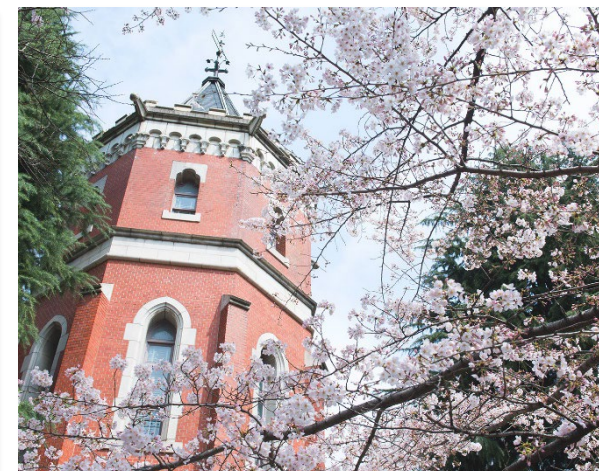
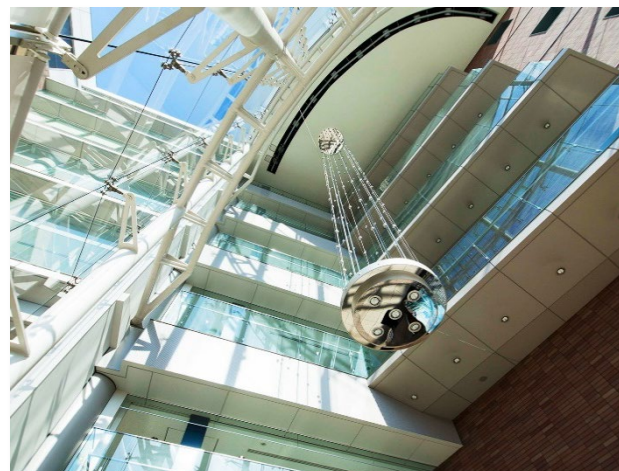




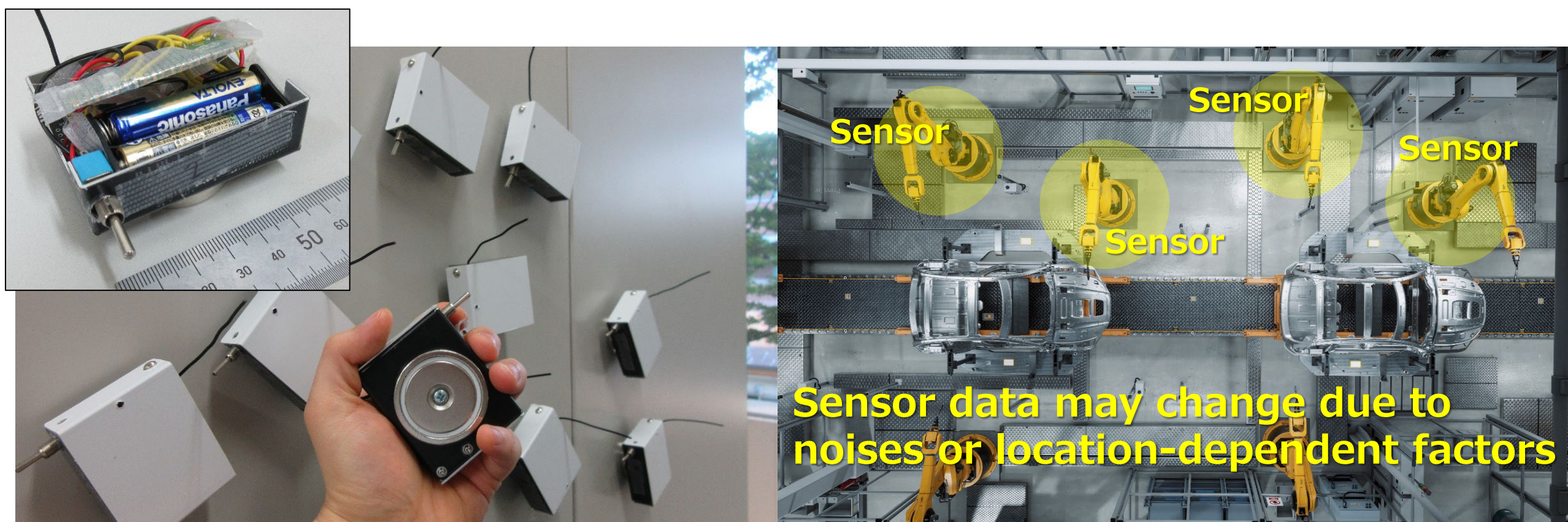
An Ultra-Fast On-Device CNN Finetuning Approach for Resource-Limited Edge Devices

Hiroki Matsutani (Keio University, Japan)



On-device learning (ODL) for IoT devices

- Motivation for neural network training at edge side
Addressing the **gap** between **pretrained model** and **deployed environment** by updating the model on-device [1,2]

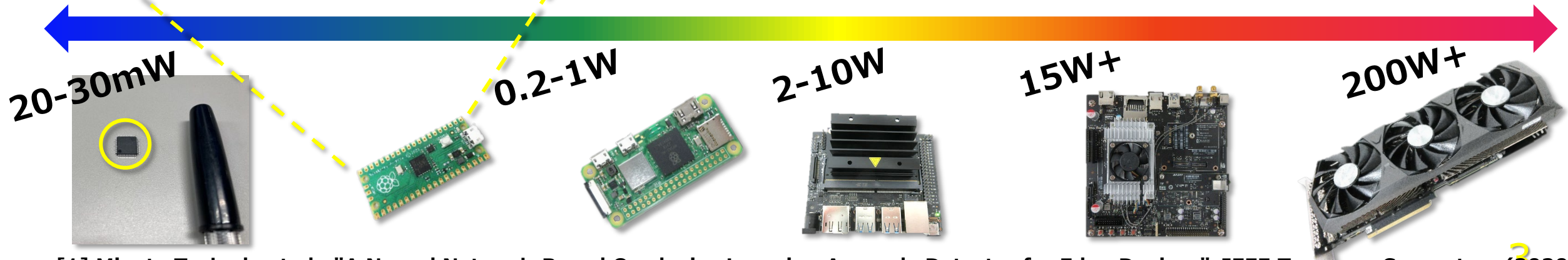
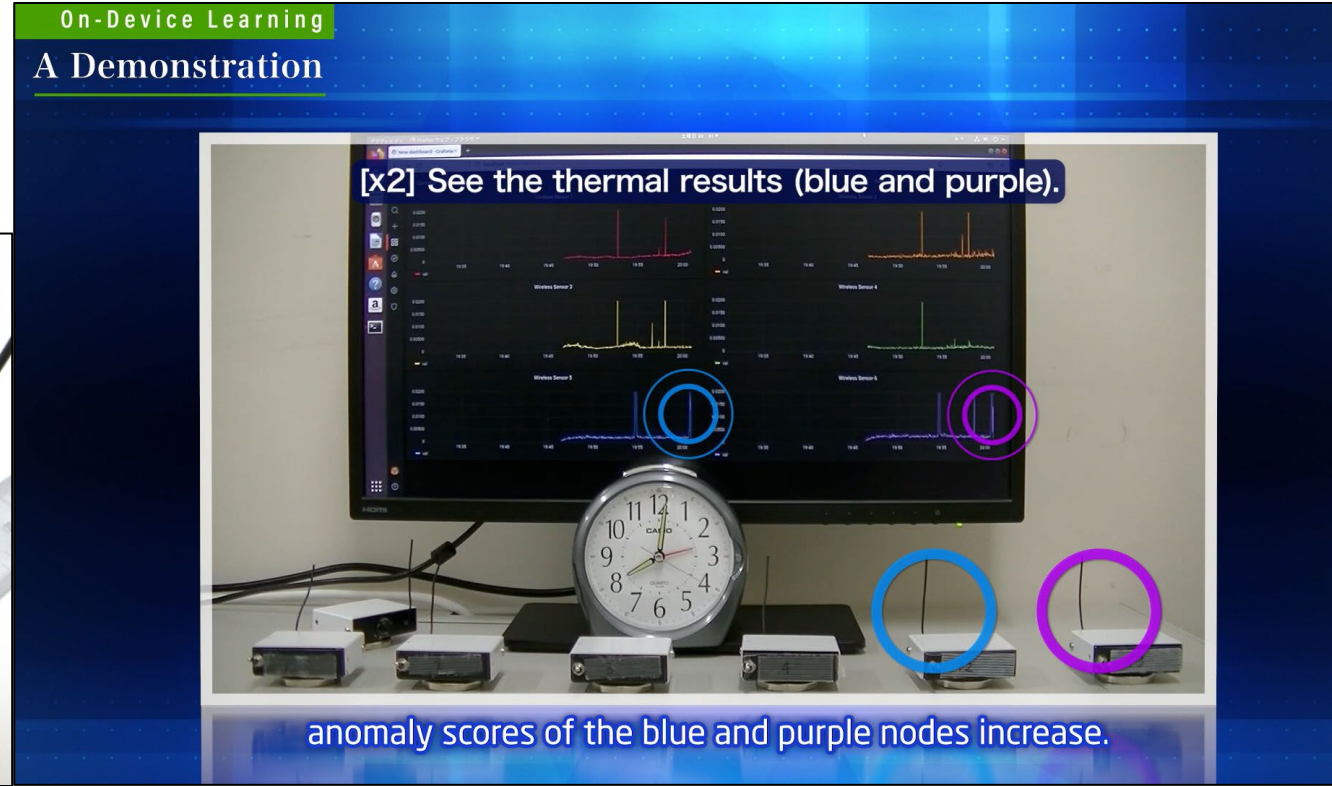
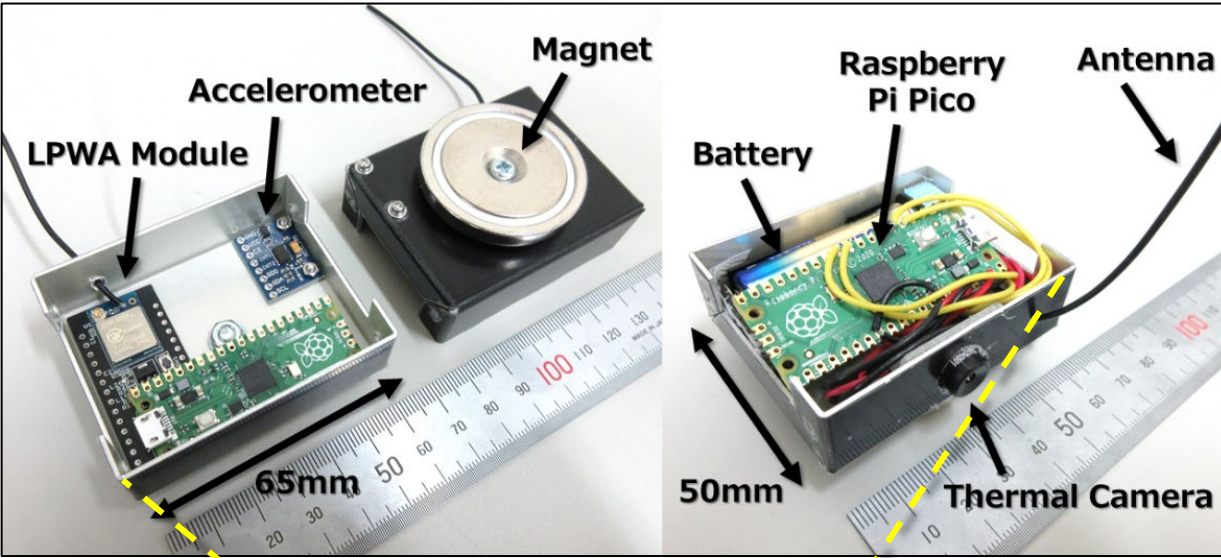


[1] Mineto Tsukada et al., "A Neural Network-Based On-device Learning Anomaly Detector for Edge Devices", IEEE Trans. on Computers (2020).

[2] Kazuki Sunaga et al., "Addressing Gap between Training Data and Deployed Environment by On-Device Learning", IEEE Micro (2023).²

On-device learning (ODL) for IoT devices

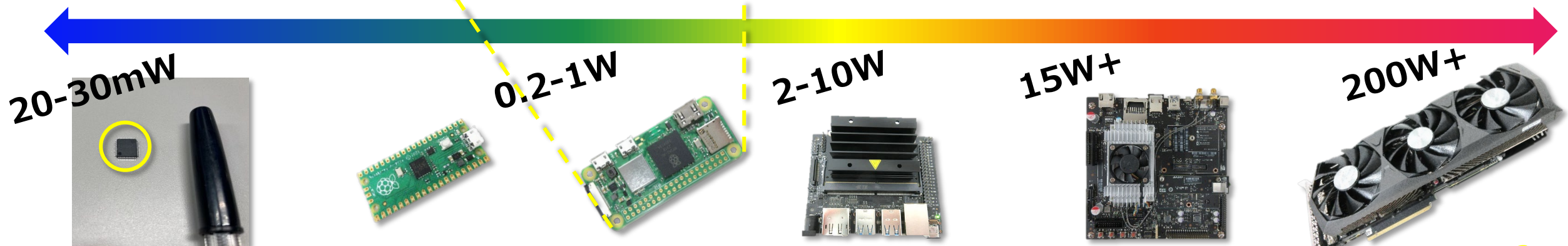
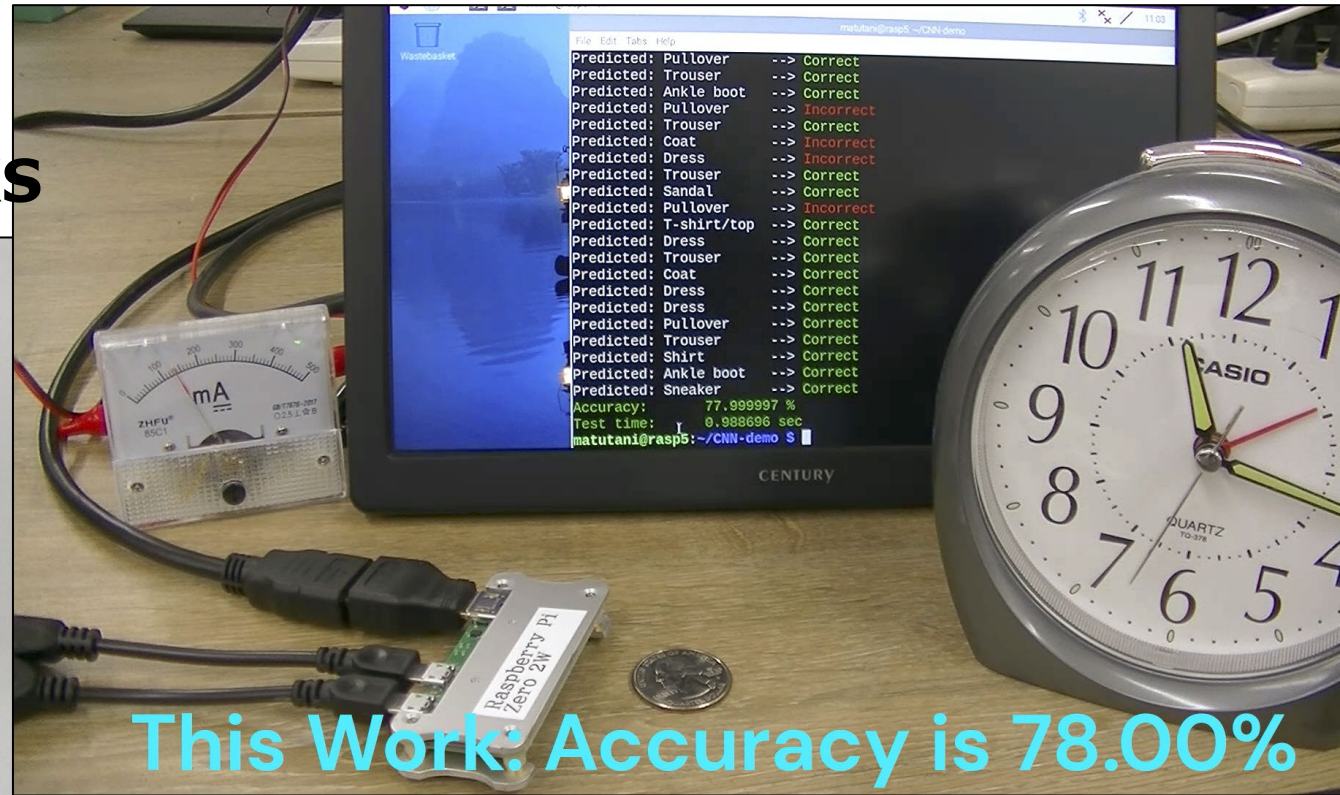
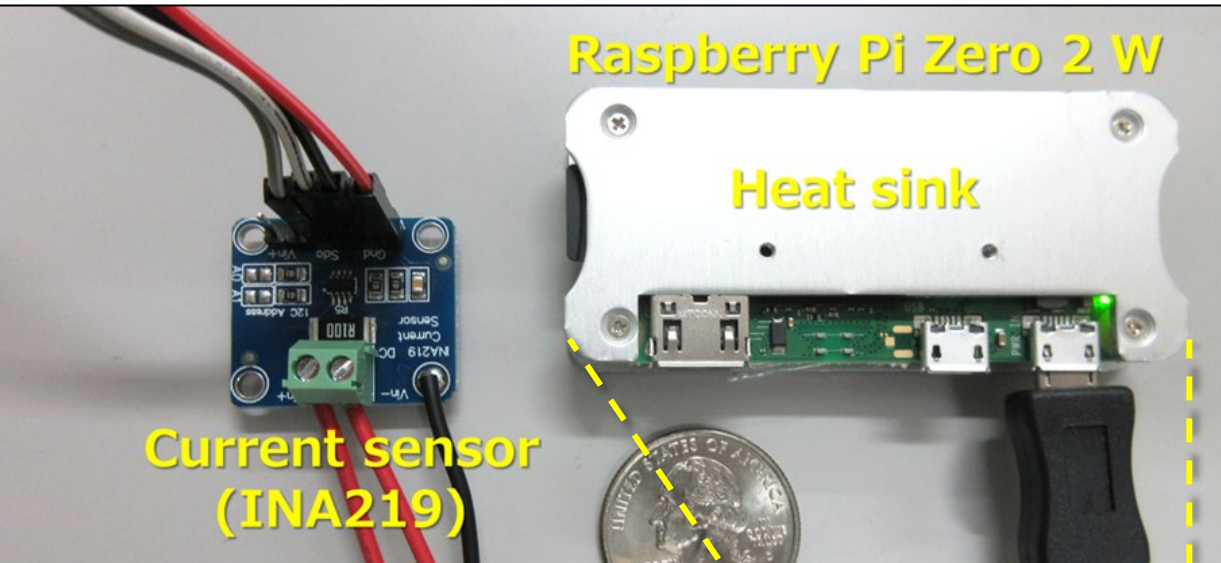
- Lightweight **sequential training** (since 2017) [1]



[1] Mineto Tsukada et al., "A Neural Network-Based On-device Learning Anomaly Detector for Edge Devices", IEEE Trans. on Computers (2020).

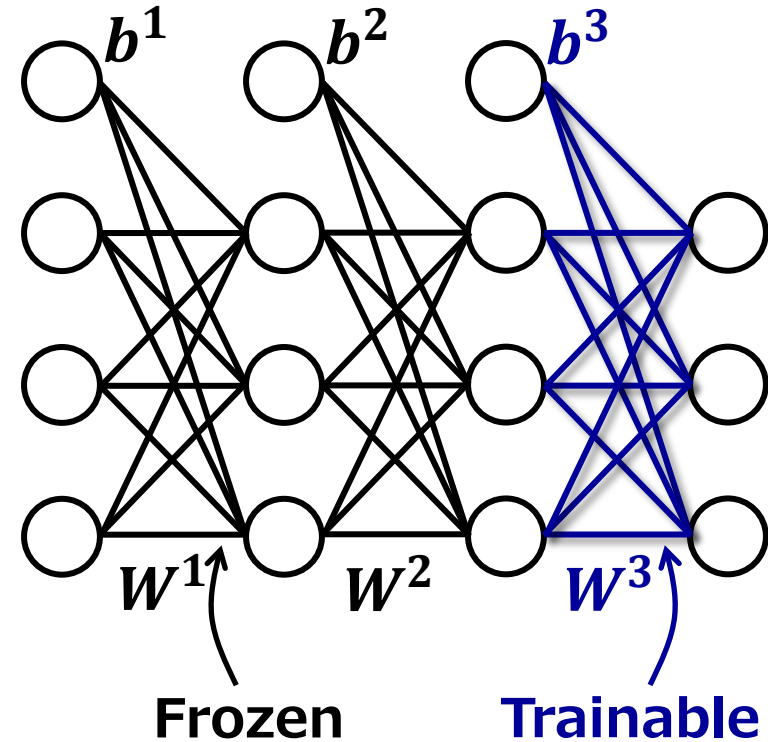
This talk: On-device finetuning for DNNs

- We focus on DNN/CNN including **vision** & **LLM** tasks



Baseline finetuning methods

FT-Last [1]

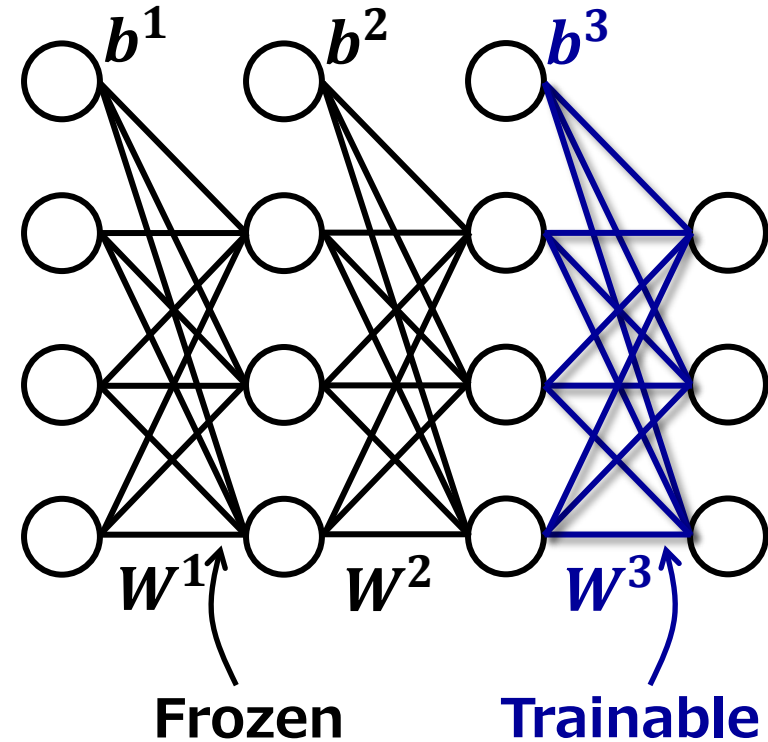


Weights of the **last layer**
(W^3 , b^3) are updated

[1] Haoyu Ren et al., "TinyOL: TinyML with Online-Learning on Microcontrollers", IJCNN'21.

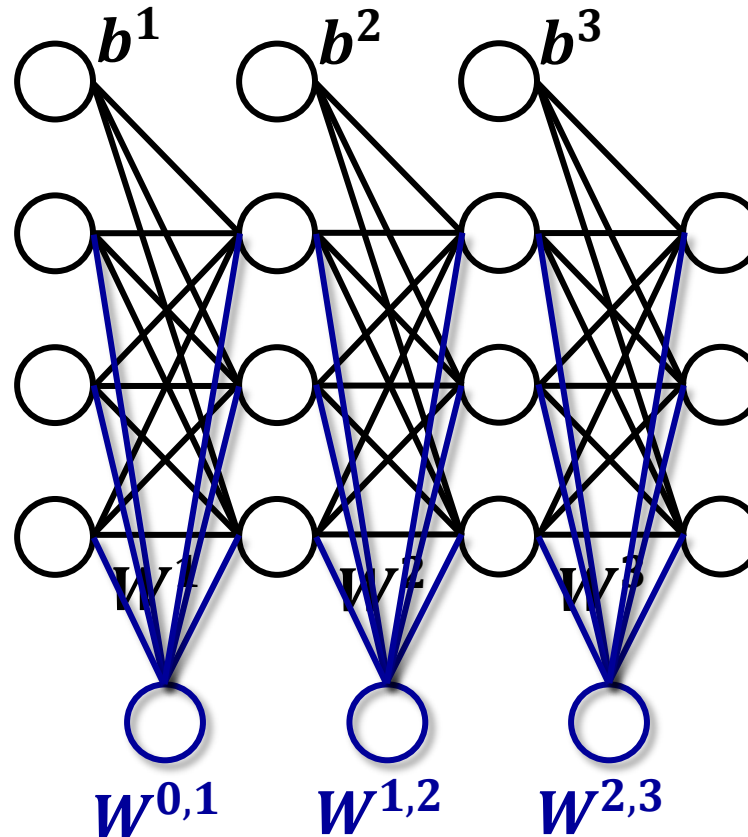
Baseline finetuning methods

FT-Last [1]



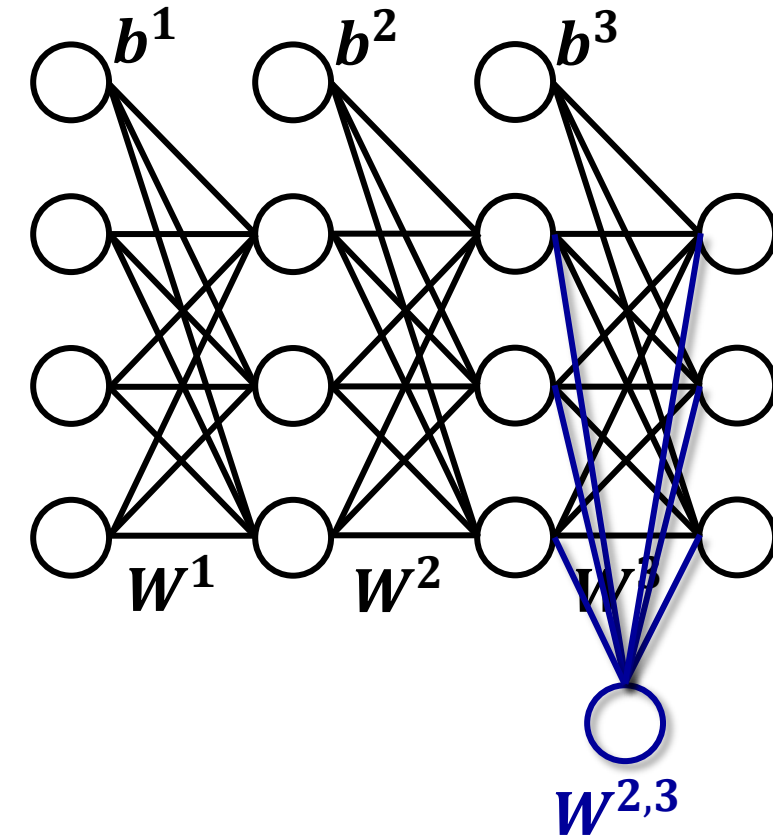
Weights of the **last layer** (W^3 , b^3) are updated

LoRA-All [2]



Trainable adapters are attached to **all layers**

LoRA-Last



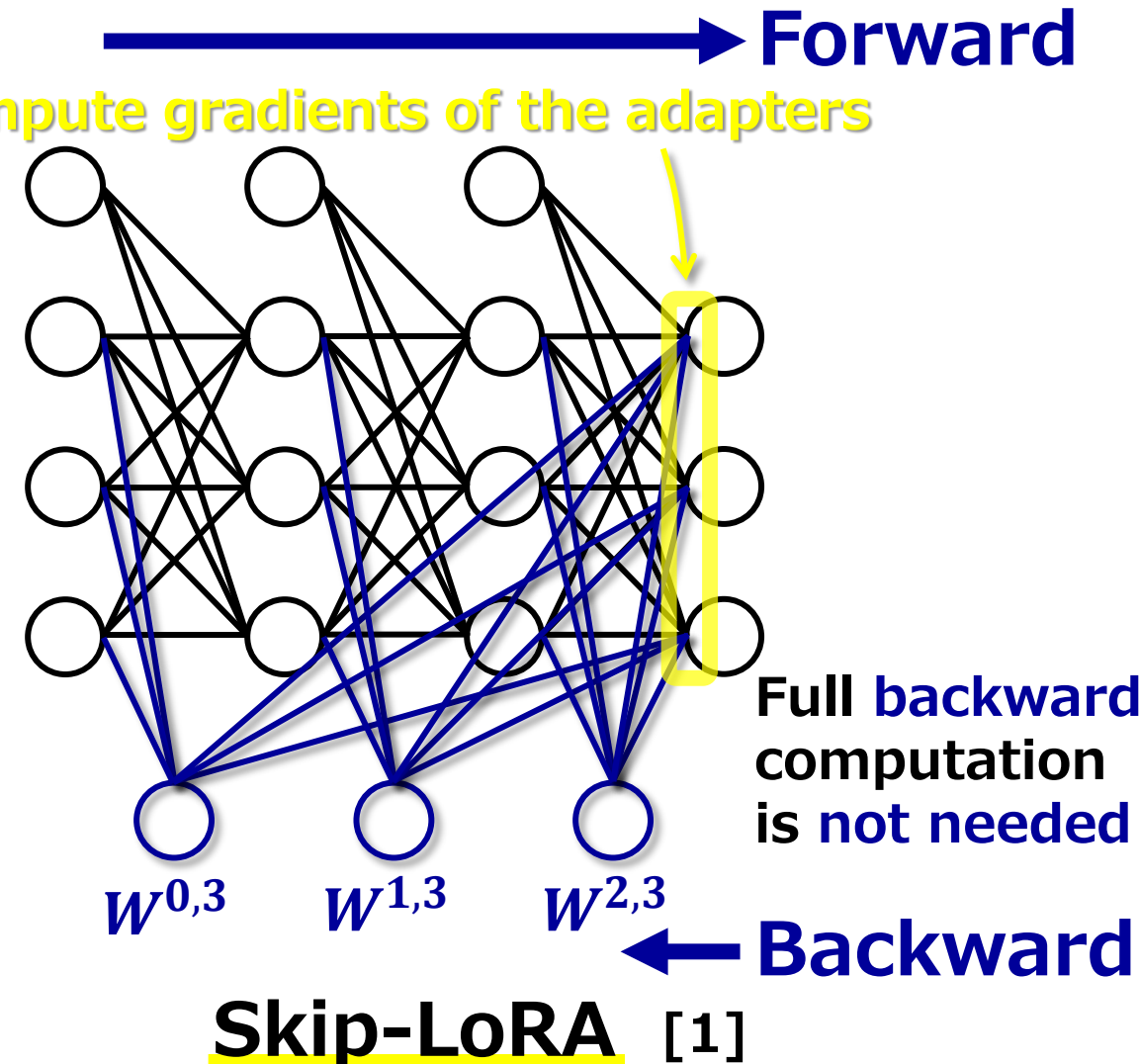
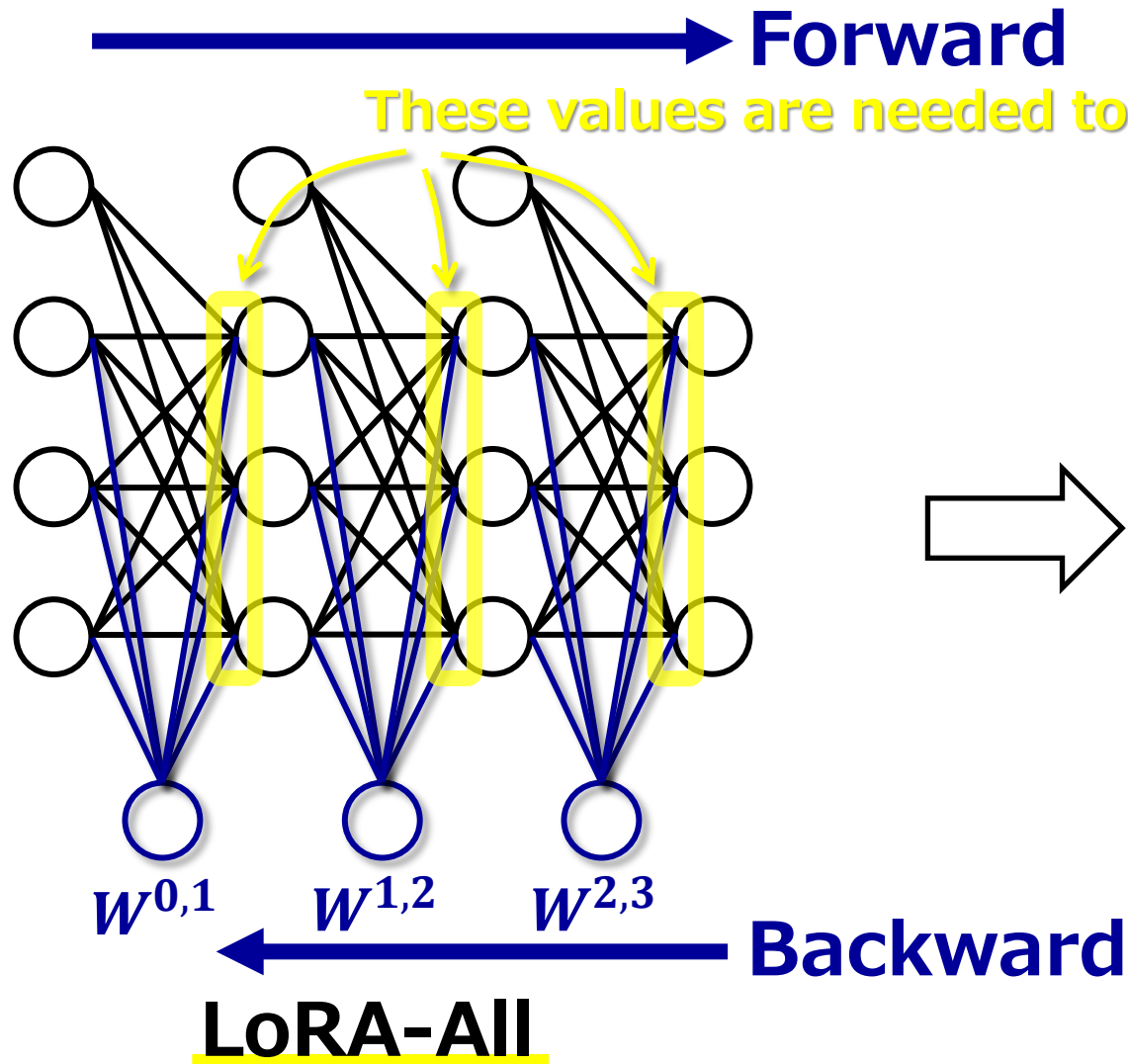
Trainable adapter is attached to the **last layer**

[1] Haoyu Ren et al., "TinyOL: TinyML with Online-Learning onMicrocontrollers", IJCNN'21.

[2] Edward J. Hu et al., "LoRA: Low-Rank Adaptation of Large LanguageModels", arXiv:2106.09685 (2021).

Our proposed approach: Skip-LoRA

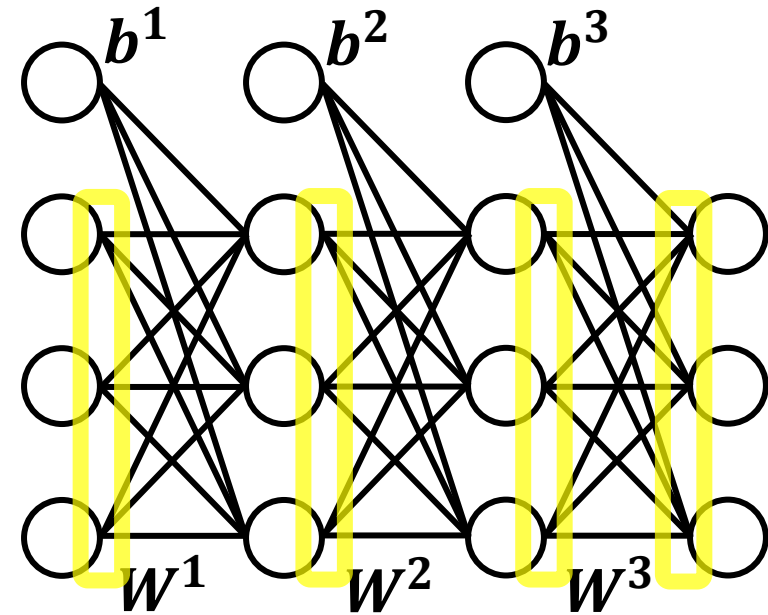
- Skip-LoRA can reduce the **backward** computation



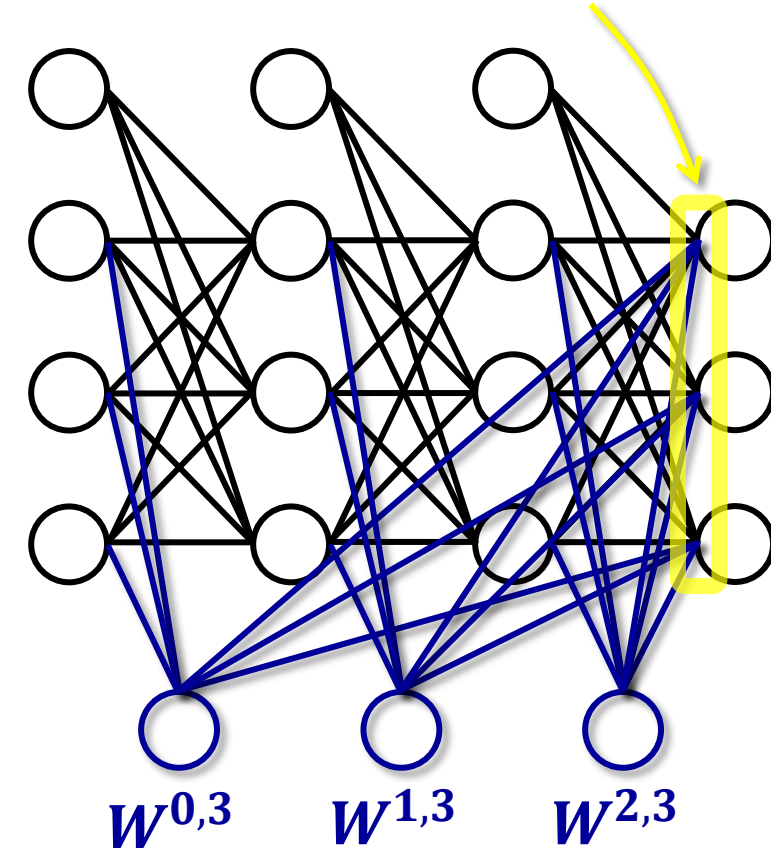
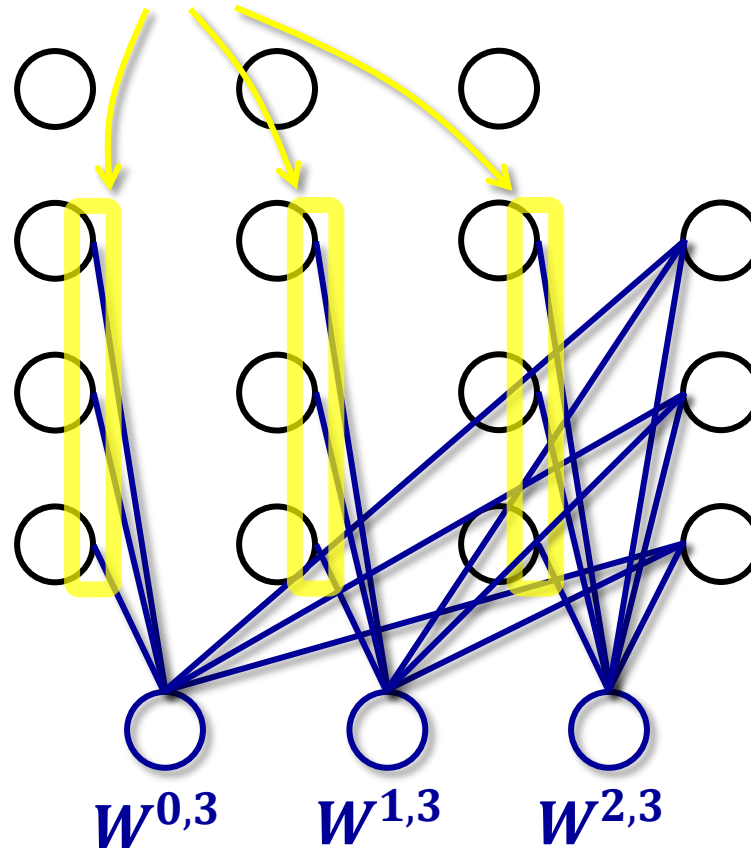
Our proposed approach: Skip2-LoRA

- Skip2-LoRA can reuse **forward** computation results

These values are needed to compute gradients of the adapters



Forward computation results of the base model are **cached**

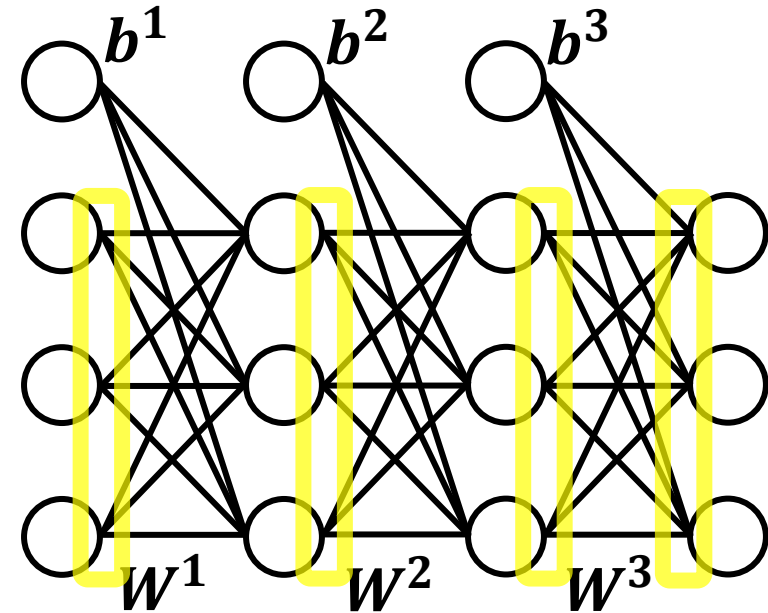


Base model + **Adapters** = **Skip-LoRA** [1]

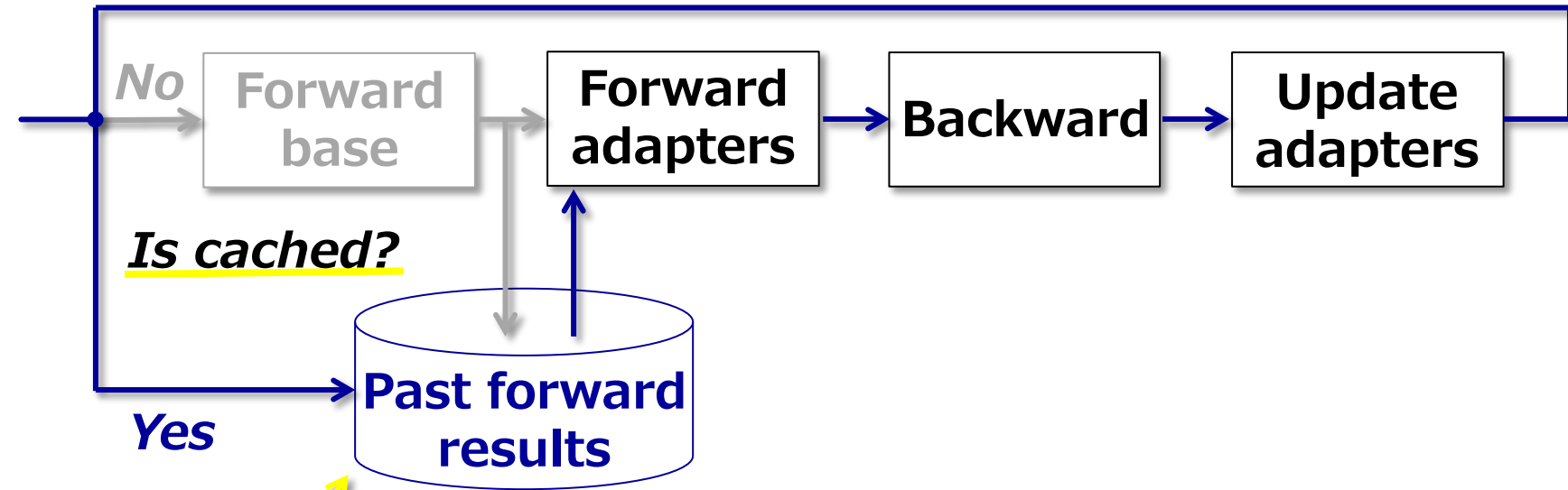
Our proposed approach: Skip2-LoRA

- Skip2-LoRA can reuse **forward** computation results

$$(\# \text{ iterations}) = (\# \text{ samples}) / (\text{Batch size}) \times (\# \text{ epochs})$$



Forward computation results of the base model are **cached**



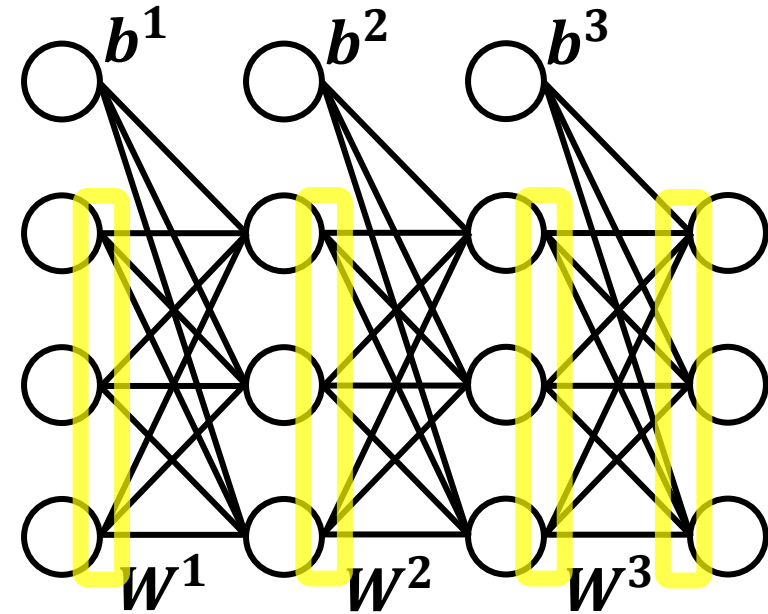
iterations is typically larger than 1;
so, we can skip most of the forward computation [1]

Base model

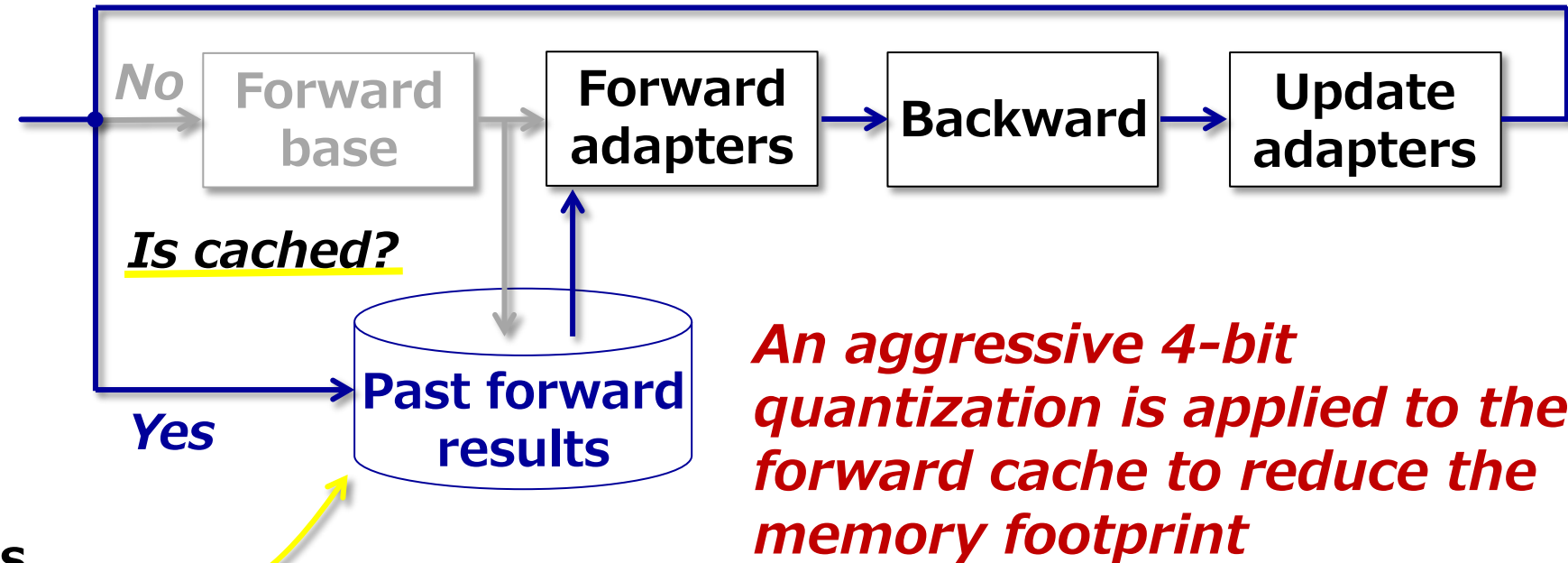
Our proposed approach: Skip2-LoRA

- Skip2-LoRA can reuse **forward** computation results

$$(\# \text{ iterations}) = (\# \text{ samples}) / (\text{Batch size}) \times (\# \text{ epochs})$$



Forward computation results of the base model are **cached**

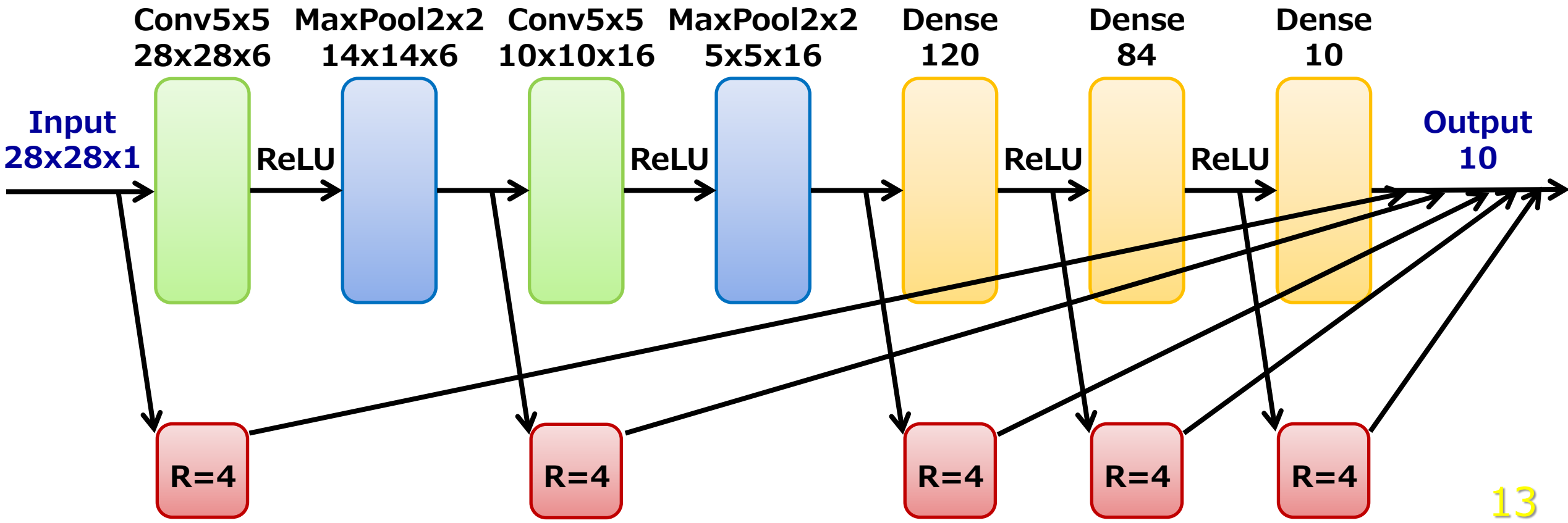
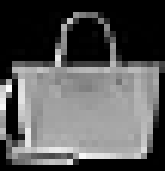


iterations is typically larger than 1;
so, we can skip most of the forward computation [1]

Skip2-LoRA for CNNs: Model

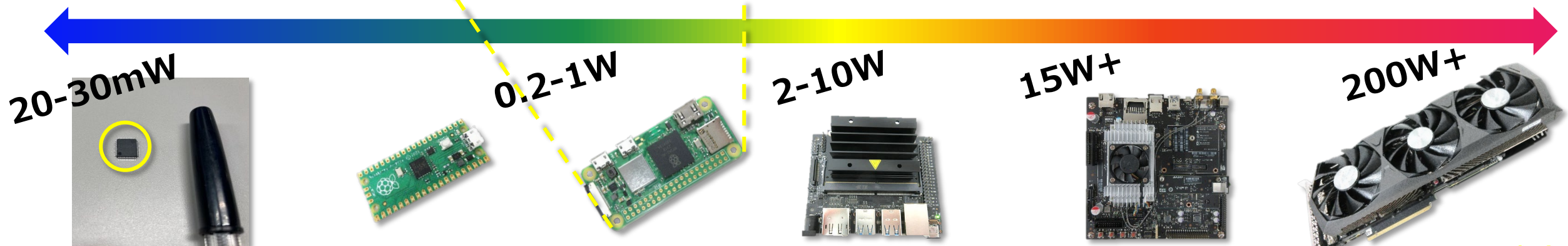
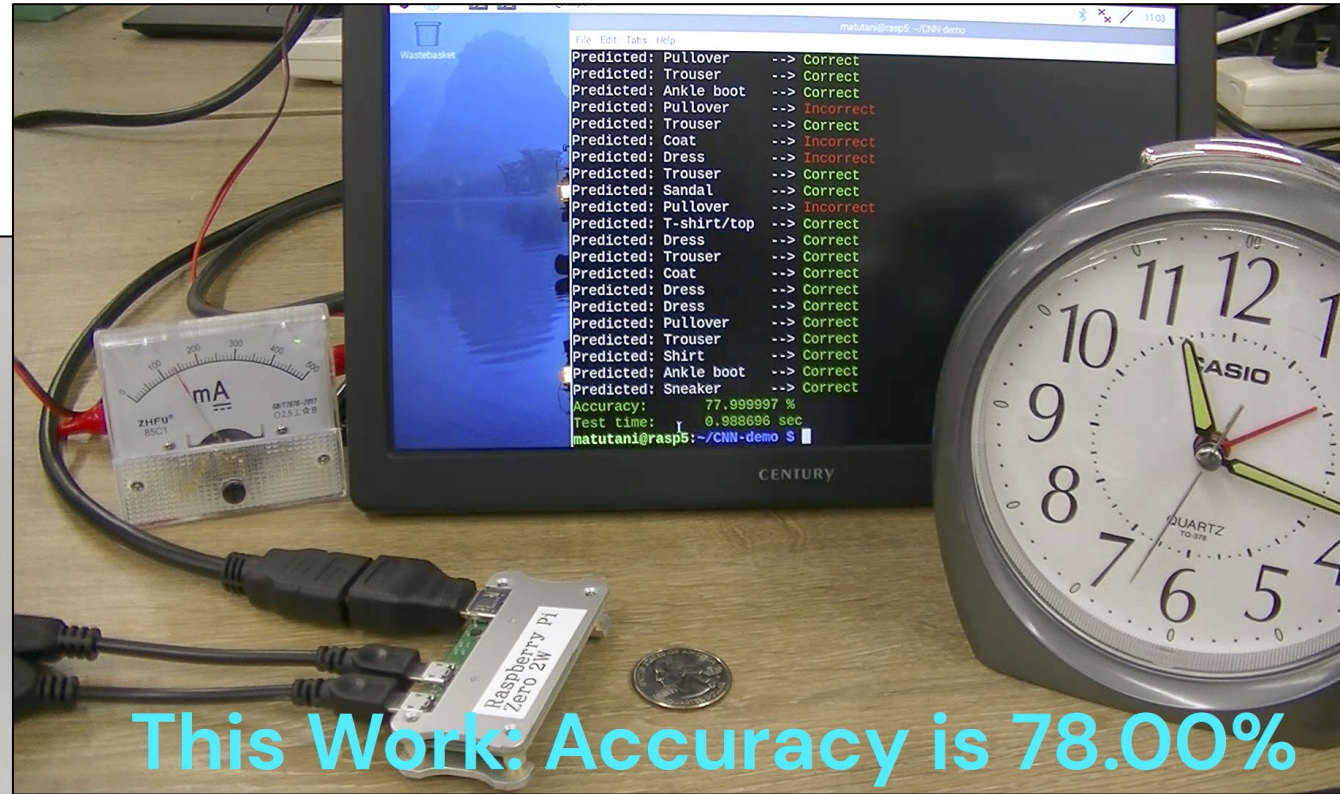
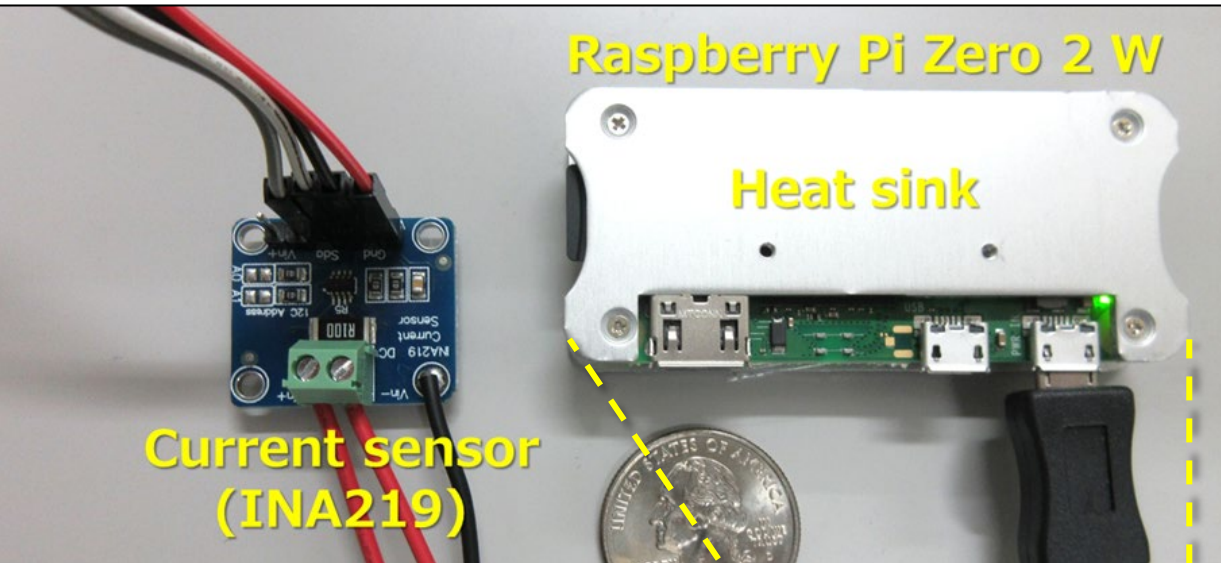
- Pretrained with fashion-MNIST
- Finetuned with **Rotated** fashion-MNIST (75 deg)
- Tested with **Rotated** fashion-MNIST (75 deg)

This gap reduces accuracy



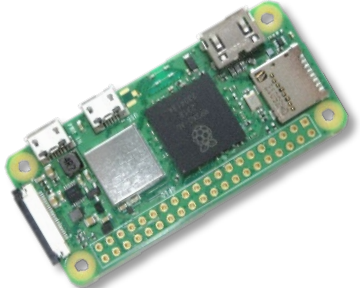
Skip2-LoRA for CNNs: Raspberry Pi

- **Raspberry Pi Zero 2W**
ARM Cortex-A53 @1GHz



Skip2-LoRA for CNNs: Results

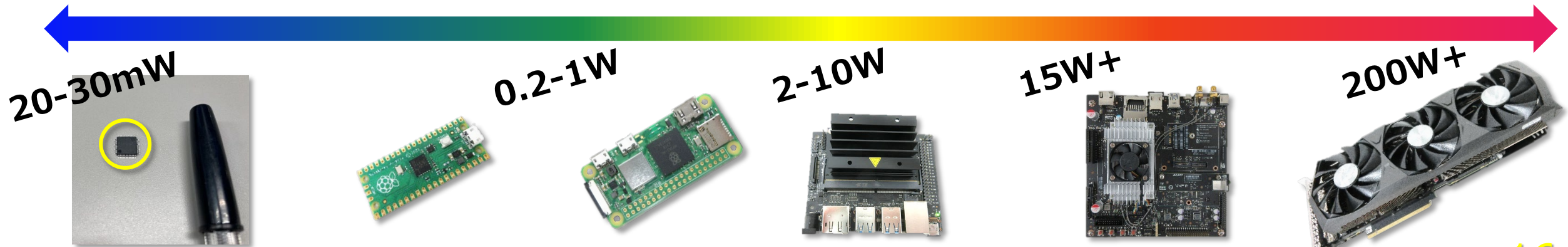
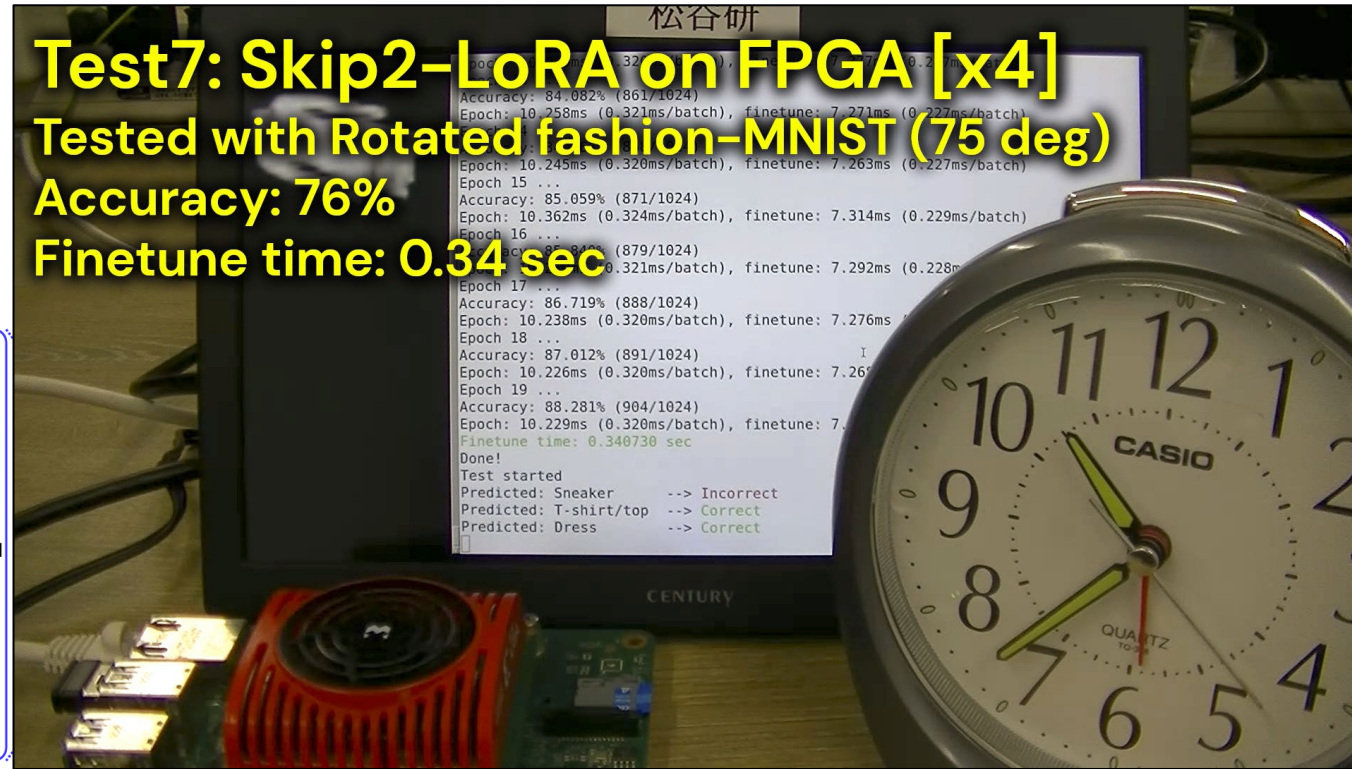
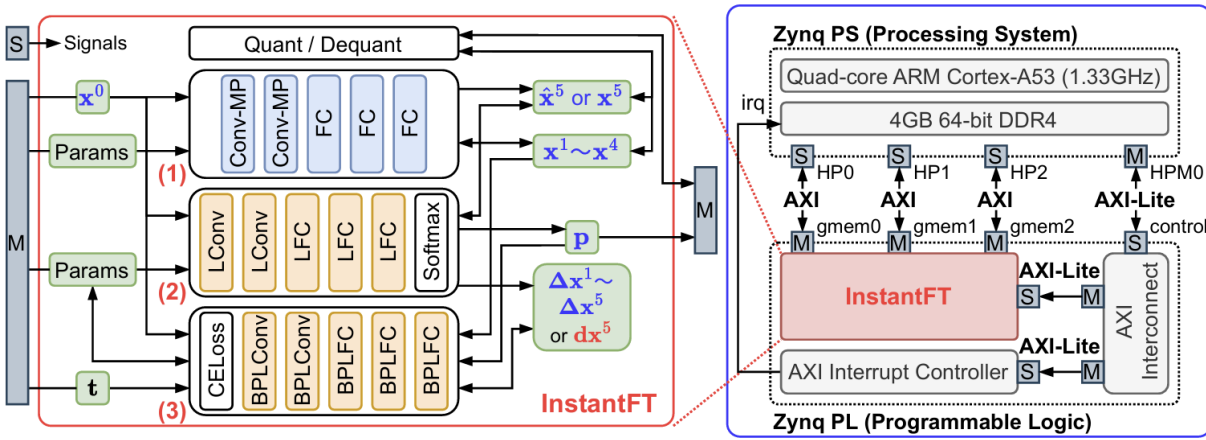
- In this work, Skip2-LoRA [1] is applied to **CNNs**
- An aggressive **4-bit quantization** is applied to the **forward cache** to reduce the memory footprint

Model	Accuracy	FT time @RPZ2	Cache size
No Finetuning (FT)	9.18 %		 Raspberry Pi Zero 2W (aka \$15 computer)
FT-Last	60.94 %	18.09 sec	
LoRA-Last	53.81 %	18.09 sec	
LoRA-All	75.59 %	114.15 sec	
Skip-LoRA	73.54 %	19.84 sec	
Skip2-LoRA	73.54 %	3.90 sec	
Quant Skip2-LoRA	74.02 %	4.27 sec	7,336 kB 1,036 kB

*Number of FT samples: 1024, Number of epochs for FT: 10

Skip2-LoRA for CNNs: Low-cost FPGA

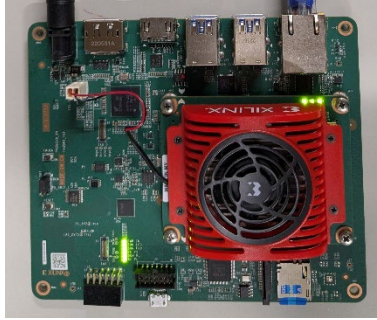
- **Low-cost FPGA board**
AMD KV260 Starter Kit
Finetune time: 0.34 sec



Skip2-LoRA for CNNs: Summary

- Ultra-fast **on-device CNN finetuning** for resource-limited edge devices

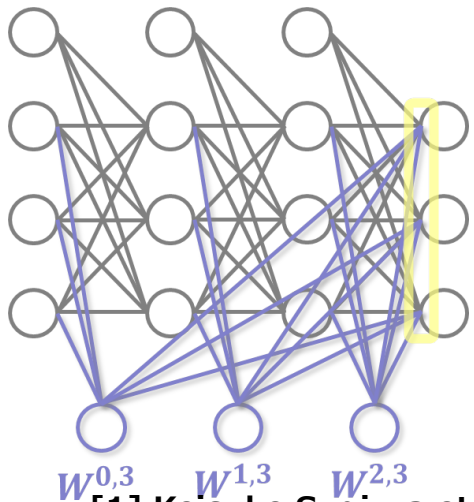
Our approach: **algorithm** and **architecture** codesign



Step 1:

New connection pattern

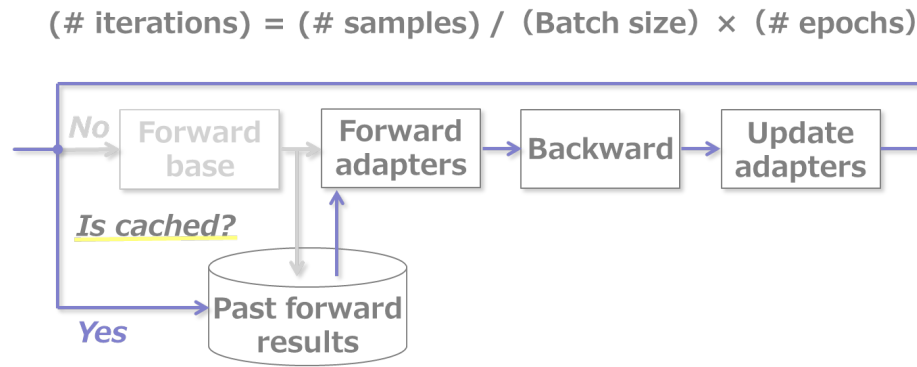
114.15 sec $\rightarrow$ 19.84 sec



Step 2:

Reusing computation results

19.84 sec $\rightarrow$ 4.27 sec



Step 3:

FPGA-based acceleration

4.27 sec $\rightarrow$ **0.34 sec**

