

Knowledge Distillation with Adversarial Samples Supporting Decision Boundary

Byeongho Heo,^{1*} Minsik Lee,^{2*} Sangdoo Yun,³ Jin Young Choi¹
{bhheo, jychoi}@snu.ac.kr, mleepaper@hanyang.ac.kr, sangdoo.yun@navercorp.com

¹Department of ECE, ASRI, Seoul National University, Korea

²Division of EE, Hanyang University, Korea

³Clova AI Research, NAVER Corp, Korea

AAAI 2019

Contents

- Knowledge Distillation
- Adversarial Samples
- Method
- Experiments

Knowledge Distillation

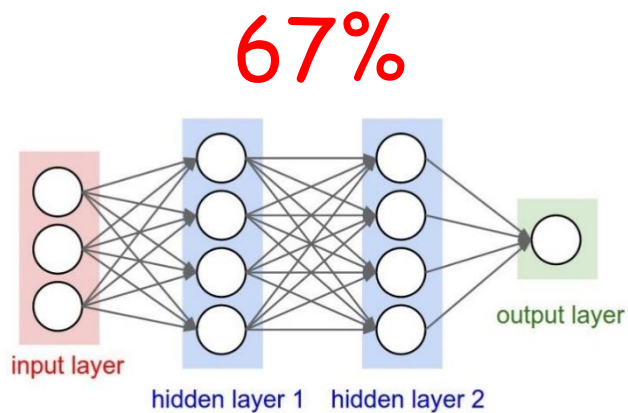
- In order to get better performance, we can use
 - More complex neural network (over-parameterization)
 - Ensemble learning

huge computation cost!

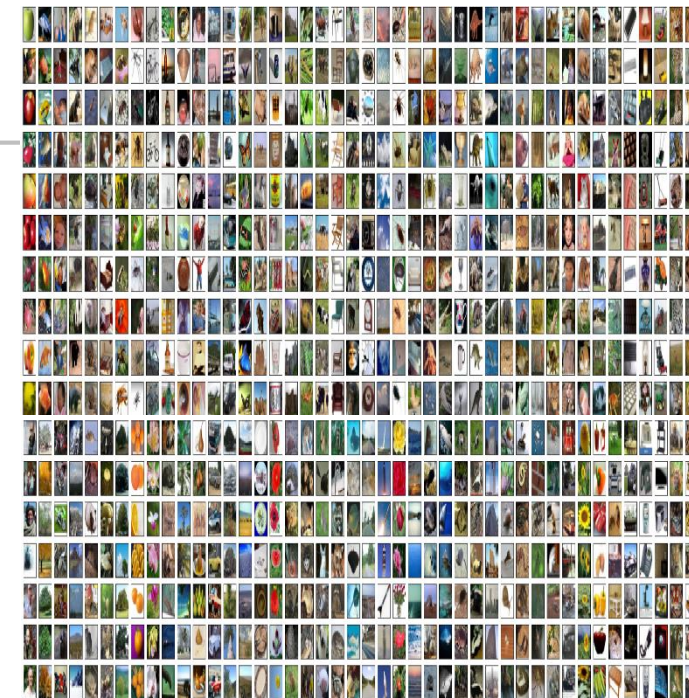
- In real-world tasks, we want
 - Small model with low complexity
 - Comparable performance

Knowledge Distillation

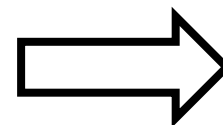
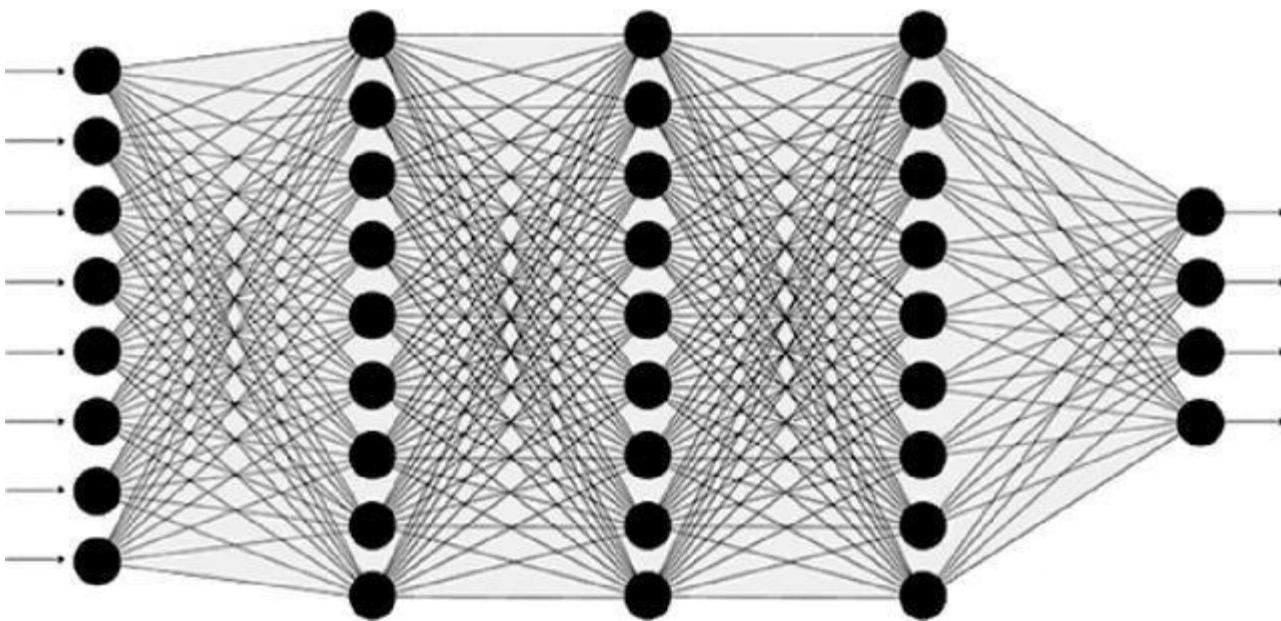
✗



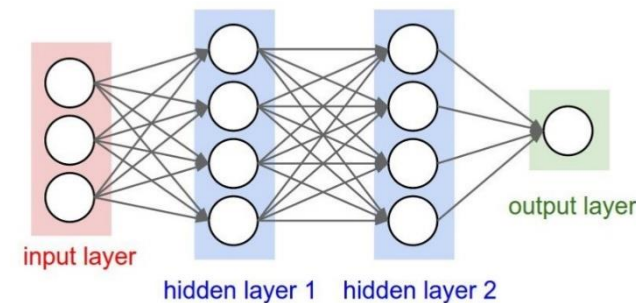
← train



98%



95%



✓

Knowledge Distillation

□ How to distillation? [1]

- A transfer set
- Minimize the cross-entropy of the distribution produced by two models

$$\min - p * \log(q)$$

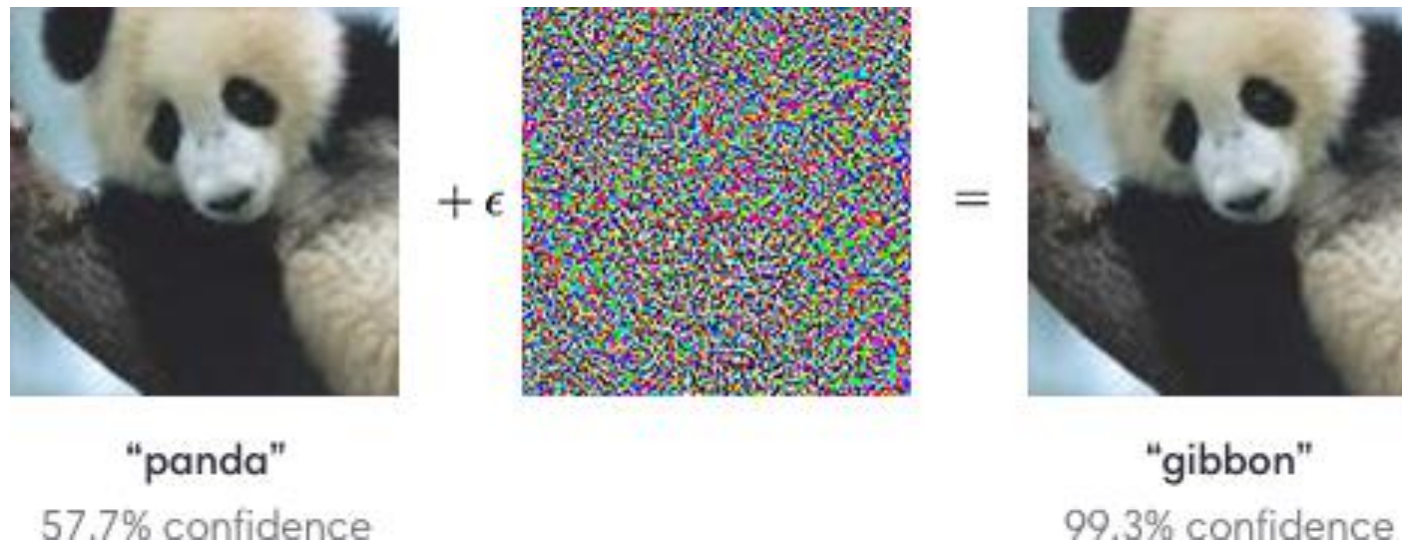
□ Why is it called distillation?

$$q_i = \frac{\exp(z_i/T)}{\sum_j \exp(z_j/T)}$$

*$T \rightarrow 0, q_i \rightarrow \text{one hot}$
 $T \rightarrow \infty, q_i \rightarrow \text{softer}$*

- T is a temperature that is normally set to 1. Using a higher value for T produces a softer probability distribution over classes.
- When the temperature T is increased, knowledge is transferred to the distilled model. After it has been trained, it uses a temperature of 1.

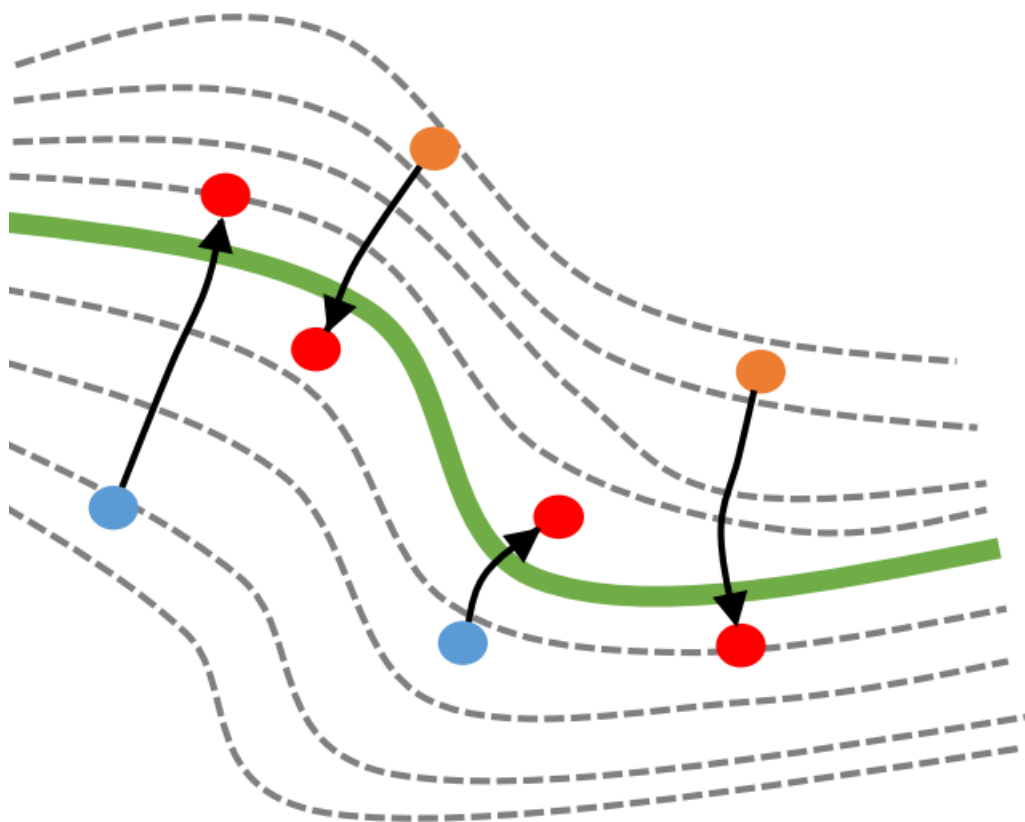
Adversarial Samples



$$x + \epsilon \operatorname{sign}(\nabla_x \mathcal{L}_{\text{CE}}(f(x), y)) = x_{\text{adv}}$$

$$\max_{\|x' - x\| \leq \epsilon} \mathbb{E}_{(x, y) \sim \mathcal{D}} [\mathcal{L}(f_{\theta}(x'), y)]$$

Adversarial Samples



Adversarial attack

$$x_{adv} = x + \epsilon \operatorname{sign}(\nabla_x \mathcal{L}_{CE}(f(x), y))$$

$$\max_{\|x' - x\| \leq \epsilon} \mathbb{E}_{(x, y) \sim \mathcal{D}} [\mathcal{L}(f_{\theta}(x'), y)]$$

Method

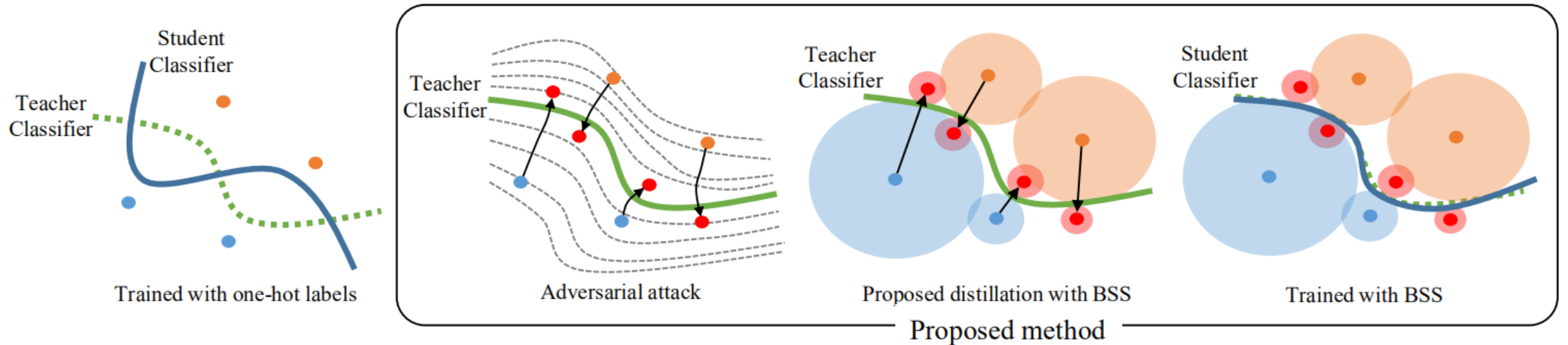


Figure 1: The concept of knowledge distillation using samples close to the decision boundary. The dots in the figure represent the training sample and the circle around a dot represents the distance to the nearest decision boundary. The samples close to the decision boundary enable more accurate knowledge transfer.

- The generalization performance of a classifier highly depends on how well the classifier learns the true decision boundary between the actual class distributions
- The key idea is about using adversarial samples near a decision boundary to transfer the knowledge related with the boundary.

Method

- ❑ Iterative Scheme to find a BBS.
 - BBS: boundary supporting samples.
- ❑ Knowledge distillation using BBS

Method

□ Iterative Scheme to find a BBS.

- The classifier f can produce classification scores for all classes, $f_b(x)$ and $f_k(x)$ denote the classification score for the base class and target class, respectively.
- The adversarial attack is to decrease $f_b(x)$ while increasing $f_k(x)$.
- The loss function for adversarial attack

$$\mathcal{L}_k(x) = f_b(x) - f_k(x)$$

This loss becomes **zero** at a point on the decision boundary, **positive** at a point within the base class, and **negative** at an adversarial point within the target class.

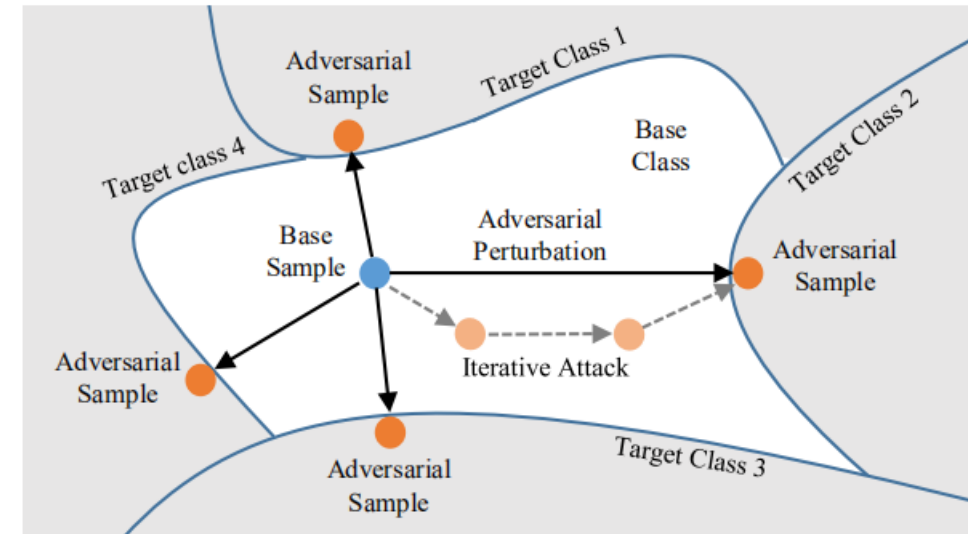


Figure 2: Iterative scheme to find BSSs for a base sample

Method

$$J(\mathbf{a}, \mathbf{b}) = -\mathbf{a}^T \log \mathbf{b}$$

□ Knowledge distillation using BBS

$$\mathcal{L}(n) = \mathcal{L}_{cls}(n) + \alpha \mathcal{L}_{KD}(n) + \beta \sum_k^K p_n^k \mathcal{L}_{BS}(n, k).$$

$$\mathcal{L}_{cls}(n) = J(\mathbf{y}_n^{true}, \sigma(f_s(\mathbf{x}_n))),$$

$$\mathcal{L}_{KD}(n) = J\left(\sigma\left(\frac{f_t(\mathbf{x}_n)}{T}\right), \sigma\left(\frac{f_s(\mathbf{x}_n)}{T}\right)\right),$$

$$\mathcal{L}_{BS}(n, k) = J\left(\sigma\left(\frac{f_t(\hat{\mathbf{x}}_n^k)}{T}\right), \sigma\left(\frac{f_s(\hat{\mathbf{x}}_n^k)}{T}\right)\right).$$

Experiments

Table 1: Comparison on CIFAR-10 dataset

Student	Original	Hinton (2015)	FITNET (2015) +Hinton	AT (2016) +Hinton	FSP (2017) +Hinton	Proposed	FSP (2017) +Proposed
ResNet 8	86.02%	86.66%	86.73%	86.86%	87.07%	<u>87.32%</u>	87.52%
ResNet 14	89.11%	89.75%	89.72%	89.84%	89.92%	90.34%	<u>90.13%</u>
ResNet 20	90.16%	90.77%	90.46%	<u>90.81%</u>	90.27%	91.23%	90.19%

Table 2: Comparison on ImageNet 32×32

	Original	Hinton (2015)	FITNET (2015) +Hinton	AT (2016) +Hinton	FSP (2017) +Hinton	Proposed	FSP (2017) +Proposed
Top1 acc	31.94%	32.43%	32.60%	32.61%	32.66%	<u>32.69%</u>	32.72%
Top5 acc	56.21%	56.99%	57.02%	57.14%	57.14%	<u>57.17%</u>	57.27%

Table 3: Comparison on TinyImageNet

	Original	Hinton (2015)	FITNET (2015) +Hinton	AT (2016) +Hinton	FSP (2017) +Hinton	Proposed	FSP (2017) +Proposed
Top1 acc	50.68%	52.35%	<u>53.52%</u>	52.74%	53.43%	52.99%	53.86%
Top5 acc	76.14%	77.67%	<u>78.10%</u>	77.82%	78.15%	<u>78.38%</u>	78.49%

Experiments

Table 4: Comparison of knowledge distillation using various types of adversarial samples.

	Baseline	Random noise	L2 minimize	FGSM (2015)	DeepFool (2016)	Proposed
Accuracy	86.66%	86.73%	86.94%	<u>87.06%</u>	86.95%	87.32%

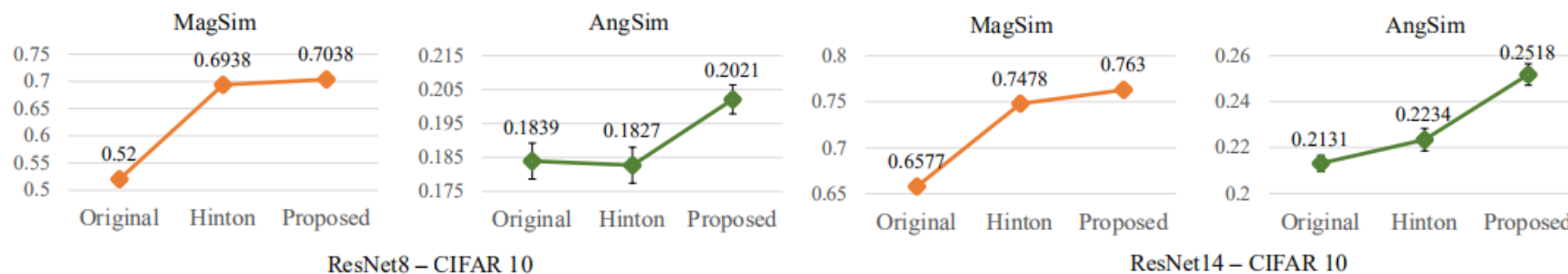


Figure 3: Evaluation of proposed method for decision boundary similarities (*MagSim*, *AngSim*).

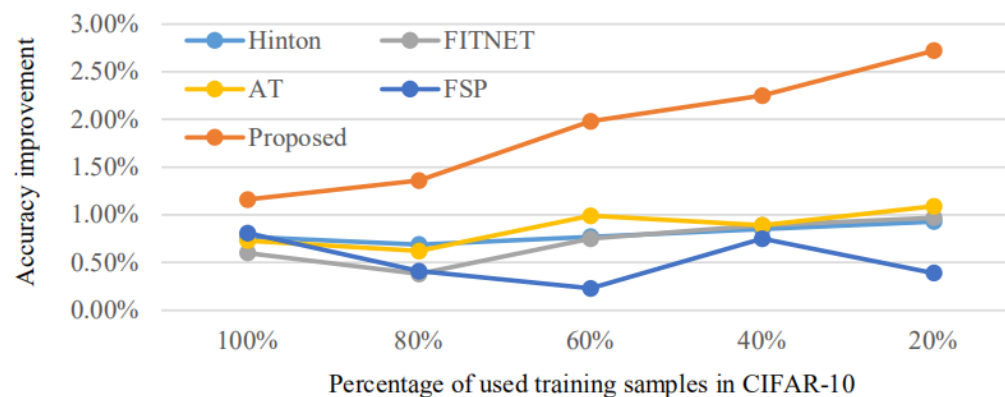


Figure 4: Generalization of the classifier. The smaller the number of training samples are, the larger improvement the proposed method shows.

Knowledge Transfer via Distillation of Activation Boundaries Formed by Hidden Neurons

Byeongho Heo,¹ Minsik Lee,² Sangdoo Yun,³ Jin Young Choi¹

{bhheo, jychoi}@snu.ac.kr, mleepaper@hanyang.ac.kr, sangdoo.yun@navercorp.com

¹Department of ECE, ASRI, Seoul National University, Korea

²Division of EE, Hanyang University, Korea

³Clova AI Research, NAVER Corp, Korea

AAAI 2019

Method

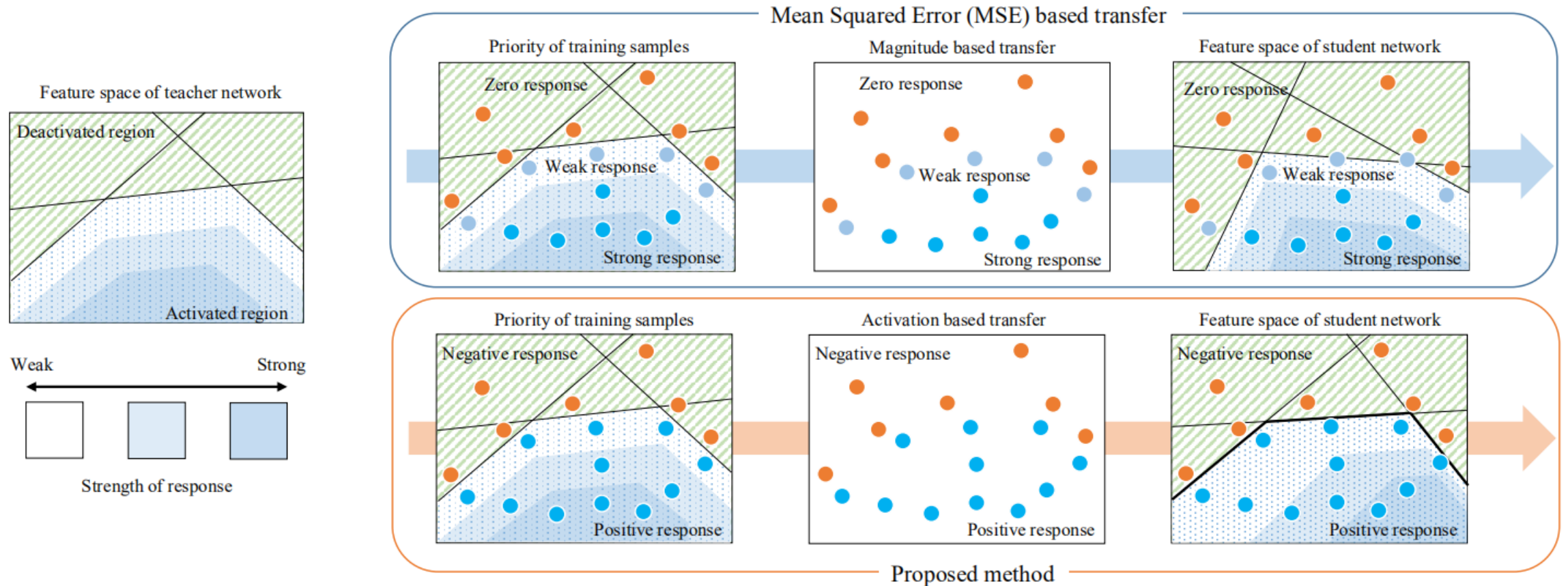


Figure 1: The concept of the proposed knowledge transfer method. The proposed method concentrates on the activation of neurons, not the magnitude of neuron responses. This concentration enables more precise transfer of the activation boundaries.

The second approach is to transfer the neuron responses for initializing the student network before the classification training. (network initialization method)

Method

- An existing method of neuron response transfer.

$$\mathcal{L}(\mathbf{I}) = \|\sigma(\mathcal{T}(\mathbf{I})) - \sigma(\mathcal{S}(\mathbf{I}))\|_2^2$$

$$\text{ReLU } \sigma(x) = \max(0, x)$$

- The part from the input image to a hidden layer is defined as function T , and a part of the student network is defined as function S .

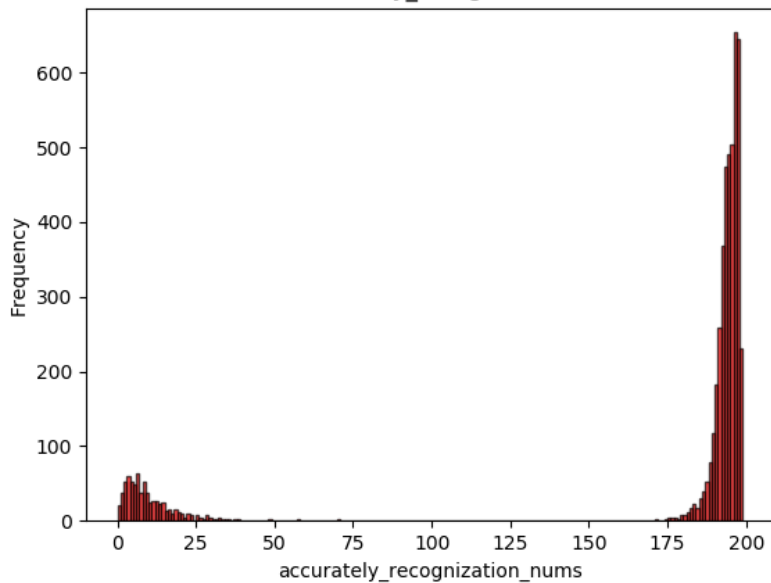
- The proposed method.

$$\mathcal{L}(\mathbf{I}) = \|\rho(\mathcal{T}(\mathbf{I})) - \rho(\mathcal{S}(\mathbf{I}))\|_1, \quad \rho(x) = \begin{cases} 1, & \text{if } x > 0 \\ 0, & \text{otherwise.} \end{cases}$$

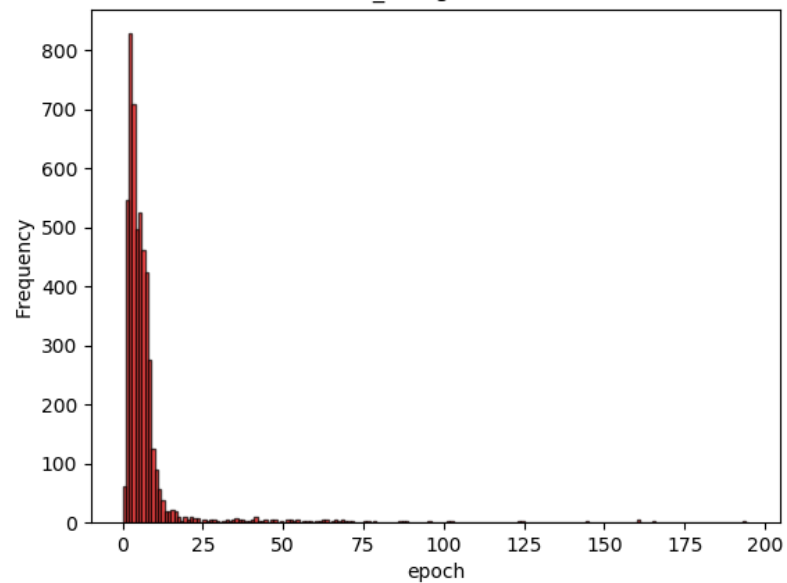
- The activation boundaries play an important role in forming the decision boundaries for classification-friendly partitioning of the feature space in each hidden layer.
- $\rho(x)$ expresses whether a neuron is activated or not.

Experiments

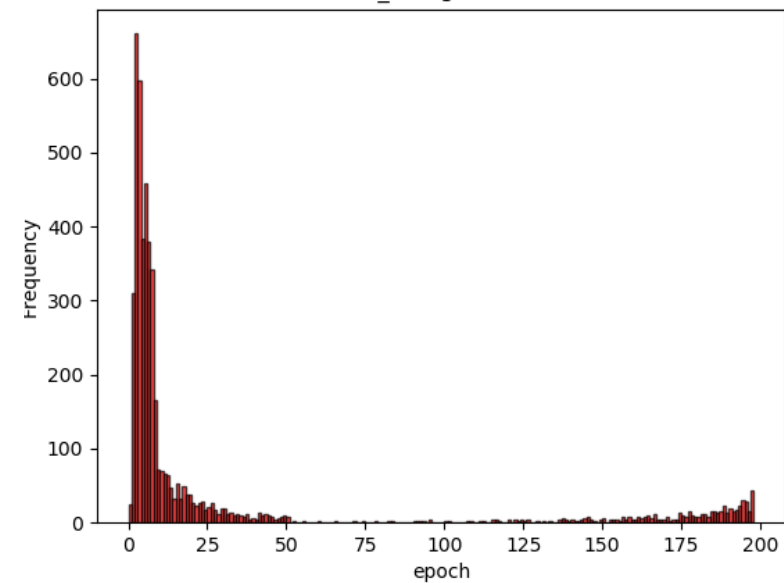
accurately_recognition



first_recognition



last_recognition



Thanks
