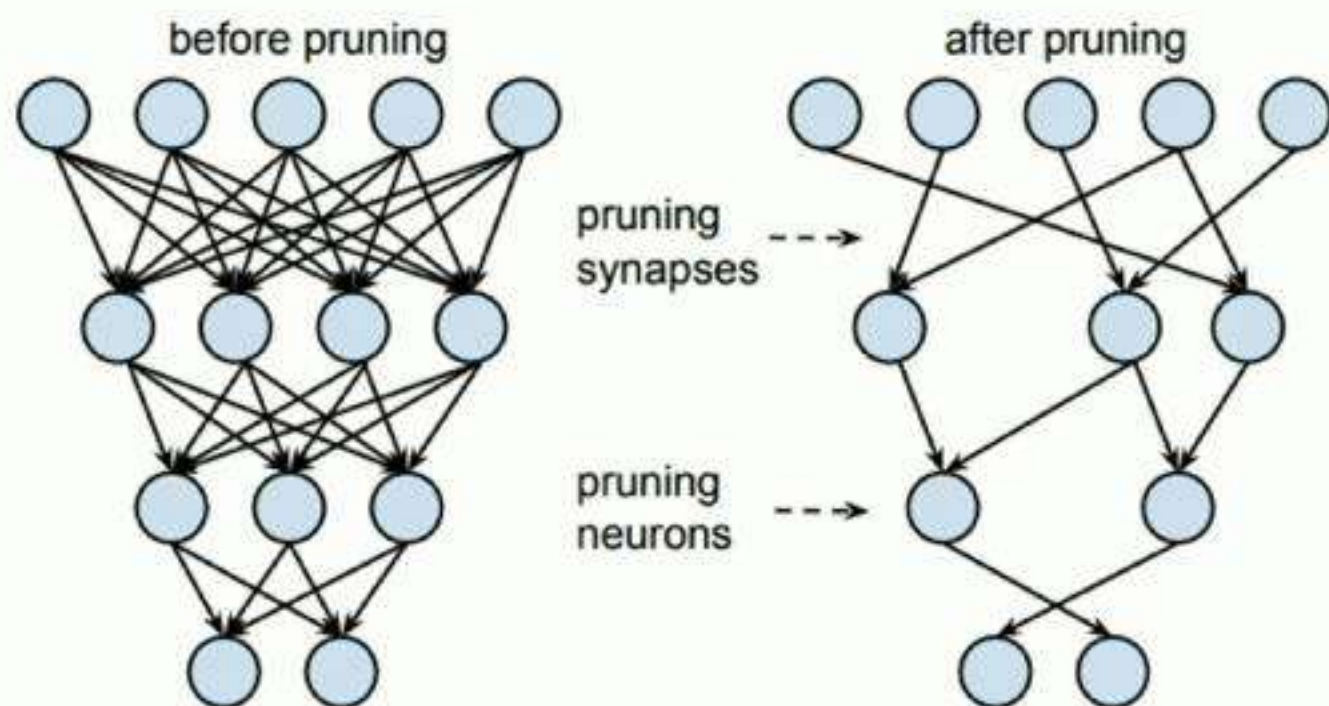


Significantly reduces the storage size of DNNs [1]
Good speedup with customized hardware [2]

[1] S. Han et. al. NeurIPS 2015

[2] S. Han et. al. ISCA 2016.

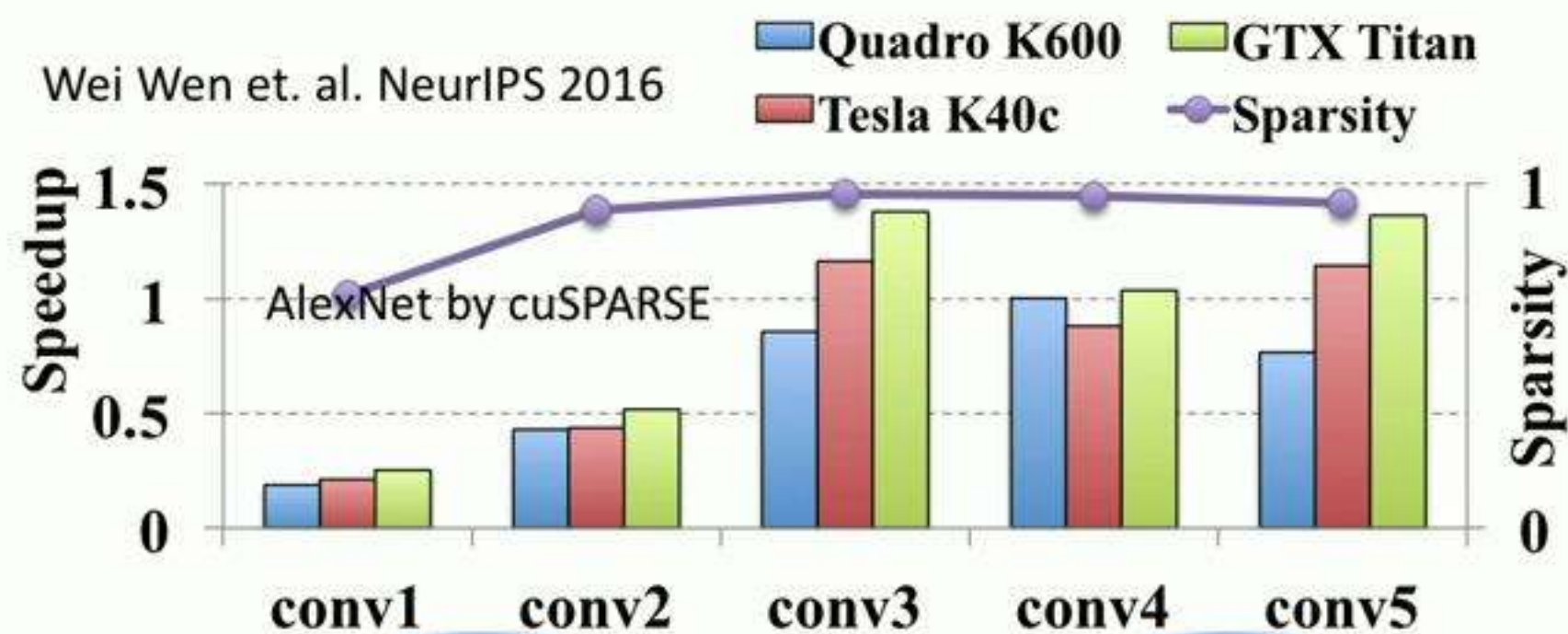
Inefficiency of Sparse Convolutional Neural Networks



Significantly reduces the storage size of DNNs [1]
Good speedup with customized hardware [2]

[1] S. Han et. al. NeurIPS 2015

[2] S. Han et. al. ISCA 2016.



Efficiency of Structured Sparsity

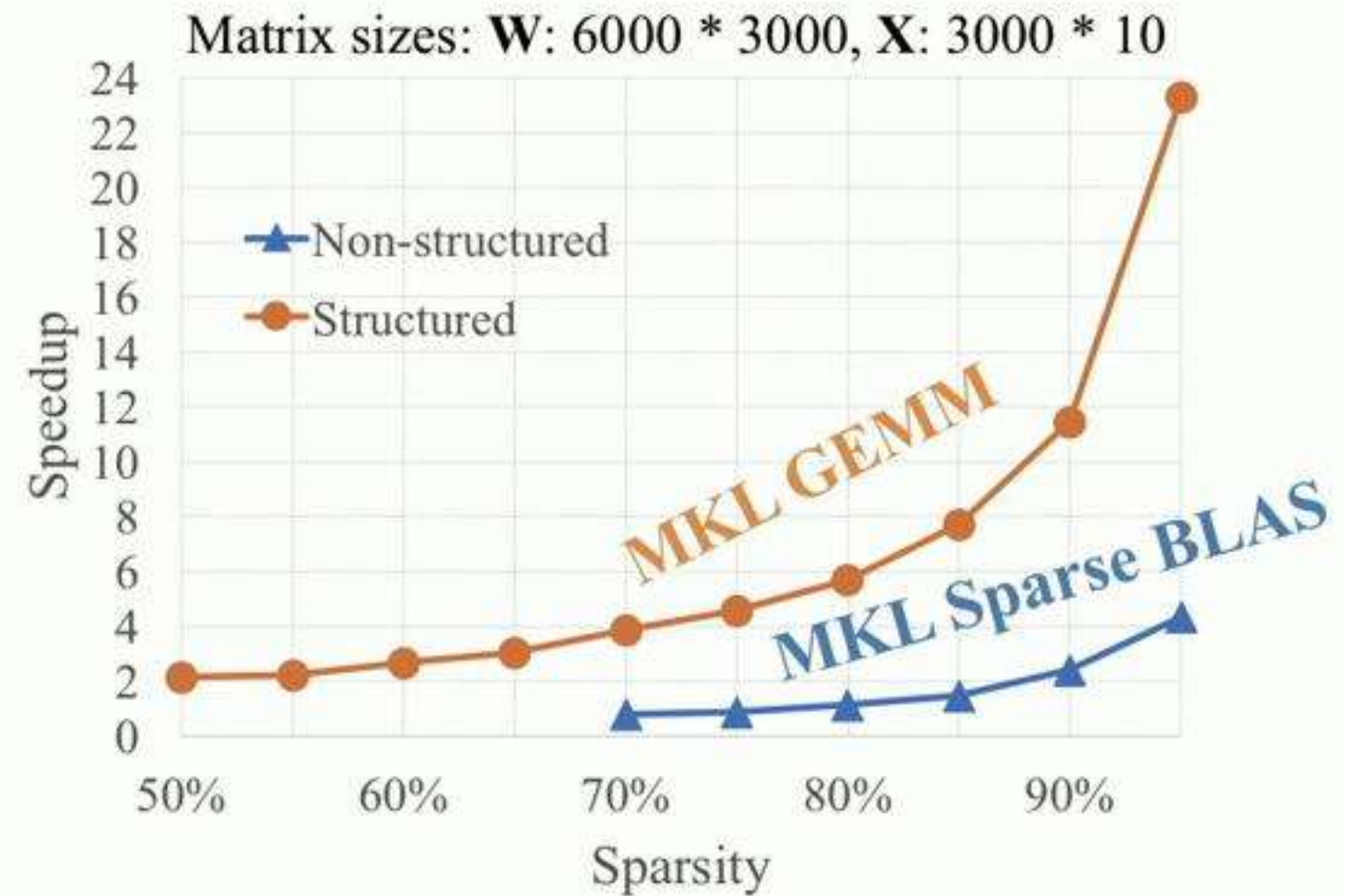
Structured sparsity



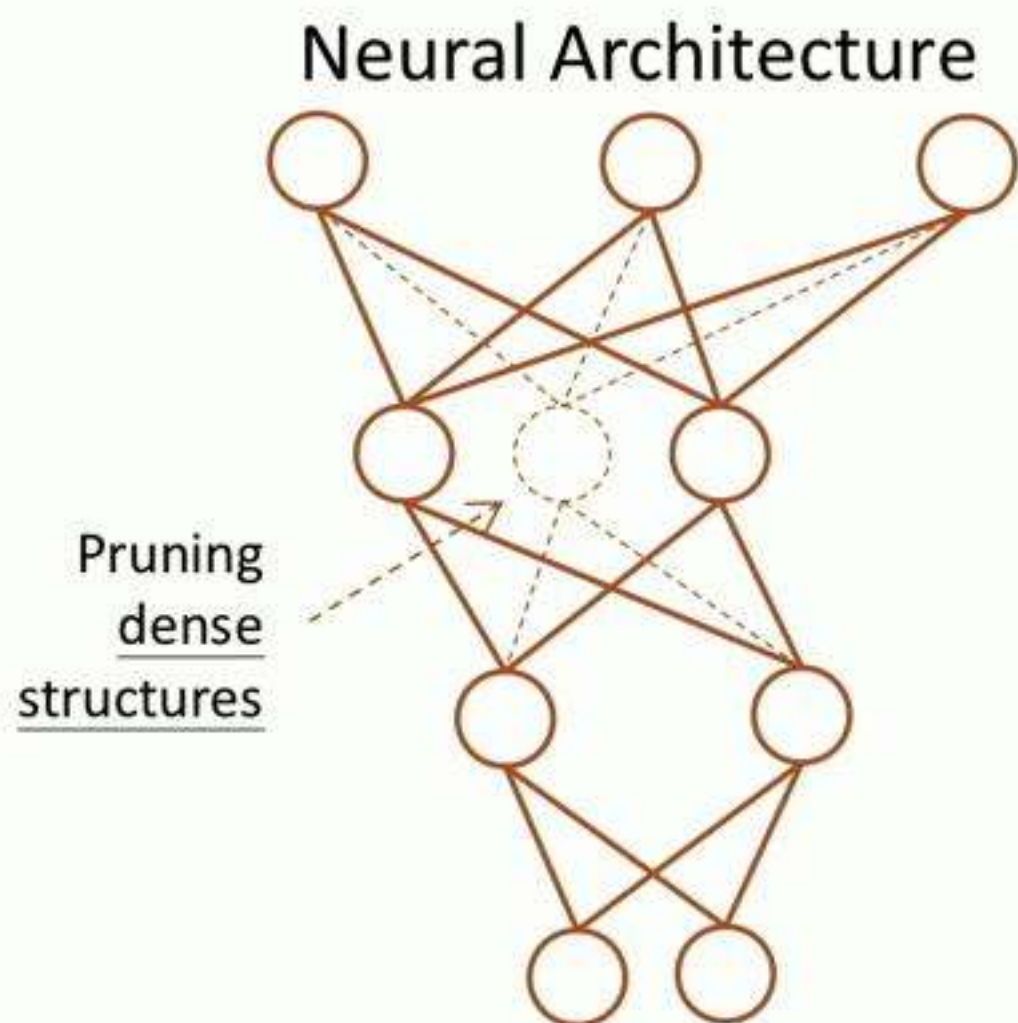
Non-structured sparsity



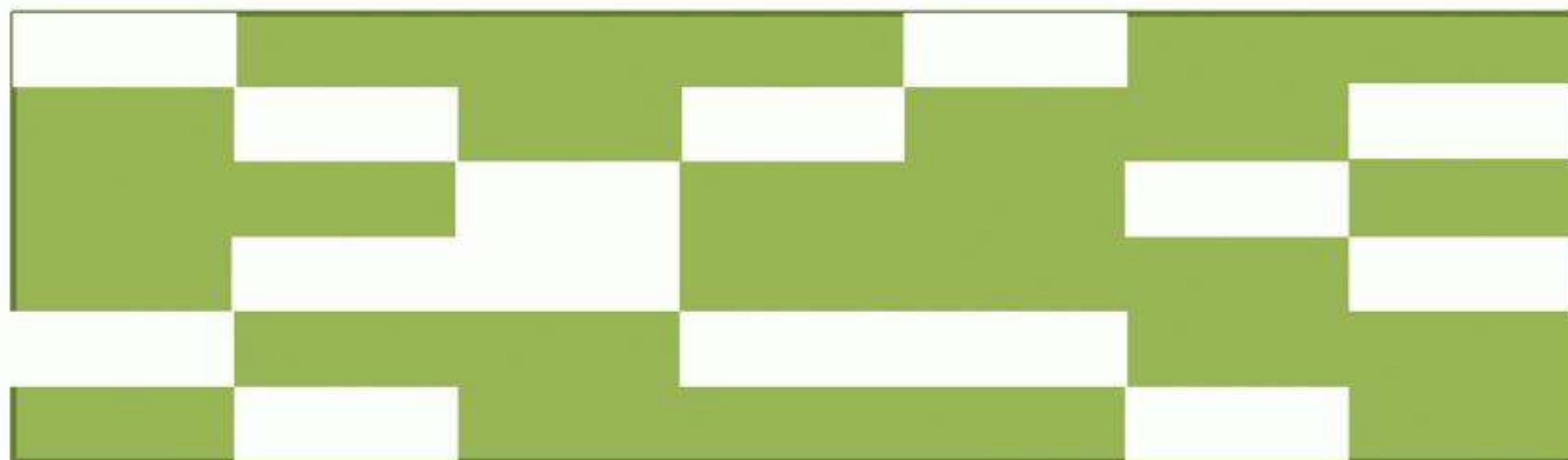
Tested on CPUs



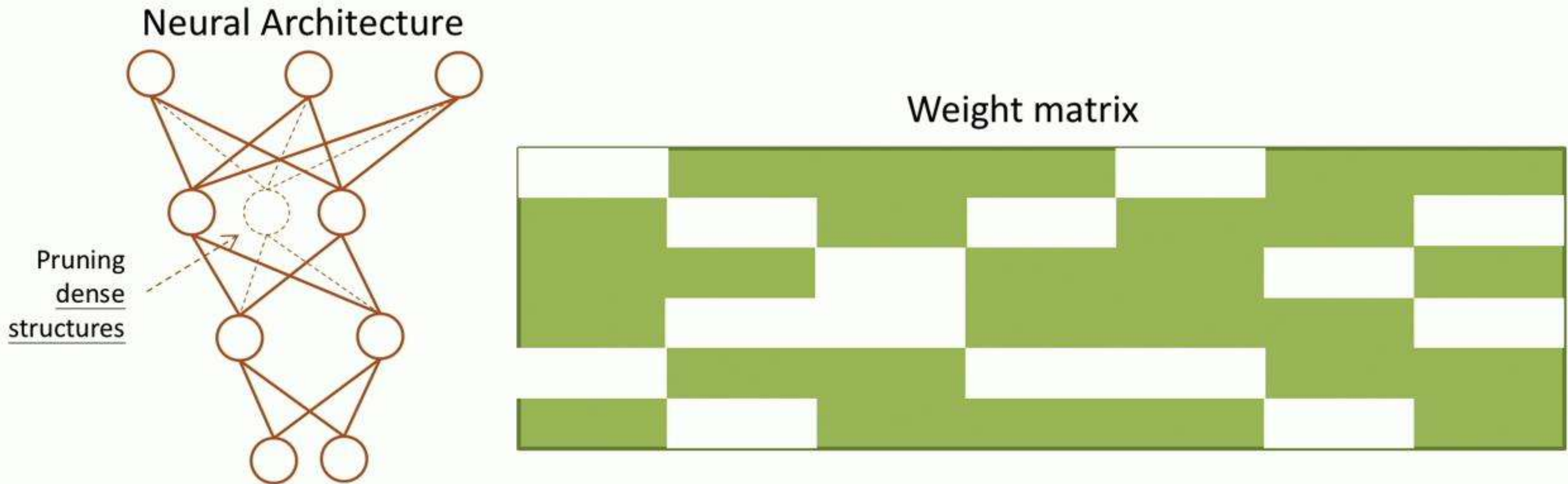
Structurally Sparse DNNs (SSDNNs)



Weight matrix



Structurally Sparse DNNs (SSDNNs)



Connections/Weights are removed group by group.

In neural architecture: a group can be a neuron, a layer, a filter, a hidden state, etc.

In weight matrix: a group can be a rectangle block, a row, a column or even the whole matrix

Group Lasso [1] Regularization is All you Need for SSDNNs

Step 1: Weights are split to G groups $\mathbf{w}^{(1..G)}$

e.g. $(w_1, w_2, w_3, w_4, w_5) \rightarrow$ group 1: $\mathbf{w}^{(1)} = (w_1, w_2, w_3)$, group 2: $\mathbf{w}^{(2)} = (w_4, w_5)$

[1] Yuan et al. 2006

Group Lasso [1] Regularization is All you Need for SSDNNs

Step 1: Weights are split to G groups $\mathbf{w}^{(1..G)}$

e.g. $(w_1, w_2, w_3, w_4, w_5) \rightarrow$ group 1: $\mathbf{w}^{(1)} = (w_1, w_2, w_3)$, group 2: $\mathbf{w}^{(2)} = (w_4, w_5)$

Step 2: Add group Lasso on each group $\mathbf{w}^{(g)}$ $\|\mathbf{w}^{(g)}\|_g = \sqrt{\sum_{i=1}^{|\mathbf{w}^{(g)}|} \left(w_i^{(g)}\right)^2}$ i.e. vector length
 $\text{sqrt}(w_1^2 + w_2^2 + w_3^2), \text{sqrt}(w_4^2 + w_5^2)$

[1] Yuan et al. 2006

Group Lasso [1] Regularization is All you Need for SSDNNs

Step 1: Weights are split to G groups $\mathbf{w}^{(1..G)}$

e.g. $(w_1, w_2, w_3, w_4, w_5) \rightarrow$ group 1: $\mathbf{w}^{(1)} = (w_1, w_2, w_3)$, group 2: $\mathbf{w}^{(2)} = (w_4, w_5)$

Step 2: Add group Lasso on each group $\mathbf{w}^{(g)}$ $\|\mathbf{w}^{(g)}\|_g = \sqrt{\sum_{i=1}^{|\mathbf{w}^{(g)}|} \left(w_i^{(g)}\right)^2}$ i.e. vector length

$\text{sqrt}(w_1^2 + w_2^2 + w_3^2), \text{sqrt}(w_4^2 + w_5^2)$

Step 3: Sum group Lasso over all groups as a regularization: $R_g(\mathbf{w}) = \sum_{g=1}^G \|\mathbf{w}^{(g)}\|_g,$

$\text{sqrt}(w_1^2 + w_2^2 + w_3^2) + \text{sqrt}(w_4^2 + w_5^2)$

[1] Yuan et al. 2006

Group Lasso [1] Regularization is All you Need for SSDNNs

Step 1: Weights are split to G groups $\mathbf{w}^{(1..G)}$

e.g. $(w_1, w_2, w_3, w_4, w_5) \rightarrow$ group 1: $\mathbf{w}^{(1)} = (w_1, w_2, w_3)$, group 2: $\mathbf{w}^{(2)} = (w_4, w_5)$

Step 2: Add group Lasso on each group $\mathbf{w}^{(g)}$ $\|\mathbf{w}^{(g)}\|_g = \sqrt{\sum_{i=1}^{|\mathbf{w}^{(g)}|} \left(w_i^{(g)}\right)^2}$ i.e. vector length

$\text{sqrt}(w_1^2 + w_2^2 + w_3^2), \text{sqrt}(w_4^2 + w_5^2)$

Step 3: Sum group Lasso over all groups as a regularization: $R_g(\mathbf{w}) = \sum_{g=1}^G \|\mathbf{w}^{(g)}\|_g,$

$\text{sqrt}(w_1^2 + w_2^2 + w_3^2) + \text{sqrt}(w_4^2 + w_5^2)$

Step 4: SGD optimizing $\arg \min_{\mathbf{w}} \{E(\mathbf{w})\} = \arg \min_{\mathbf{w}} \{E_D(\mathbf{w}) + \lambda_g \cdot R_g(\mathbf{w})\}$

[1] Yuan et al. 2006

Group Lasso [1] Regularization is All you Need for SSDNNs

Step 1: Weights are split to G groups $\mathbf{w}^{(1..G)}$

e.g. $(w_1, w_2, w_3, w_4, w_5) \rightarrow$ group 1: $\mathbf{w}^{(1)} = (w_1, w_2, w_3)$, group 2: $\mathbf{w}^{(2)} = (w_4, w_5)$

Step 2: Add group Lasso on each group $\mathbf{w}^{(g)}$ $\|\mathbf{w}^{(g)}\|_g = \sqrt{\sum_{i=1}^{|\mathbf{w}^{(g)}|} \left(w_i^{(g)}\right)^2}$ i.e. vector length

$\text{sqrt}(w_1^2 + w_2^2 + w_3^2), \text{sqrt}(w_4^2 + w_5^2)$

Step 3: Sum group Lasso over all groups as a regularization: $R_g(\mathbf{w}) = \sum_{g=1}^G \|\mathbf{w}^{(g)}\|_g,$

$\text{sqrt}(w_1^2 + w_2^2 + w_3^2) + \text{sqrt}(w_4^2 + w_5^2)$

Step 4: SGD optimizing $\arg \min_{\mathbf{w}} \{E(\mathbf{w})\} = \arg \min_{\mathbf{w}} \{E_D(\mathbf{w}) + \lambda_g \cdot R_g(\mathbf{w})\}$

We refer to our method as **Structured Sparsity Learning (SSL)**

[1] Yuan et al. 2006

Structured Sparsity Learning (SSL)

- Why can SSL learn to remove weights group by group?

The gradient-based explanation:

$$\boxed{\mathbf{w}^{(g)}} \leftarrow \mathbf{w}^{(g)} - \eta \left(\boxed{\frac{\partial E_D(\mathbf{w})}{\partial \mathbf{w}^{(g)}}} + \lambda_g \frac{\mathbf{w}^{(g)}}{\|\mathbf{w}^{(g)}\|_2} \right)$$

A group of
weights

Regular gradients
to minimize error



Structured Sparsity Learning (SSL)

- Why can SSL learn to remove weights group by group?

The gradient-based explanation:

$$\mathbf{w}^{(g)}$$

A group of
weights

$$\leftarrow \mathbf{w}^{(g)} - \eta \left(\frac{\partial E_D(\mathbf{w})}{\partial \mathbf{w}^{(g)}} + \lambda_g \frac{\mathbf{w}^{(g)}}{\|\mathbf{w}^{(g)}\|_2} \right)$$

Regular gradients
to minimize error

$$\frac{\mathbf{w}^{(g)}}{\|\mathbf{w}^{(g)}\|_2}$$

Additional gradients
to learn sparsity

A diagram illustrating the additional gradient term. A blue vector labeled $\mathbf{w}^{(g)}$ points upwards and to the right. A shorter orange vector labeled $-\mathbf{w}^{(g)} / \|\mathbf{w}^{(g)}\|_2$ points downwards and to the left, representing the negative normalized vector.

Structured Sparsity Learning (SSL)

- Why can SSL learn to remove weights group by group?

The gradient-based explanation:

$$\mathbf{w}^{(g)}$$

A group of weights

$$\leftarrow \mathbf{w}^{(g)} - \eta \left(\frac{\partial E_D(\mathbf{w})}{\partial \mathbf{w}^{(g)}} + \lambda_g \frac{\mathbf{w}^{(g)}}{\|\mathbf{w}^{(g)}\|_2} \right)$$

Regular gradients to minimize error

$$\frac{\mathbf{w}^{(g)}}{\|\mathbf{w}^{(g)}\|_2}$$

Additional gradients to learn sparsity

$$-\mathbf{w}^{(g)} / \|\mathbf{w}^{(g)}\|_2$$

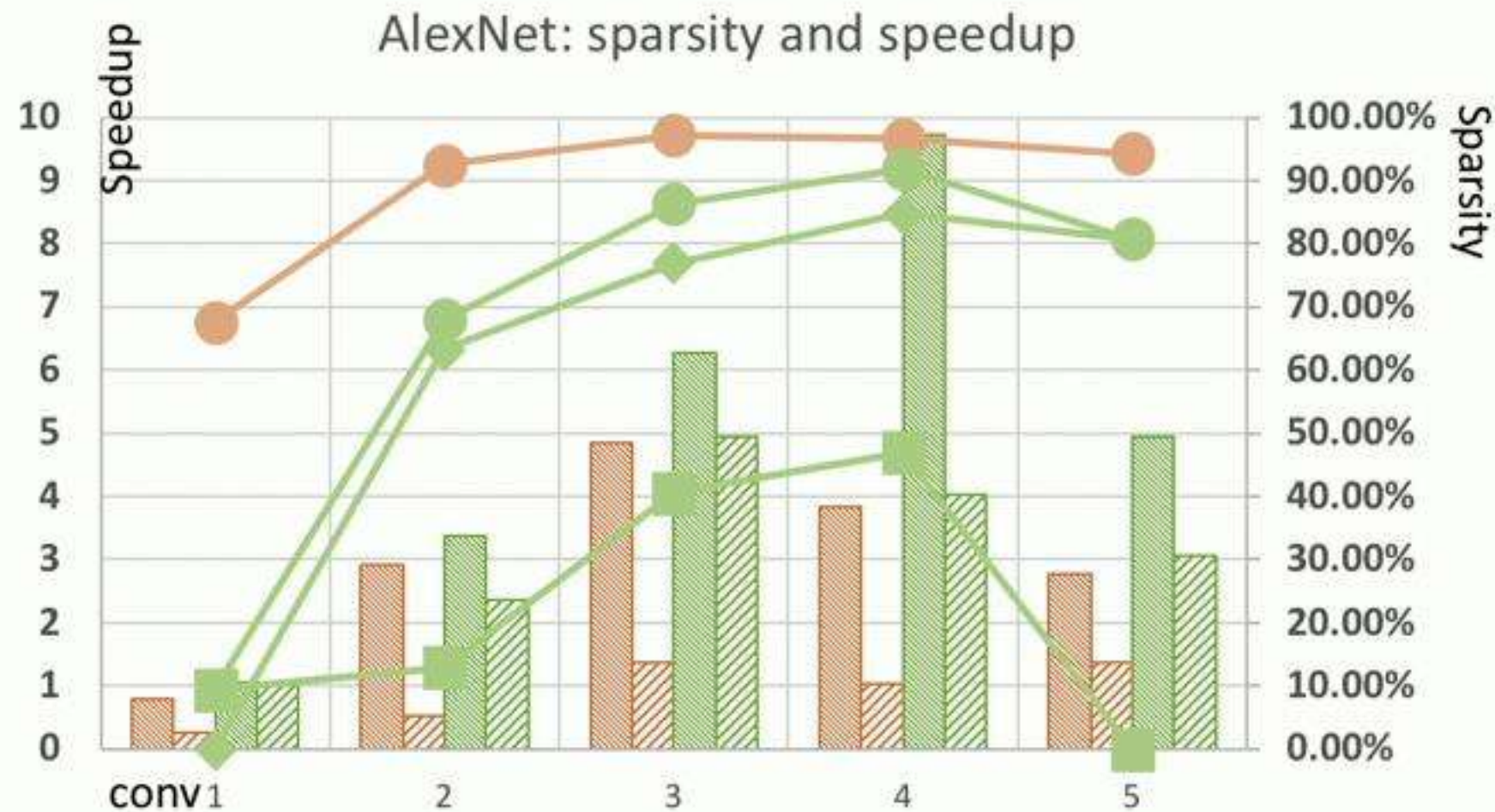
Many groups can be pushed to zeros

Structurally Sparse AlexNet

Structured Sparsity (removing columns and rows):



Non-structured sparsity



Structured Sparsity Learning (SSL)

- Why can SSL learn to remove weights group by group?

The gradient-based explanation:

$$\mathbf{w}^{(g)}$$

A group of weights

$$\leftarrow \mathbf{w}^{(g)}$$

$$-$$

$$\eta \left(\frac{\partial E_D(\mathbf{w})}{\partial \mathbf{w}^{(g)}} \right)$$

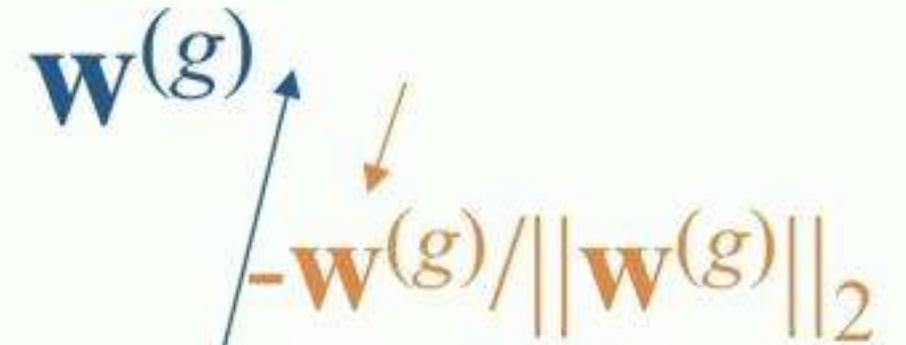
Regular gradients to minimize error

$$+$$

$$\lambda_g$$

$$\frac{\mathbf{w}^{(g)}}{\|\mathbf{w}^{(g)}\|_2}$$

Additional gradients to learn sparsity



$\mathbf{w}^{(g)}$

$-\mathbf{w}^{(g)} / \|\mathbf{w}^{(g)}\|_2$

Many groups can be pushed to zeros

Structurally Sparse *AlexNet*

Structured Sparsity (removing columns and rows):



Non-structured sparsity

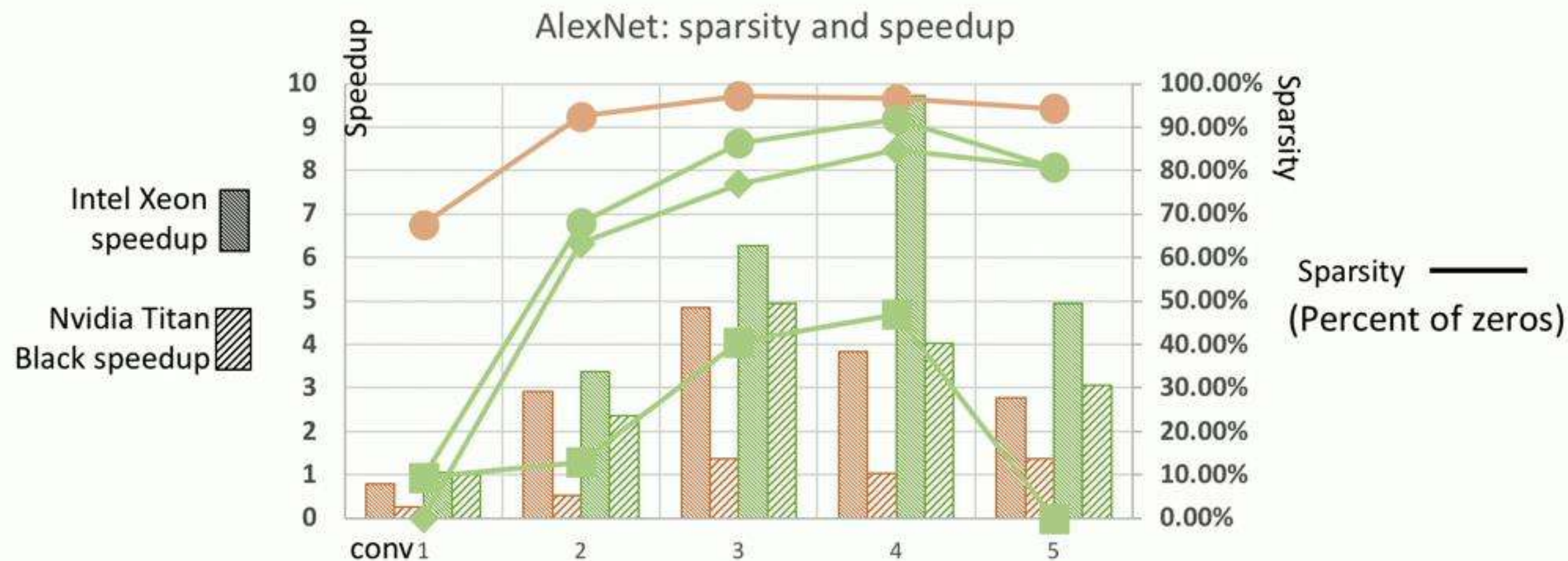


Structurally Sparse AlexNet

Structured Sparsity (removing columns and rows):



Non-structured sparsity

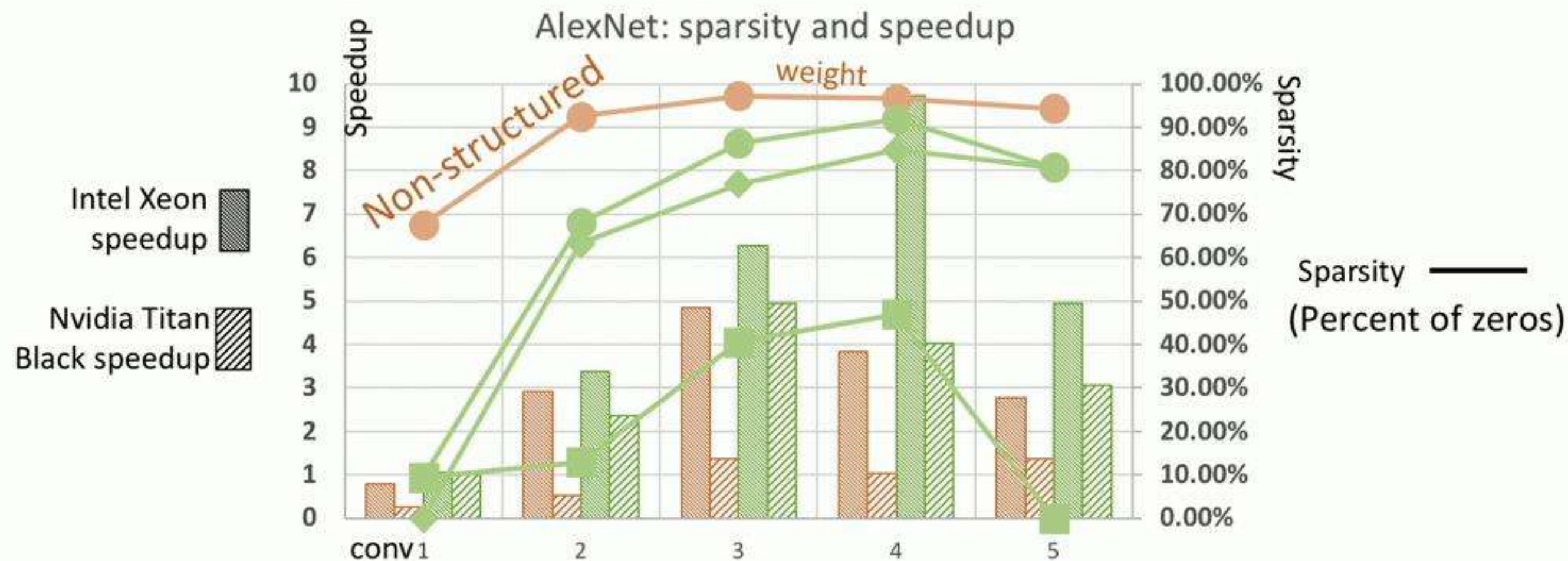


Structurally Sparse AlexNet

Structured Sparsity (removing columns and rows):



Non-structured sparsity

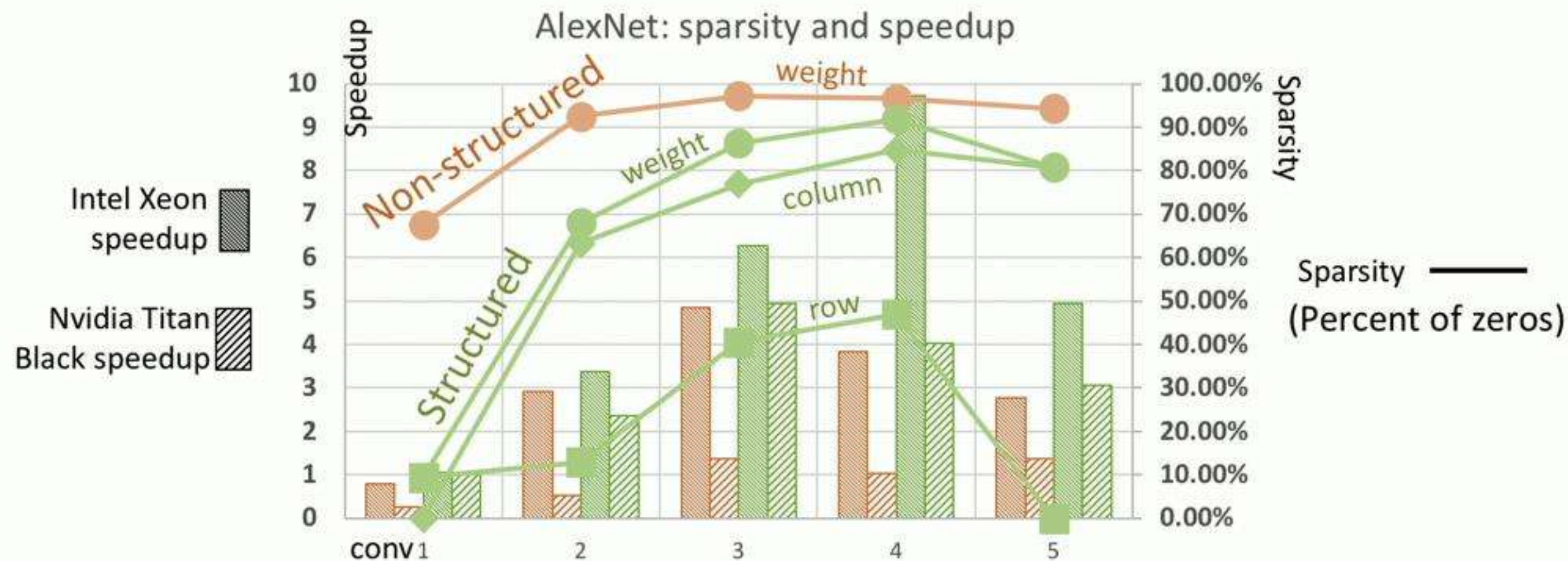


Structurally Sparse AlexNet

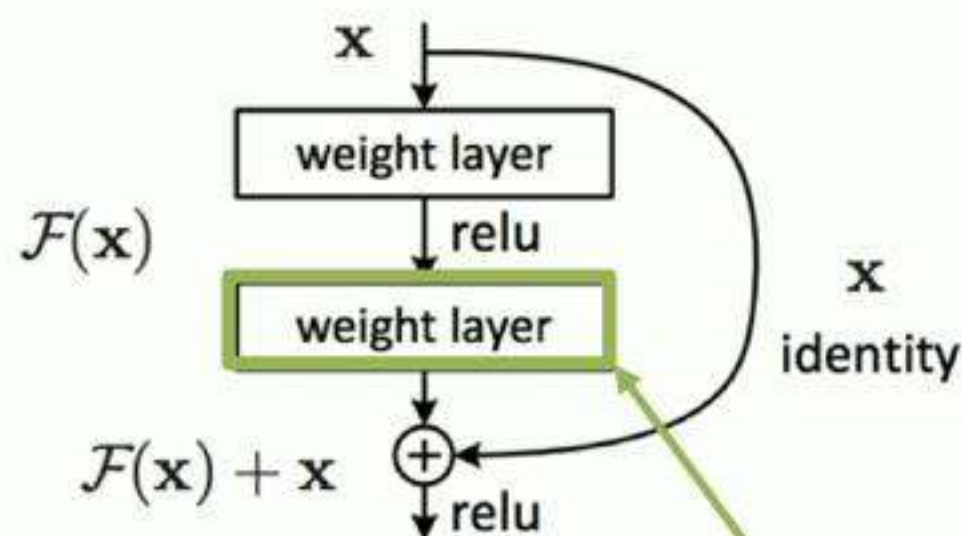
Structured Sparsity (removing columns and rows):



Non-structured sparsity



Structurally Sparse *ResNets* – Removing Layers

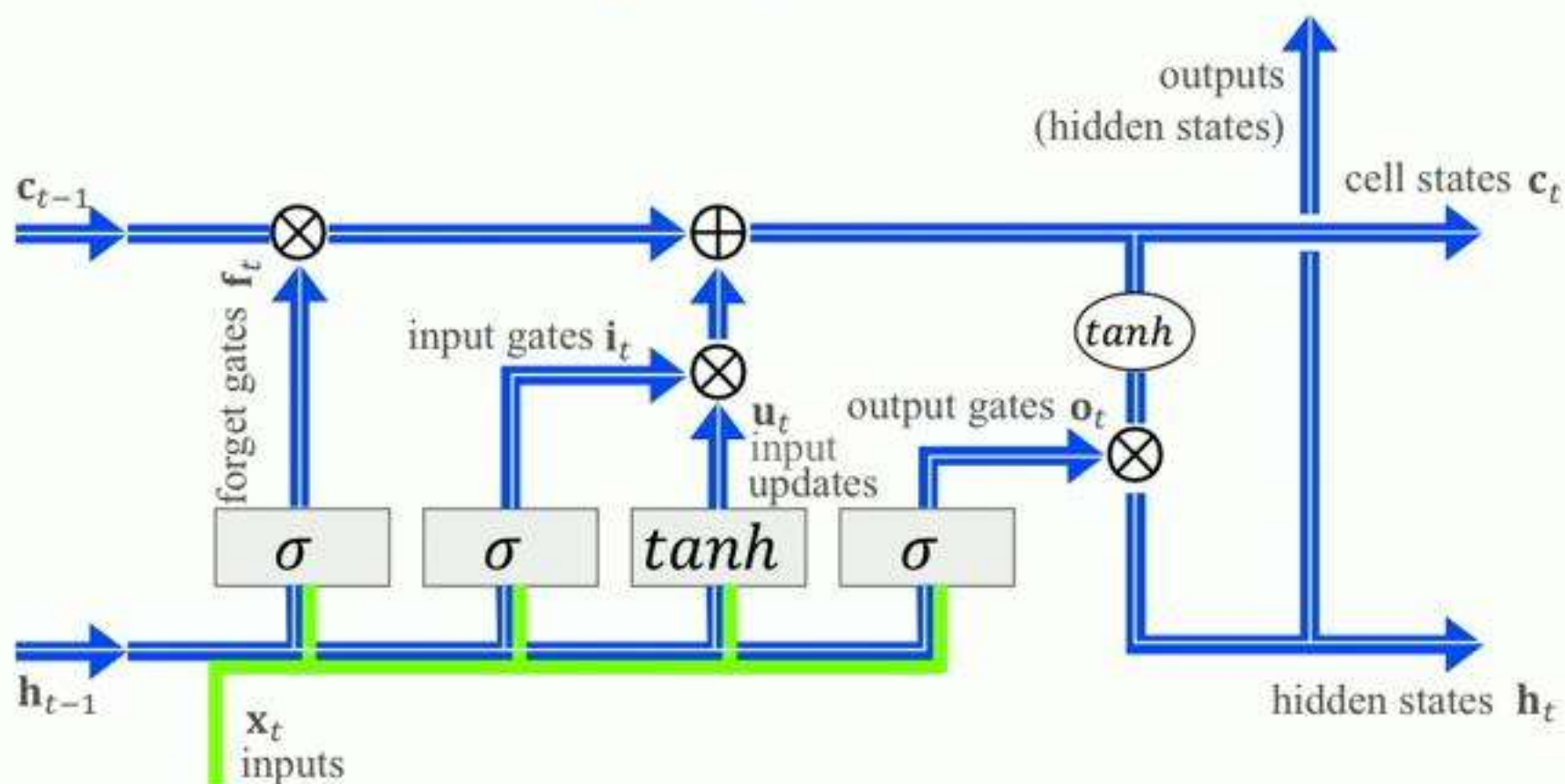


K. He et al. CVPR 2016.

One layer is
one group.

	# layers	error
ResNet	20	8.82%
SSL-ResNet	14	8.54%
Reduced	6/20	
ResNet	32	7.51%
SSL-ResNet	18	7.40%
Reduced	14/32	

Structurally Sparse LSTMs: Removing Hidden Structures



$$\mathbf{i}_t = \sigma(\mathbf{x}_t \cdot \mathbf{W}_{xi} + \mathbf{h}_{t-1} \cdot \mathbf{W}_{hi} + \mathbf{b}_i)$$

$$\mathbf{f}_t = \sigma(\mathbf{x}_t \cdot \mathbf{W}_{xf} + \mathbf{h}_{t-1} \cdot \mathbf{W}_{hf} + \mathbf{b}_f)$$

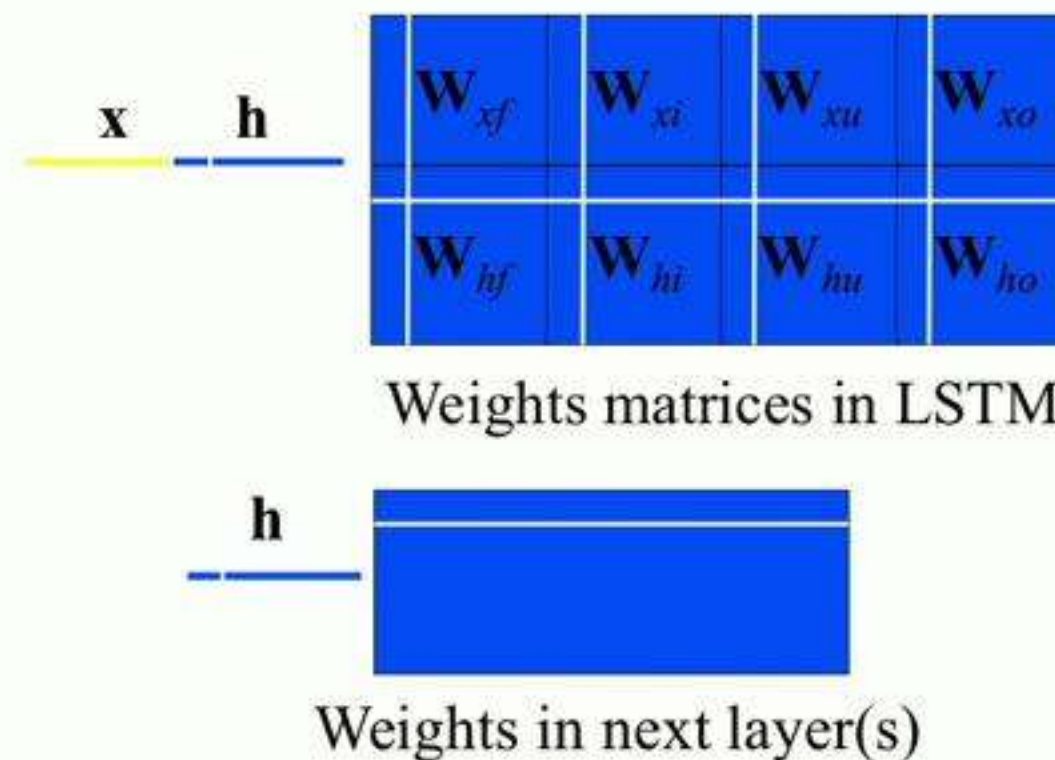
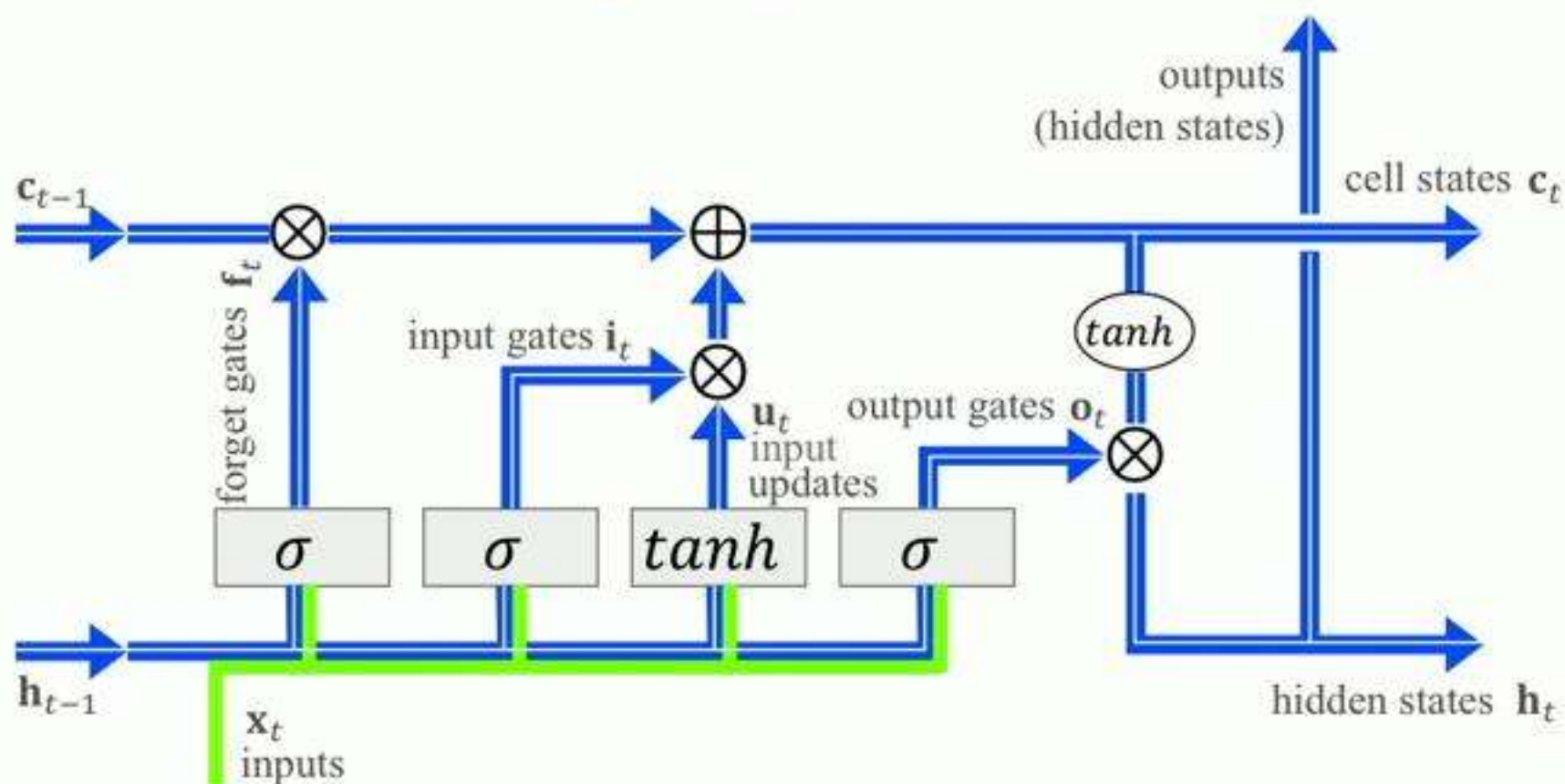
$$\mathbf{o}_t = \sigma(\mathbf{x}_t \cdot \mathbf{W}_{xo} + \mathbf{h}_{t-1} \cdot \mathbf{W}_{ho} + \mathbf{b}_o)$$

$$\mathbf{u}_t = \tanh(\mathbf{x}_t \cdot \mathbf{W}_{xu} + \mathbf{h}_{t-1} \cdot \mathbf{W}_{hu} + \mathbf{b}_u)$$

$$\mathbf{c}_t = \mathbf{f}_t \odot \mathbf{c}_{t-1} + \mathbf{i}_t \odot \mathbf{u}_t$$

$$\mathbf{h}_t = \mathbf{o}_t \odot \tanh(\mathbf{c}_t)$$

Structurally Sparse LSTMs: Removing Hidden Structures



$$\begin{aligned} i_t &= \sigma(x_t \cdot W_{xi} + h_{t-1} \cdot W_{hi} + b_i) \\ f_t &= \sigma(x_t \cdot W_{xf} + h_{t-1} \cdot W_{hf} + b_f) \\ o_t &= \sigma(x_t \cdot W_{xo} + h_{t-1} \cdot W_{ho} + b_o) \\ u_t &= \tanh(x_t \cdot W_{xu} + h_{t-1} \cdot W_{hu} + b_u) \\ c_t &= f_t \odot c_{t-1} + i_t \odot u_t \\ h_t &= o_t \odot \tanh(c_t) \end{aligned}$$

White strips = one group of weights in SSL

Removing one group = reducing hidden size by one

Structurally Sparse *LSTMs*

Method	Test perplexity	Hidden size of (1 st , 2 nd) LSTMs	Real-time speedup
Zaremba et. al. 2014	78.57	(1500, 1500)	1.00X
Ours (sparsifying Zaremba et. al. 2014)	78.64	(373, 315)	10.59X
	76.03	(381, 535)	7.10X
Train a smaller one	85.66	(373, 315)	10.59X

Structurally Sparse *LSTMs*

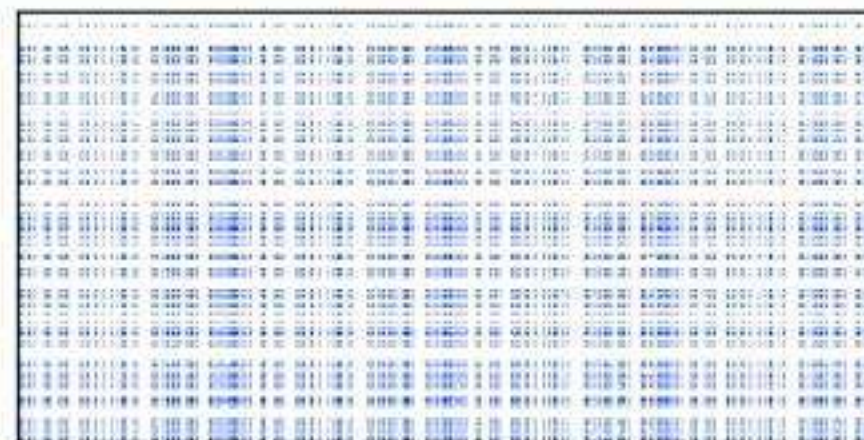
Method	Test perplexity	Hidden size of (1 st , 2 nd) LSTMs	Real-time speedup
Zaremba et. al. 2014	78.57	(1500, 1500)	1.00X
Ours (sparsifying Zaremba et. al. 2014)	78.64	(373, 315)	10.59X
	76.03	(381, 535)	7.10X
Train a smaller one	85.66	(373, 315)	10.59X

1. Reduce size while maintaining perplexity

LSTM 1



LSTM 2



Output



Learned structurally sparse LSTMs

Structurally Sparse *LSTMs*

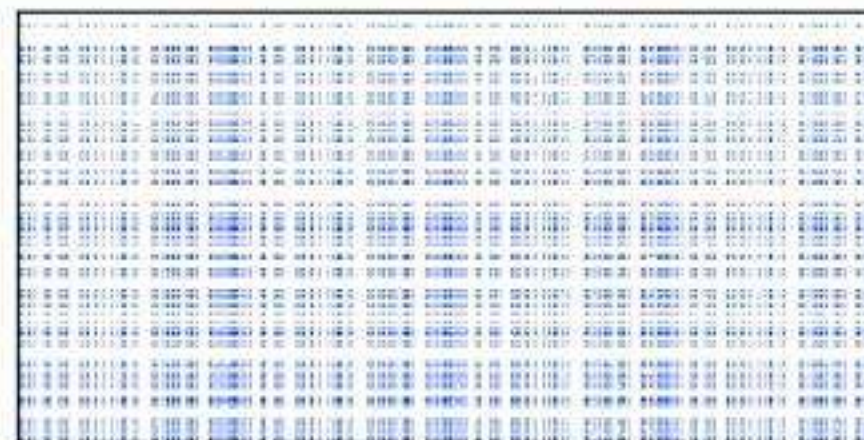
Method	Test perplexity	Hidden size of (1 st , 2 nd) LSTMs	Real-time speedup
Zaremba et. al. 2014	78.57	(1500, 1500)	1.00X
Ours (sparsifying Zaremba et. al. 2014)	78.64	(373, 315)	10.59X
	76.03	(381, 535)	7.10X
Train a smaller one	85.66	(373, 315)	10.59X

1. Reduce size while maintaining perplexity
2. Achieve lower perplexity
3. Better than training a smaller one

LSTM 1



LSTM 2



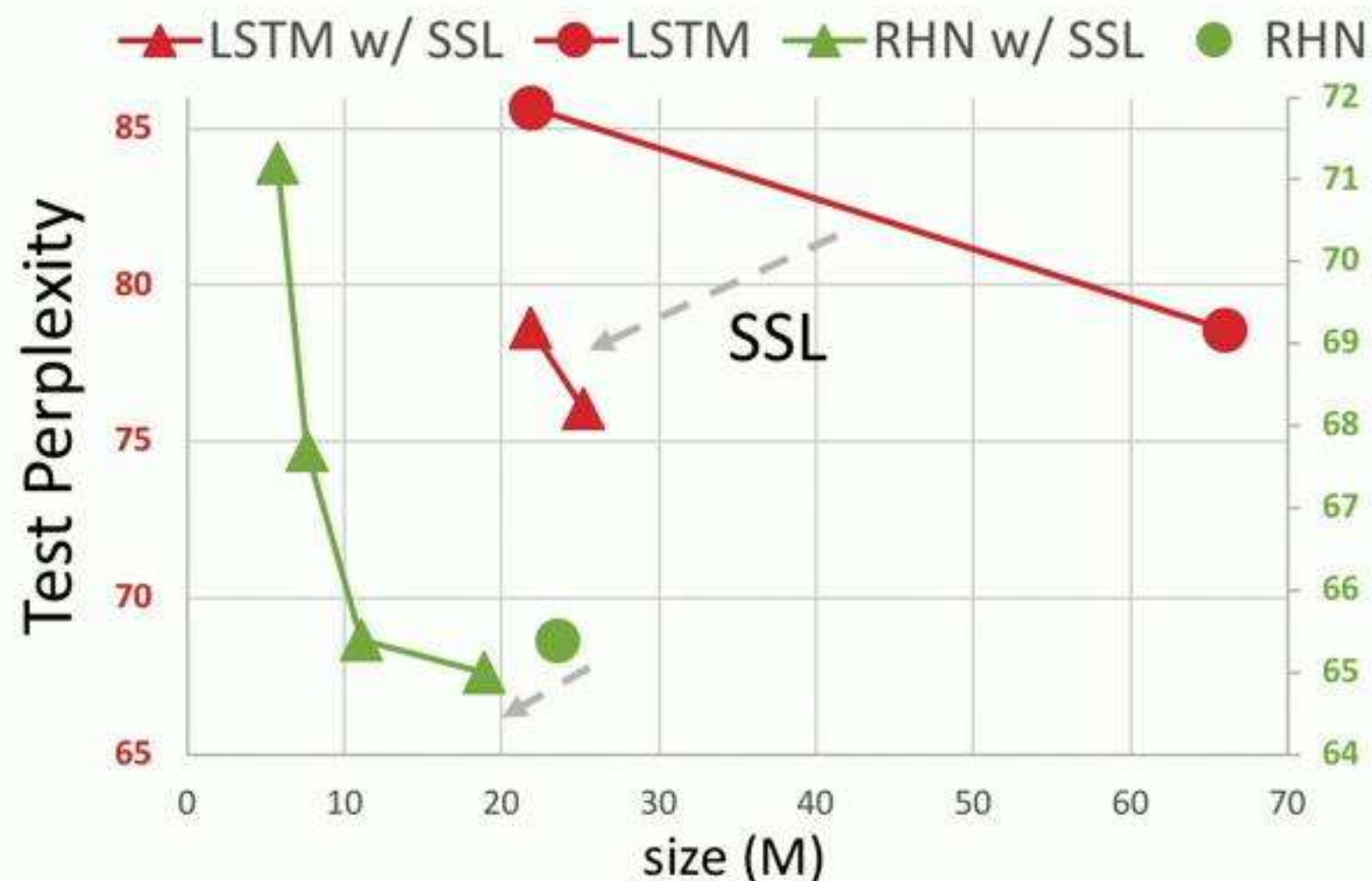
Output



Learned structurally sparse LSTMs

Generalization vs. Model Complexity under SSL

Language Modeling



☺ Start from redundant models and sparsify by SSL

RHN: Zilly et al. 2017. BiDAF: Seo et al. 2016

Structurally Sparse *LSTMs*

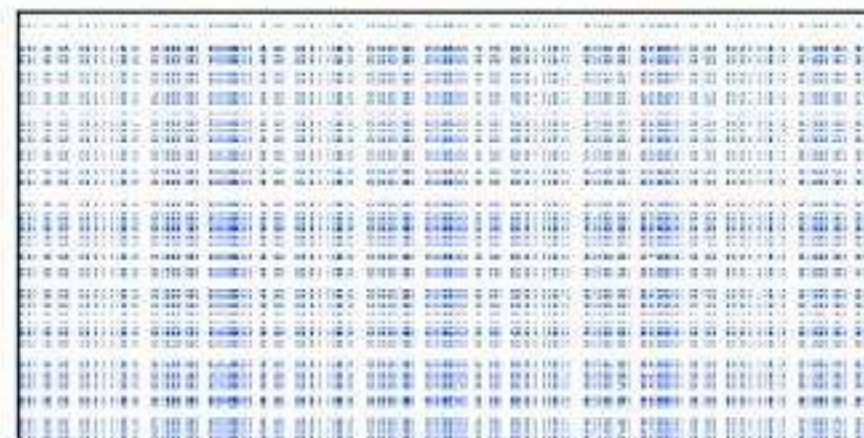
Method	Test perplexity	Hidden size of (1 st , 2 nd) LSTMs	Real-time speedup
Zaremba et. al. 2014	78.57	(1500, 1500)	1.00X
Ours (sparsifying Zaremba et. al. 2014)	78.64	(373, 315)	10.59X
	76.03	(381, 535)	7.10X
Train a smaller one	85.66	(373, 315)	10.59X

1. Reduce size while maintaining perplexity
2. Achieve lower perplexity
3. Better than training a smaller one

LSTM 1



LSTM 2



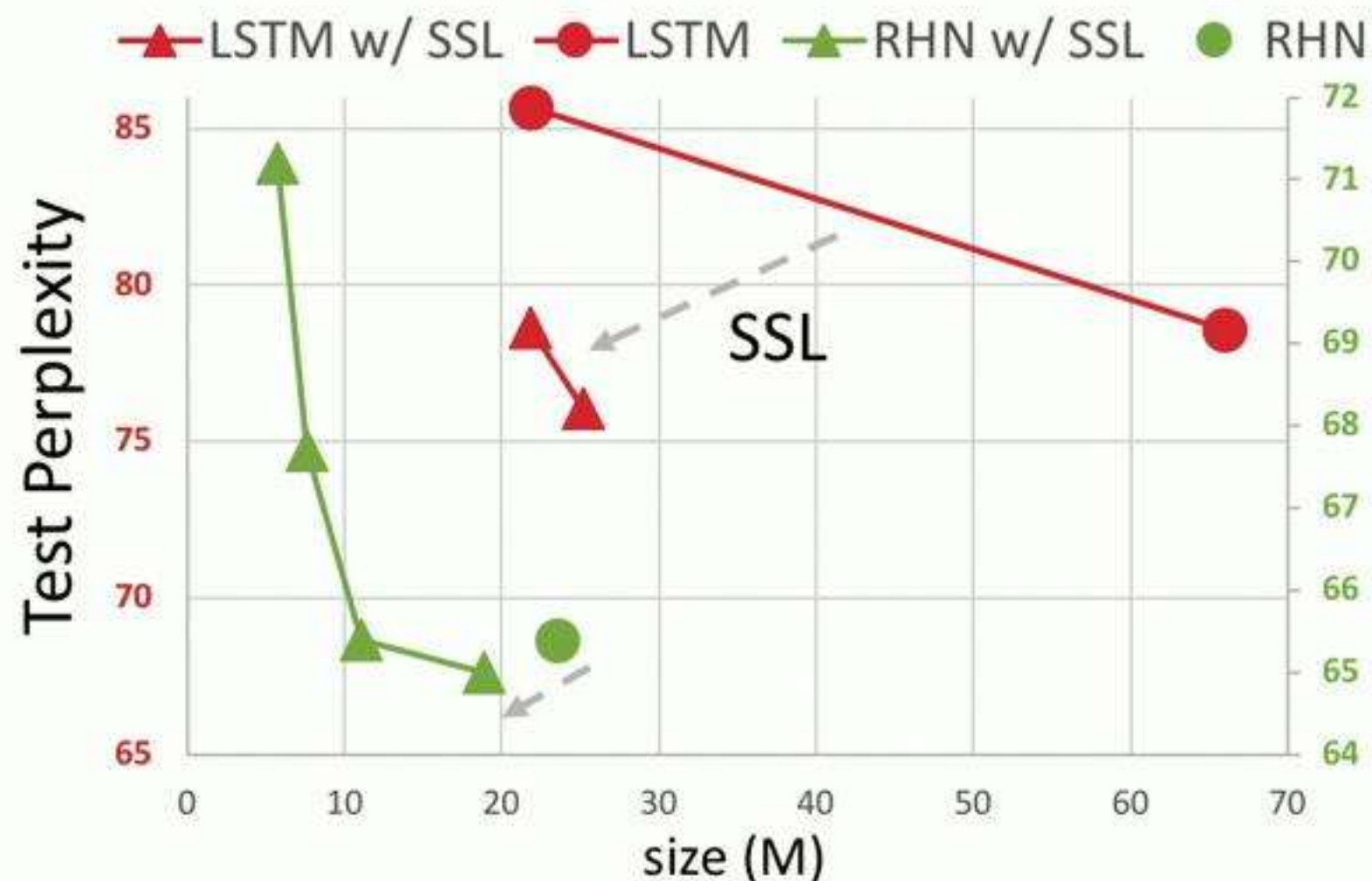
Output



Learned structurally sparse LSTMs

Generalization vs. Model Complexity under SSL

Language Modeling

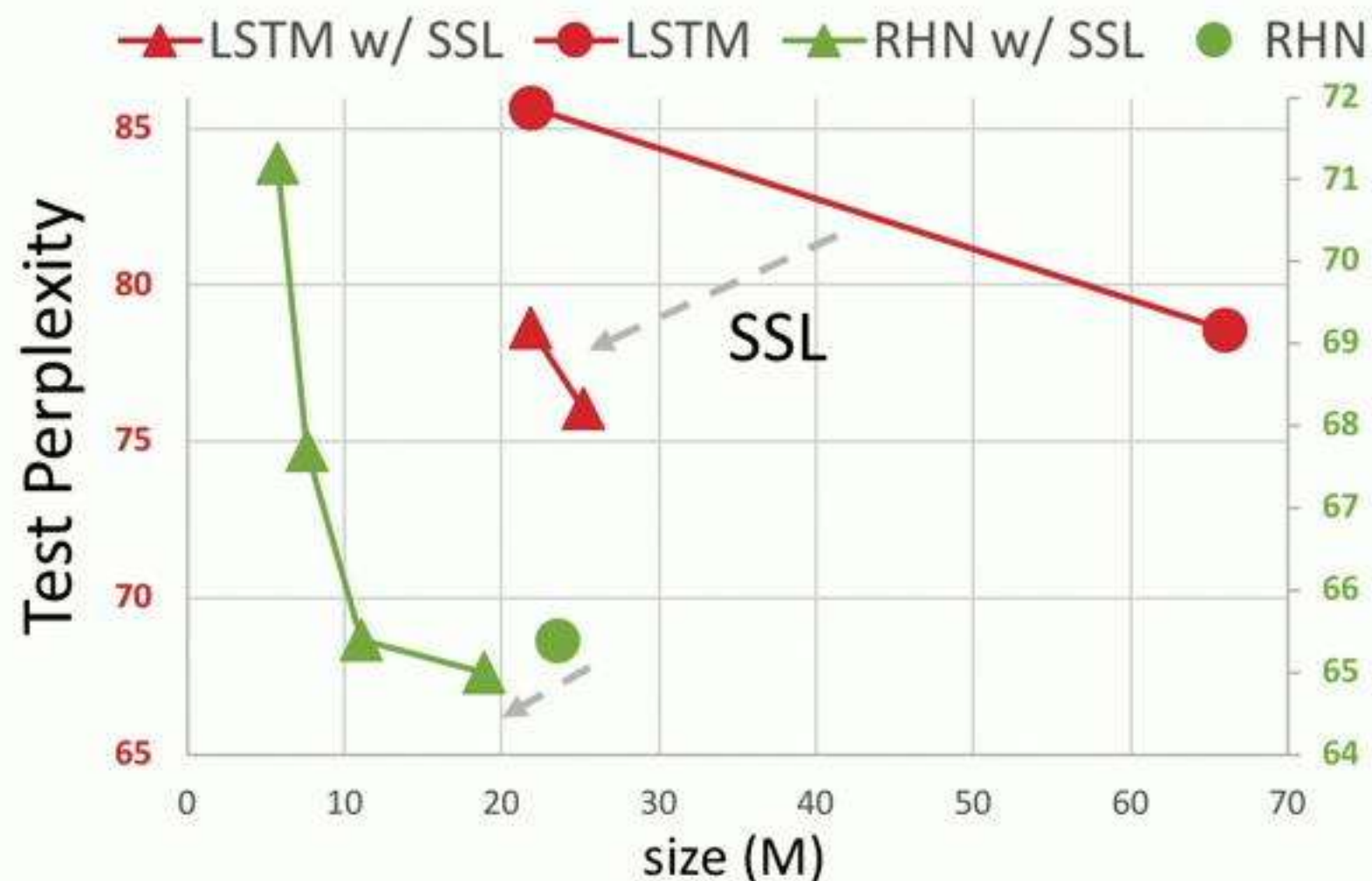


😊 Start from redundant models and sparsify by SSL

RHN: Zilly et al. 2017. BiDAF: Seo et al. 2016

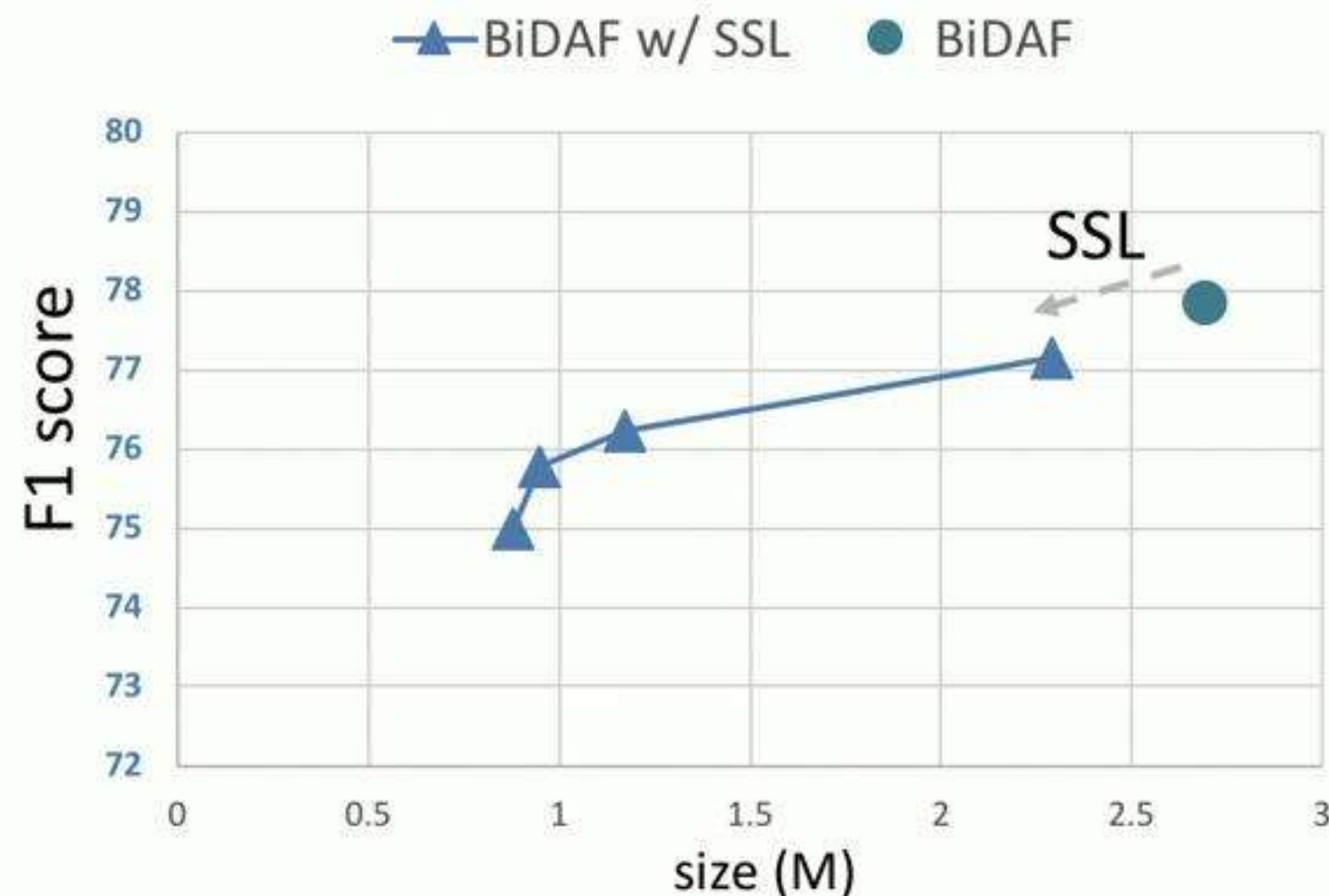
Generalization vs. Model Complexity under SSL

Language Modeling



☺ Start from redundant models and sparsify by SSL

Question Answering



☺ Trade-off for compact models with SSL

RHN: Zilly et al. 2017. BiDAF: Seo et al. 2016

Summary of Previous Research

- More scalable distributed deep learning with stochastic low-precision gradients
 - *Effectiveness verified by AI productions*
 - *Can reduce research and production cycle by reducing evaluation/training time of each new model*
 - *More scalable to train larger models*
- Sparse deep neural networks
 - *Enable more ubiquitous AI on edge devices (such as mobile phones, drones, self-driving cars, VR devices, etc)*
 - *Performance gain by starting from a large model and sparsify it down*

Future Research

- Connections to previous research and going beyond
- Scaling Natural Language Processing
- Automated machine learning

SSL Implications and Beyond

Lottery Ticket Hypothesis

(Frankle et al 2019)

“dense, randomly-initialized, feed-forward networks contain subnetworks (winning tickets) that—when trained in isolation— reach test accuracy comparable to the original network in a similar number of iterations. ”

SSL identifies winning tickets in the early stage?

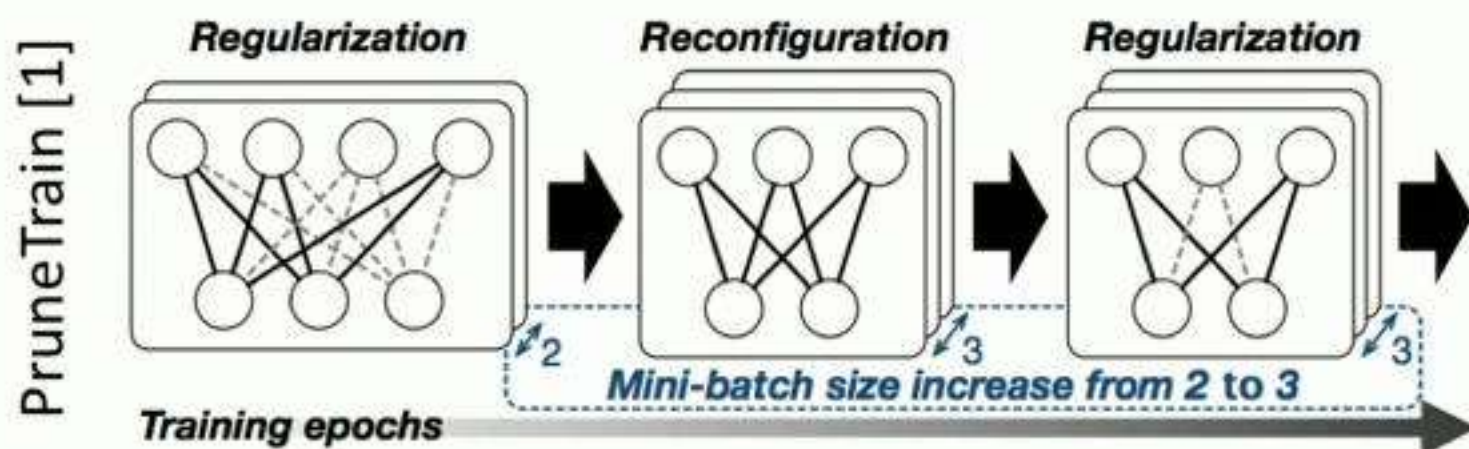
SSL Implications and Beyond

Lottery Ticket Hypothesis

(Frankle et al 2019)

“dense, randomly-initialized, feed-forward networks contain subnetworks (winning tickets) that—when trained in isolation— reach test accuracy comparable to the original network in a similar number of iterations. ”

SSL identifies winning tickets in the early stage?



39% training time reduction of ResNet50 for ImageNet.

[1] Lym, S, E Choukse, S Zangeneh, **W Wen**, M Erez, and S Shanghavi. "PruneTrain: Gradual Structured Pruning from Scratch for Faster Neural Network Training." *International Conference for High Performance Computing, Networking, Storage and Analysis (SC)* (2019).

Best Student Paper Finalist.

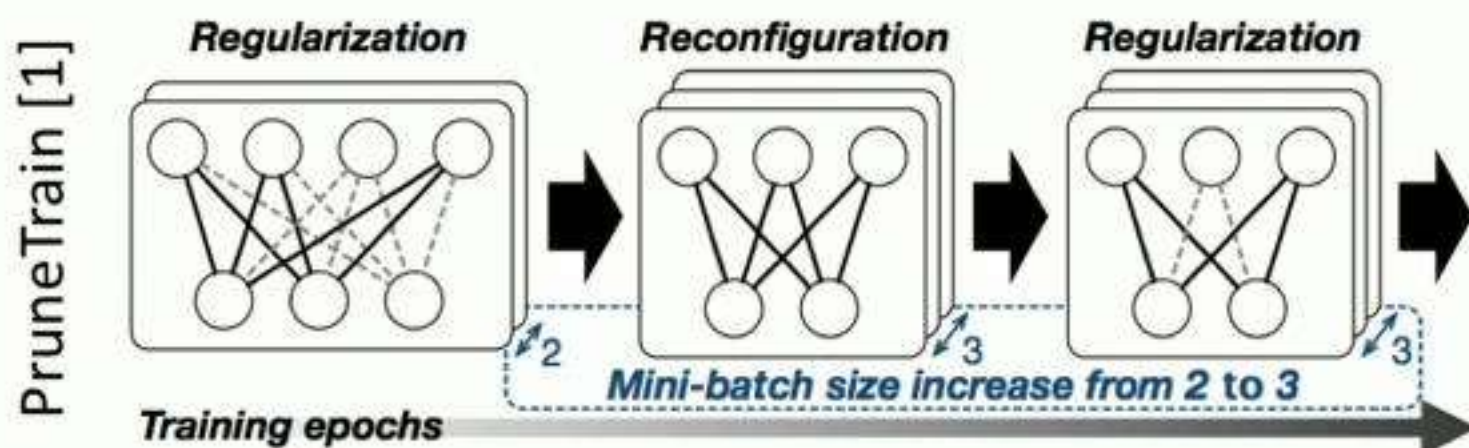
SSL Implications and Beyond

Lottery Ticket Hypothesis

(Frankle et al 2019)

“dense, randomly-initialized, feed-forward networks contain subnetworks (winning tickets) that—when trained in isolation— reach test accuracy comparable to the original network in a similar number of iterations.”

SSL identifies winning tickets in the early stage?

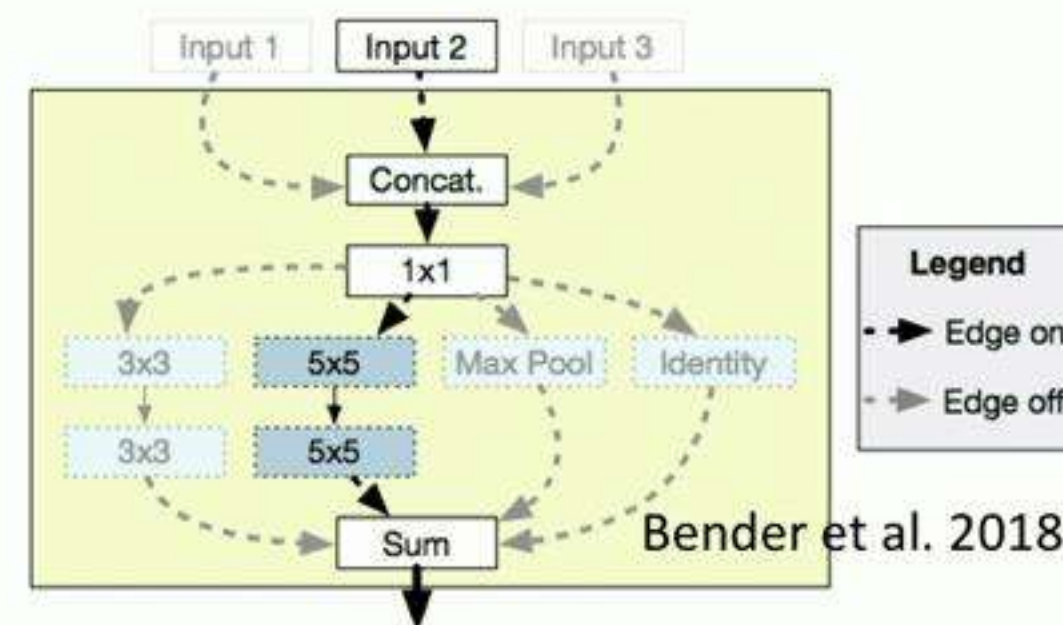


39% training time reduction of ResNet50 for ImageNet.

[1] Lym, S, E Choukse, S Zangeneh, **W Wen**, M Erez, and S Shanghavi. "PruneTrain: Gradual Structured Pruning from Scratch for Faster Neural Network Training." *International Conference for High Performance Computing, Networking, Storage and Analysis (SC)* (2019).

Best Student Paper Finalist.

AutoML with weight sharing

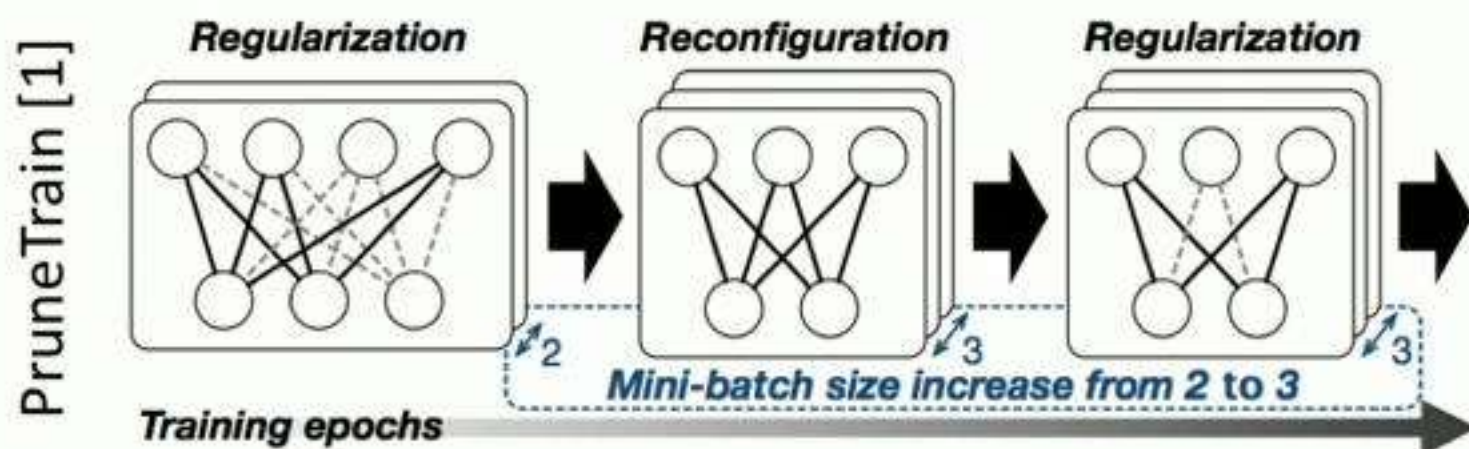


SSL Implications and Beyond

Lottery Ticket Hypothesis (Frankle et al 2019)

“dense, randomly-initialized, feed-forward networks contain subnetworks (winning tickets) that—when trained in isolation— reach test accuracy comparable to the original network in a similar number of iterations.”

SSL identifies winning tickets in the early stage?

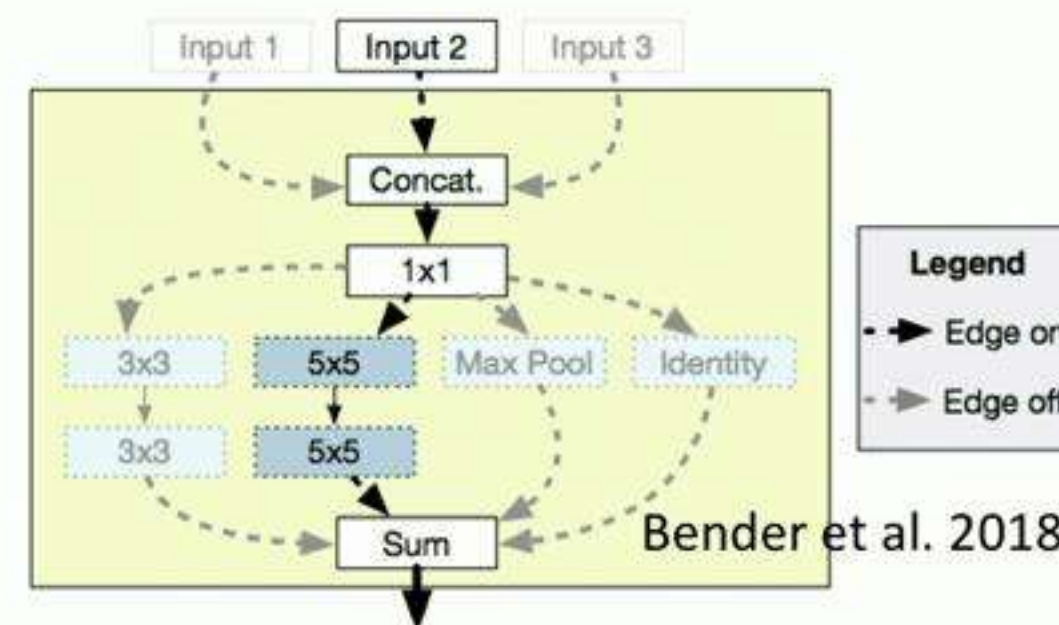


39% training time reduction of ResNet50 for ImageNet.

[1] Lym, S, E Choukse, S Zangeneh, **W Wen**, M Erez, and S Shanghavi. "PruneTrain: Gradual Structured Pruning from Scratch for Faster Neural Network Training." *International Conference for High Performance Computing, Networking, Storage and Analysis (SC)* (2019).

Best Student Paper Finalist.

AutoML with weight sharing



AutoML by SSL?

SSL can prune unimportant branches.

Scaling Natural Language Processing

#L: layer, #H: hidden size, #A: attention head

Devlin et al. 2018

Hyperparams				Dev Set Accuracy		
#L	#H	#A	LM (ppl)	MNLI-m	MRPC	SST-2
3	768	12	5.84	77.9	79.8	88.4
6	768	3	5.24	80.6	82.2	90.7
6	768	12	4.68	81.9	84.8	91.3
12	768	12	3.99	84.4	86.7	92.9
12	1024	16	3.54	85.7	86.9	93.3
24	1024	16	3.23	86.6	87.8	93.7

Scaling Natural Language Processing

#L: layer, #H: hidden size, #A: attention head

Devlin et al. 2018

Hyperparams			Dev Set Accuracy			
#L	#H	#A	LM (ppl)	MNLI-m	MRPC	SST-2
3	768	12	5.84	77.9	79.8	88.4
6	768	3	5.24	80.6	82.2	90.7
6	768	12	4.68	81.9	84.8	91.3
12	768	12	3.99	84.4	86.7	92.9
12	1024	16	3.54	85.7	86.9	93.3
24	1024	16	3.23	86.6	87.8	93.7

- Scaling NLP models always gives us performance gain
- VggNet-era NLP
 - Hard to further scale up because of computation/memory cost
 - We need more scalable training methods and models

Scaling Natural Language Processing

#L: layer, #H: hidden size, #A: attention head

Devlin et al. 2018

Hyperparams			Dev Set Accuracy			
#L	#H	#A	LM (ppl)	MNLI-m	MRPC	SST-2
3	768	12	5.84	77.9	79.8	88.4
6	768	3	5.24	80.6	82.2	90.7
6	768	12	4.68	81.9	84.8	91.3
12	768	12	3.99	84.4	86.7	92.9
12	1024	16	3.54	85.7	86.9	93.3
24	1024	16	3.23	86.6	87.8	93.7

- Scaling NLP models always gives us performance gain
- VggNet-era NLP
 - Hard to further scale up because of computation/memory cost
 - We need more scalable training methods and models

Direction 1: Design more scalable training methods

- Give an architecture, make its training faster
 - TernGrad, SSL, more?

Direction 2: Design more scalable new models

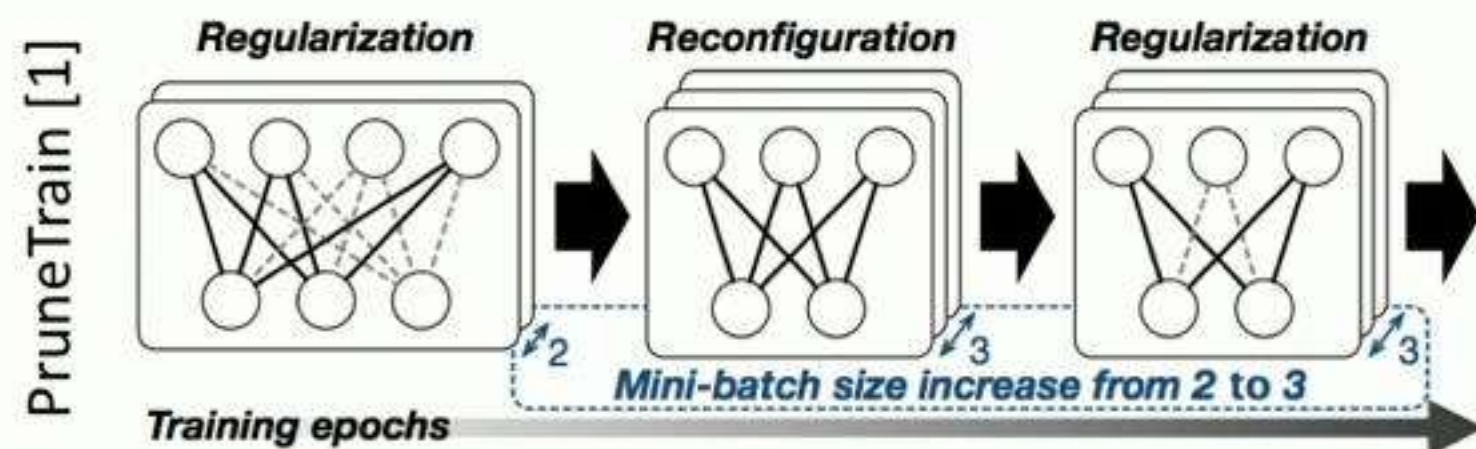
- Design a fast, accurate and compact base model and scale up
- Design scalable new models by AutoML

SSL Implications and Beyond

Lottery Ticket Hypothesis (Frankle et al 2019)

“dense, randomly-initialized, feed-forward networks contain subnetworks (winning tickets) that—when trained in isolation— reach test accuracy comparable to the original network in a similar number of iterations.”

SSL identifies winning tickets in the early stage?

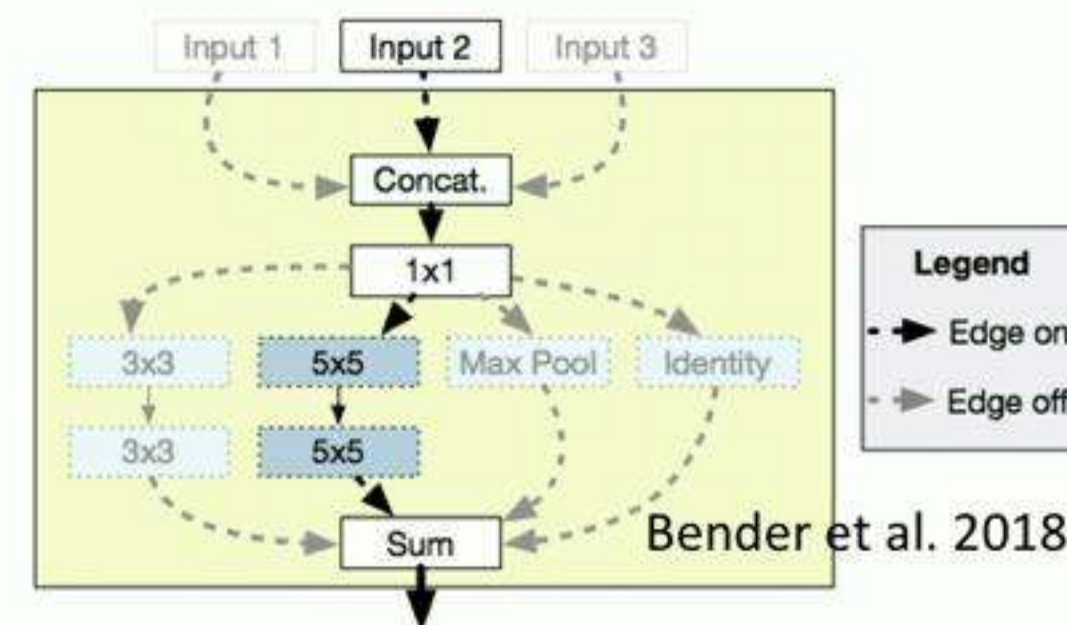


39% training time reduction of ResNet50 for ImageNet.

[1] Lym, S, E Choukse, S Zangeneh, **W Wen**, M Erez, and S Shanghavi. "PruneTrain: Gradual Structured Pruning from Scratch for Faster Neural Network Training." *International Conference for High Performance Computing, Networking, Storage and Analysis (SC)* (2019).

Best Student Paper Finalist.

AutoML with weight sharing



Bender et al. 2018

AutoML by SSL?

SSL can prune unimportant branches.

Scaling Natural Language Processing

#L: layer, #H: hidden size, #A: attention head

Devlin et al. 2018

Hyperparams			Dev Set Accuracy			
#L	#H	#A	LM (ppl)	MNLI-m	MRPC	SST-2
3	768	12	5.84	77.9	79.8	88.4
6	768	3	5.24	80.6	82.2	90.7
6	768	12	4.68	81.9	84.8	91.3
12	768	12	3.99	84.4	86.7	92.9
12	1024	16	3.54	85.7	86.9	93.3
24	1024	16	3.23	86.6	87.8	93.7

- Scaling NLP models always gives us performance gain
- VggNet-era NLP
 - Hard to further scale up because of computation/memory cost
 - We need more scalable training methods and models

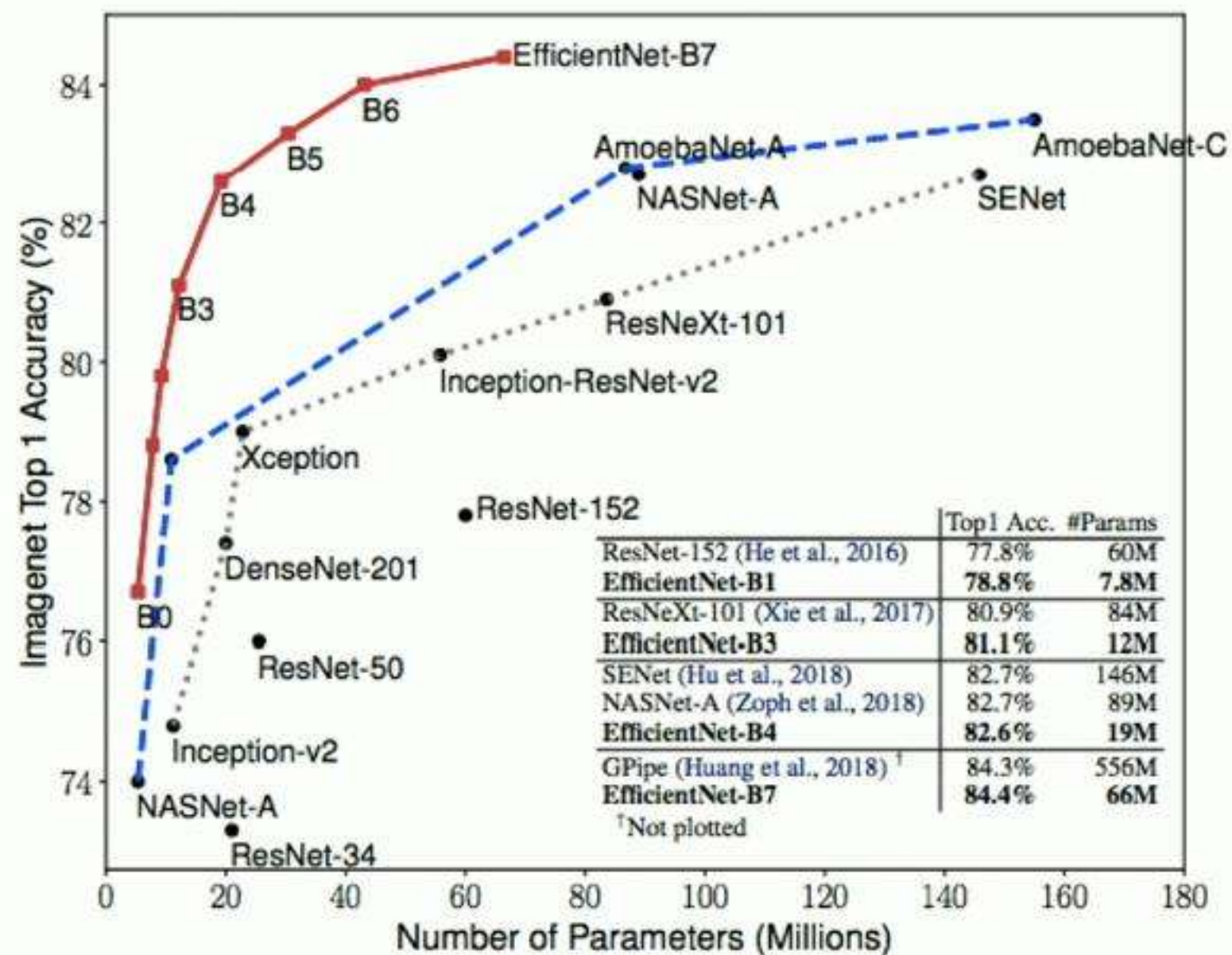
Direction 1: Design more scalable training methods

- Give an architecture, make its training faster
 - TernGrad, SSL, more?

Direction 2: Design more scalable new models

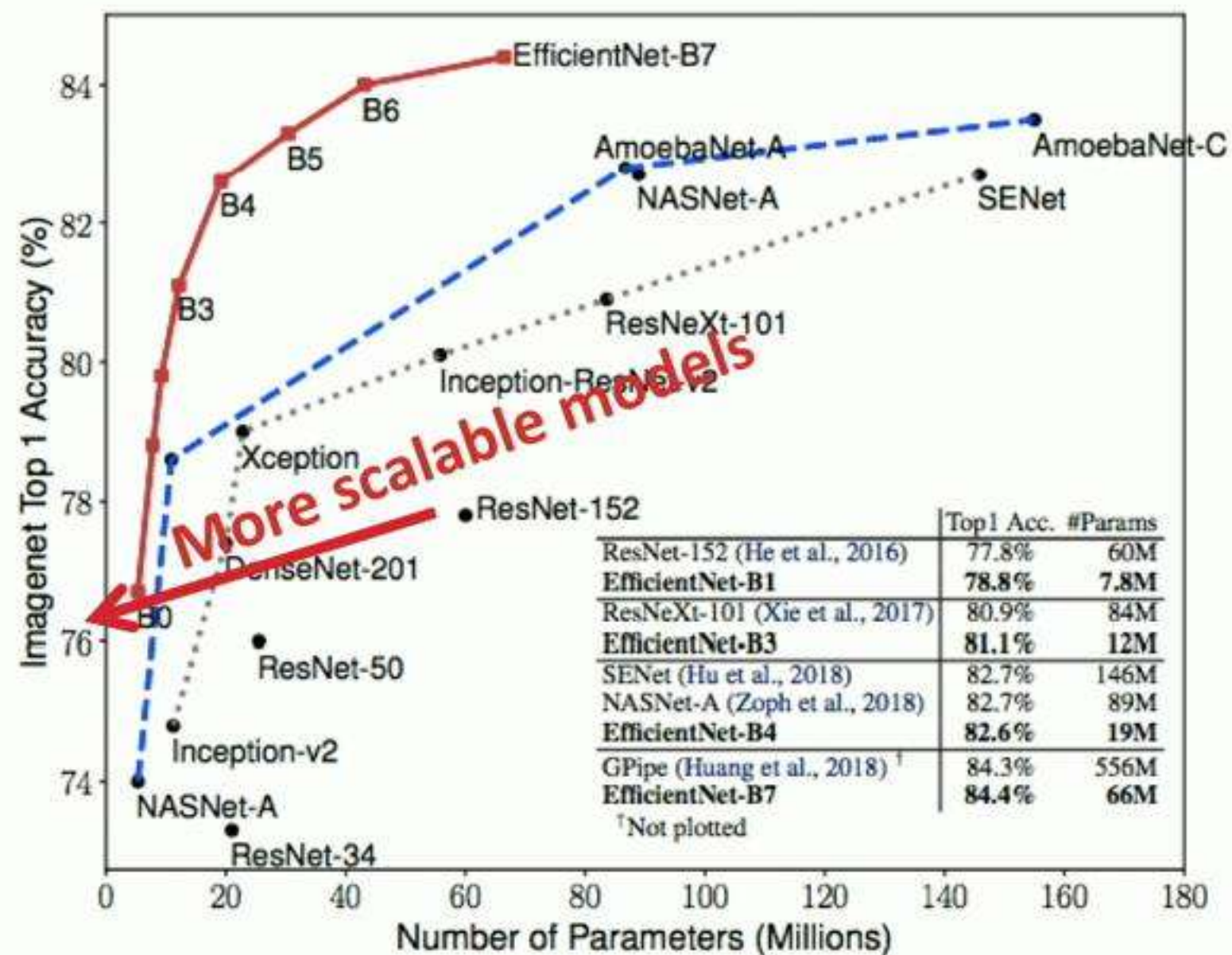
- Design a fast, accurate and compact base model and scale up
- Design scalable new models by AutoML

Automated Machine Learning (AutoML)



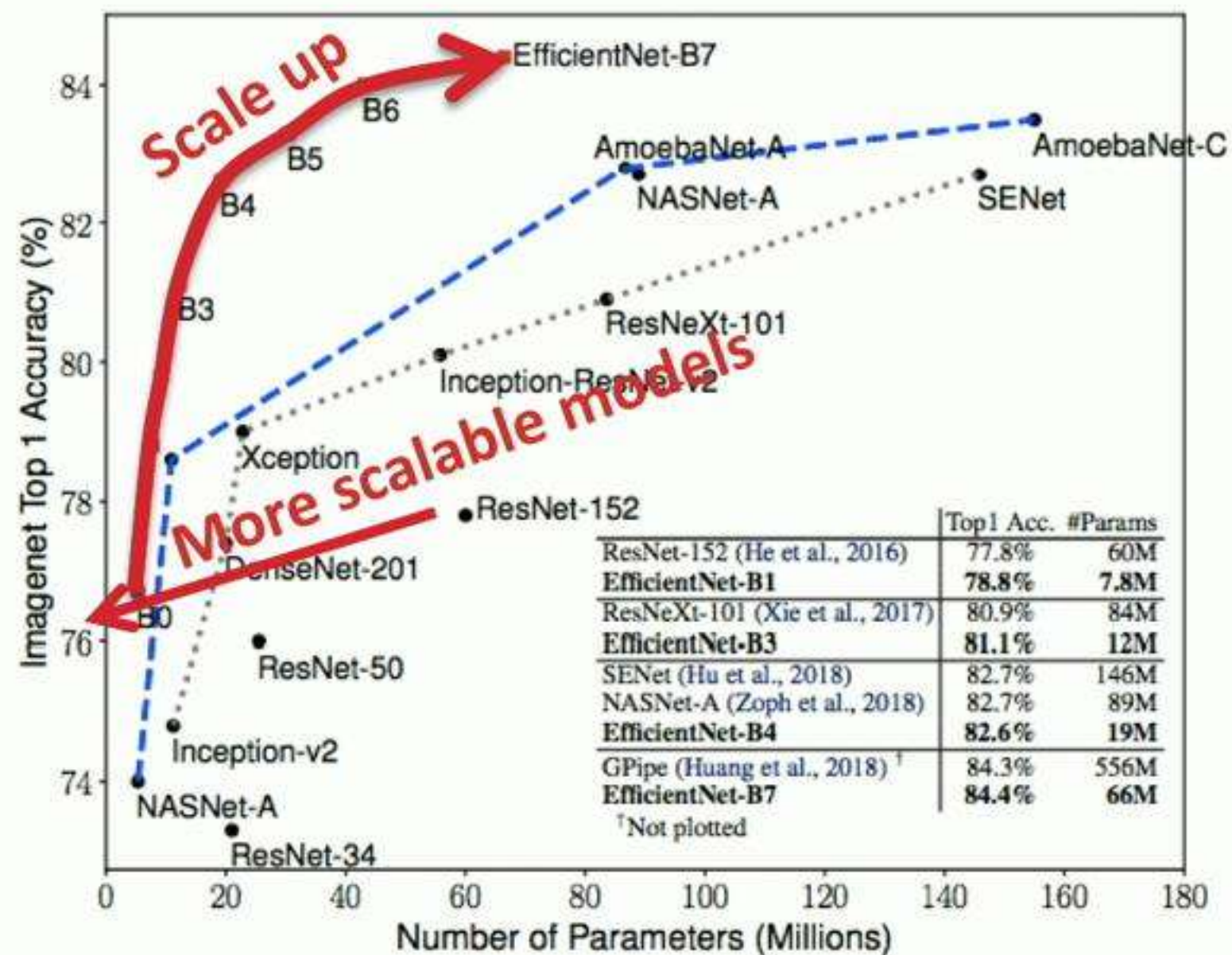
Tan & Le, 2019

Automated Machine Learning (AutoML)



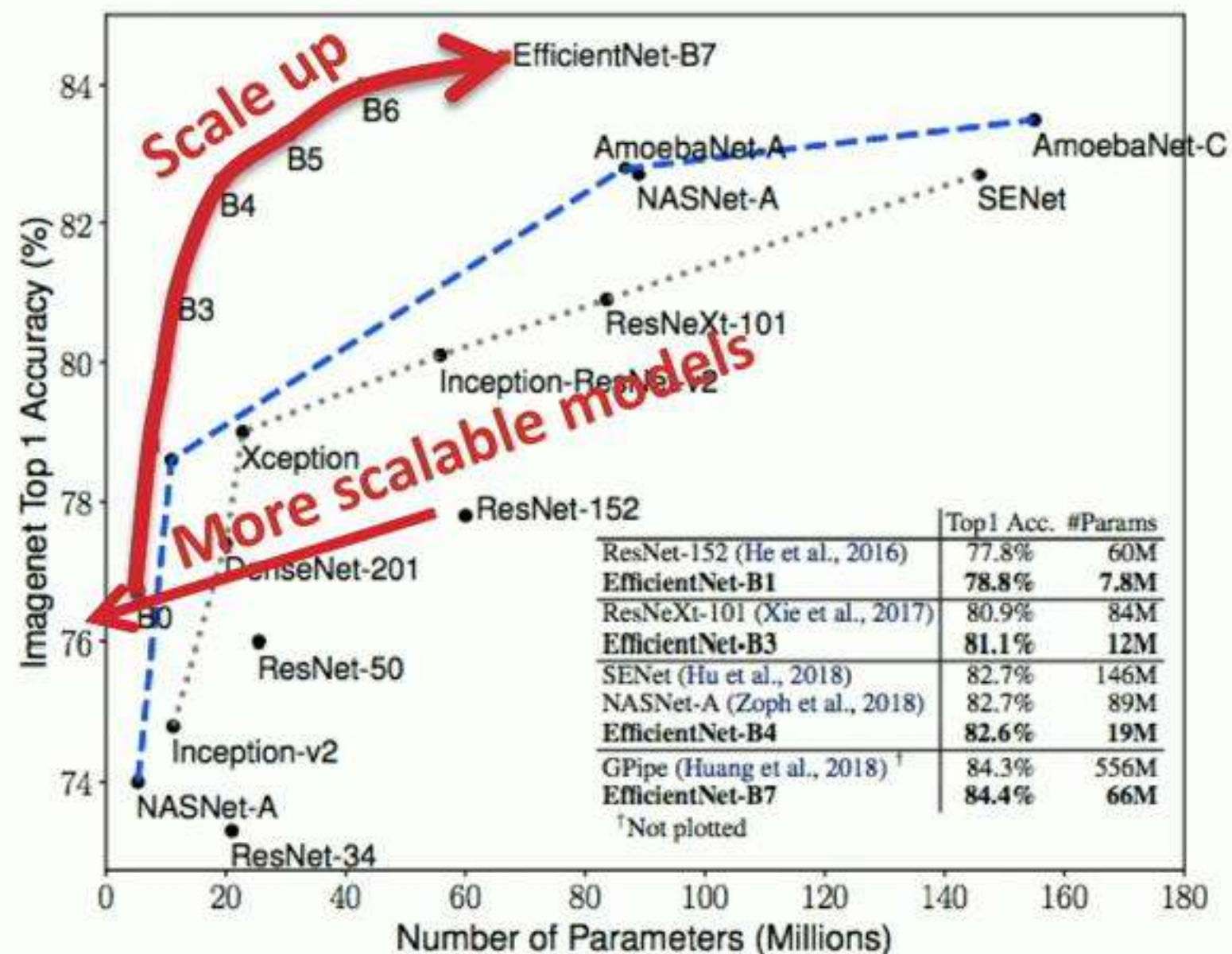
Tan & Le, 2019

Automated Machine Learning (AutoML)



Tan & Le, 2019

Automated Machine Learning (AutoML)

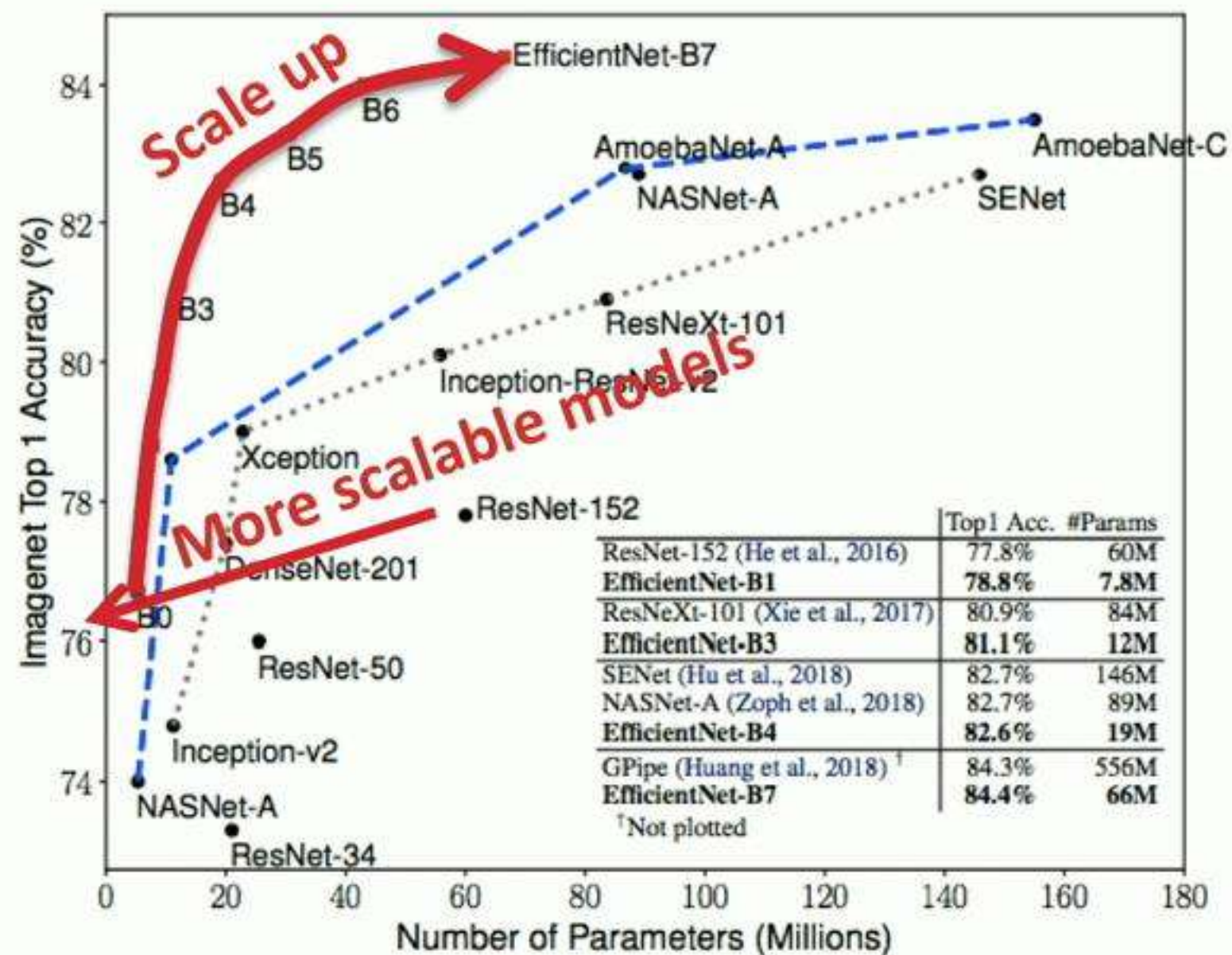


Problems in AutoML:

1. computation/memory intensive
2. Rely on human designed operations and search combination only

Tan & Le, 2019

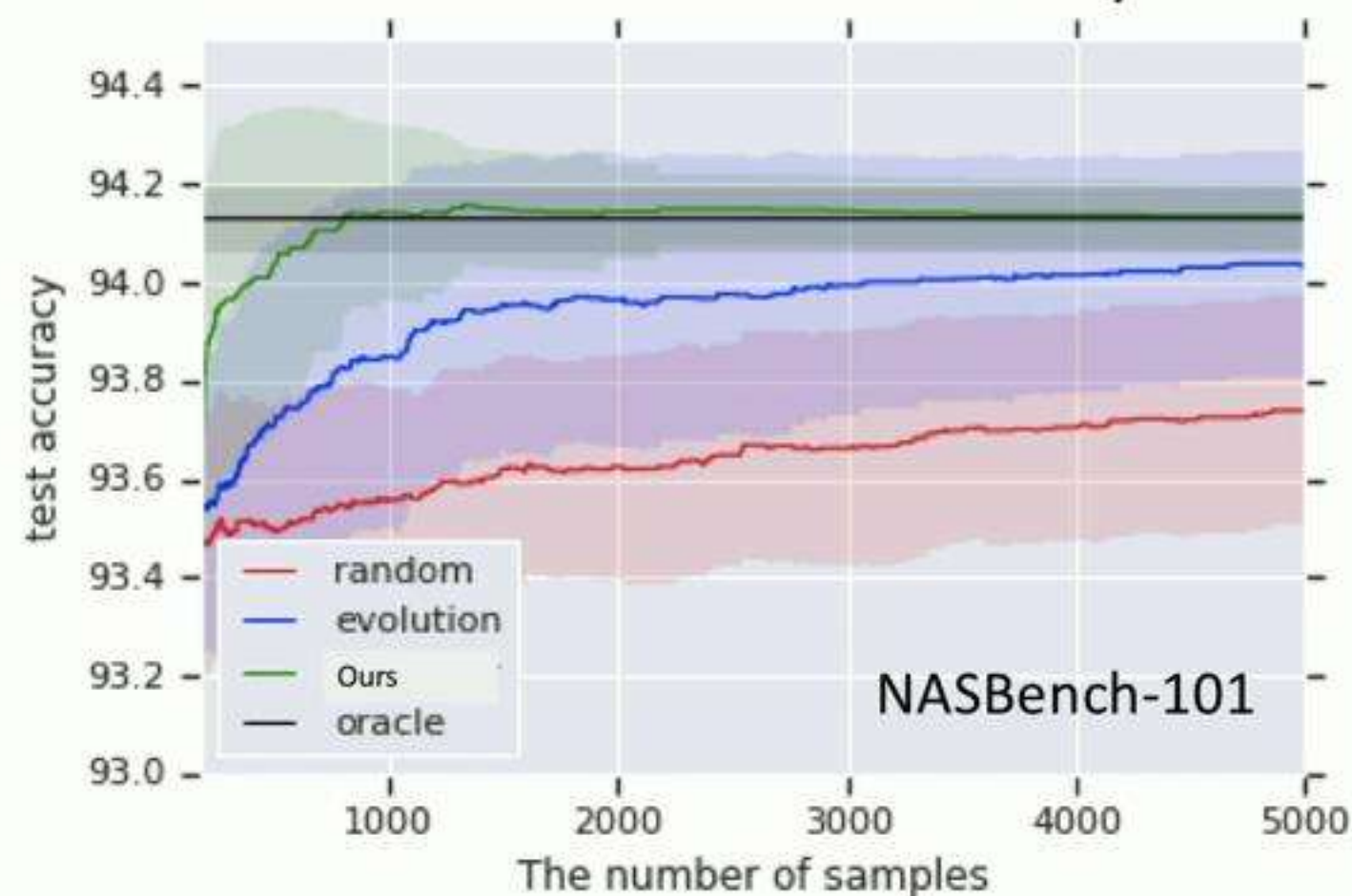
Automated Machine Learning (AutoML)



Tan & Le, 2019

Problems in AutoML:

1. computation/memory intensive
2. Rely on human designed operations and search combination only



References

- Pranav Dar, 25 Open Datasets for Deep Learning Every Data Scientist Must Work With, <https://www.analyticsvidhya.com/blog/2018/03/comprehensive-collection-deep-learning-datasets/>. 2018
- Canziani, Alfredo, Adam Paszke, and Eugenio Culurciello. "An analysis of deep neural network models for practical applications." *arXiv preprint arXiv:1605.07678* (2016).
- Silver, David, Aja Huang, Chris J. Maddison, Arthur Guez, Laurent Sifre, George Van Den Driessche, Julian Schrittwieser et al. "Mastering the game of Go with deep neural networks and tree search." *nature* 529, no. 7587 (2016): 484.
- Huang, Gao, Zhuang Liu, Laurens Van Der Maaten, and Kilian Q. Weinberger. "Densely Connected Convolutional Networks." In *CVPR*, vol. 1, no. 2, p. 3. 2017.
- Dean, Jeffrey, Greg Corrado, Rajat Monga, Kai Chen, Matthieu Devin, Mark Mao, Andrew Senior et al. "Large scale distributed deep networks." In *Advances in neural information processing systems*, pp. 1223-1231. 2012.
- Goyal, Priya, Piotr Dollár, Ross Girshick, Pieter Noordhuis, Lukasz Wesolowski, Aapo Kyrola, Andrew Tulloch, Yangqing Jia, and Kaiming He. "Accurate, large minibatch SGD: training imagenet in 1 hour." *arXiv preprint arXiv:1706.02677* (2017).



THANKS! Q&A?

Wei Wen, Duke University
@wei_wen_
www.pittnuts.com