

# Solving Semi-Supervised Few-Shot Learning from an Auto-Annotation Perspective

Tian Liu<sup>1</sup>, Anwesha Basu<sup>1</sup>, James Caverlee<sup>1</sup>, Shu Kong<sup>2</sup>

<sup>1</sup>Texas A&M University, <sup>2</sup>University of Macau

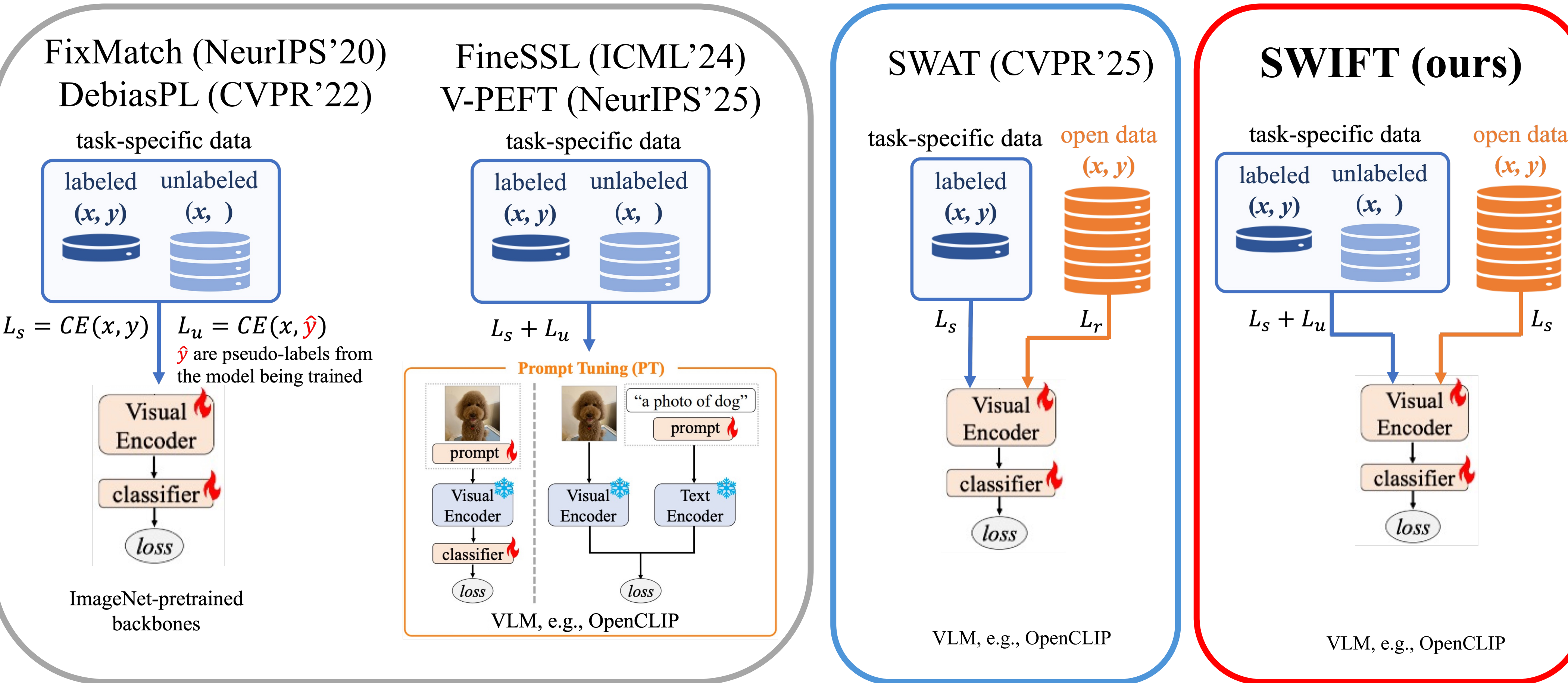


paper & code



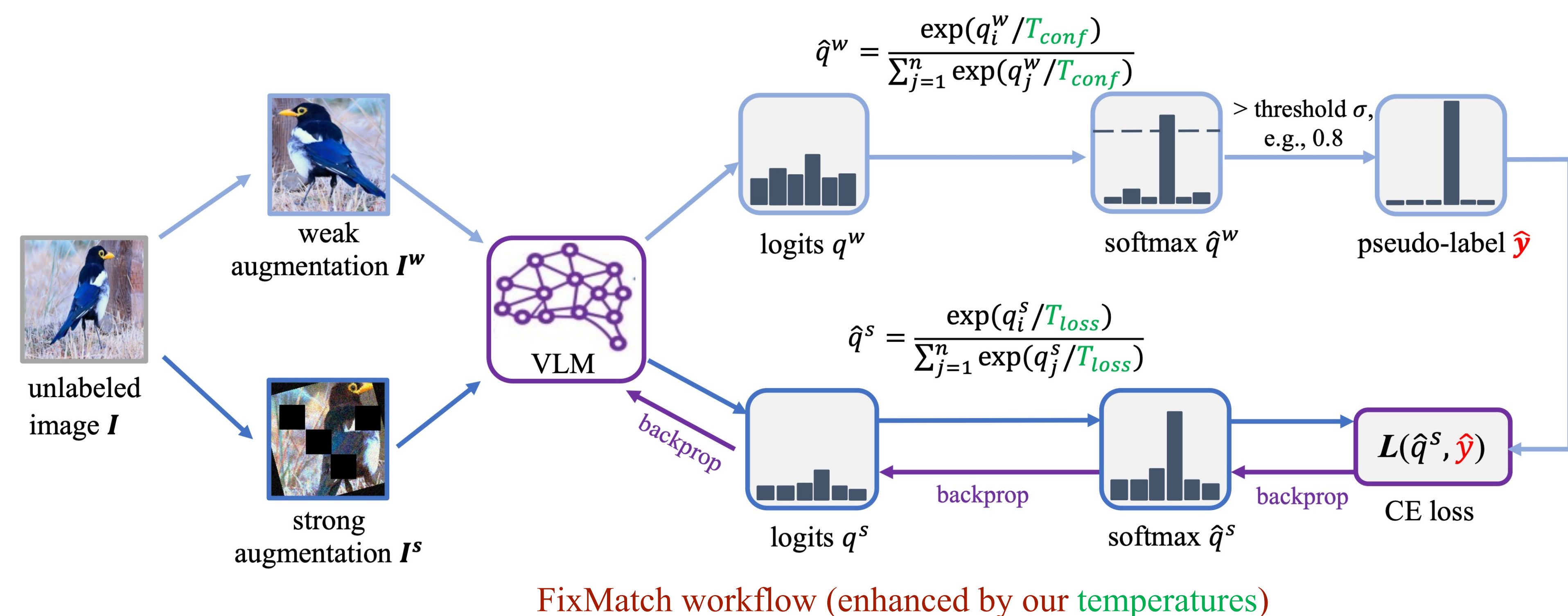
## Motivation

- Semi-Supervised Few-Shot Learning (SSFSL) mirrors real-world *auto-annotation*, aiming to learn a model over a few labeled and abundant unlabeled task-specific examples so to automatically annotate the unlabeled ones.
- Existing SSL or SSFSL methods either finetune an ImageNet-pretrained backbone or prompt a frozen VLM, while recent FSL method has achieved significant progress by *finetuning* a VLM **and retrieving task-relevant open data**.
- To bridge the gap, we are motivated to explore these techniques for SSFSL.



## Challenges

- We start by finetuning the VLM OpenCLIP ViT-B/32 with the representative SSL method *FixMatch* (NeurIPS'20) using the default hyperparameters (e.g., confidence threshold 0.8).
- Yet, it barely rivals with zero-shot methods **REAL** (CVPR'24) and significantly underperforms the few-shot finetuning method **FS-FT** (CVPR'25), which does not exploit unlabeled data.

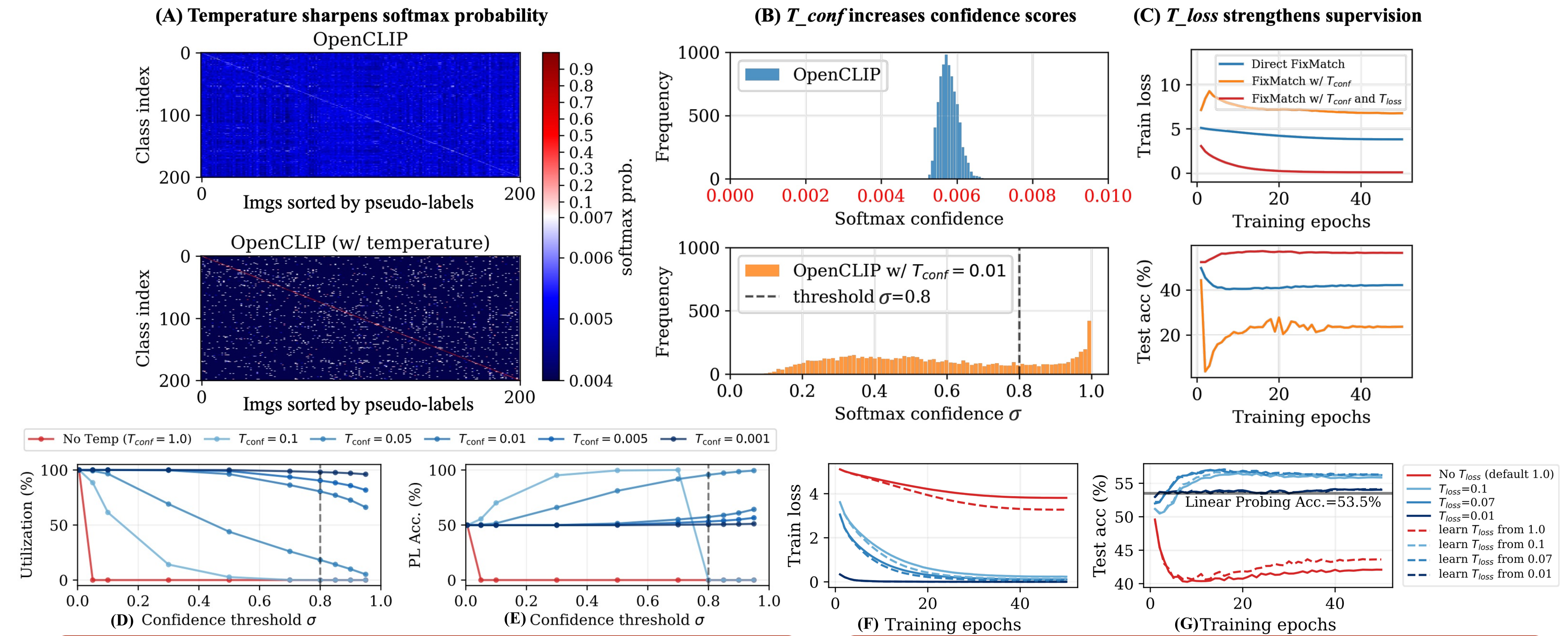


FixMatch workflow (enhanced by our temperatures)

| method        | training data | backbone        | mean acc. over five datasets |                       |                       |
|---------------|---------------|-----------------|------------------------------|-----------------------|-----------------------|
|               |               |                 | 4-shot                       | 8-shot                | 16-shot               |
| FS-FT [42]    | L             | VLM-OpenCLIP    | 61.0                         | 65.8                  | 69.1                  |
| FixMatch [57] | L + U         | INet-pretrained | 30.7                         | 45.3                  | 60.1                  |
| FixMatch      | L + U         | VLM-OpenCLIP    | 39.3 <sup>-21.7</sup>        | 49.9 <sup>-15.9</sup> | 57.2 <sup>-11.9</sup> |
| DebiasPL [67] | L + U         | INet-pretrained | 34.5                         | 49.0                  | 62.7                  |
| DebiasPL      | L + U         | VLM-OpenCLIP    | 39.6 <sup>-21.4</sup>        | 49.8 <sup>-16.0</sup> | 57.1 <sup>-12.0</sup> |

## Failure Diagnostics & Remedies by Temperatures

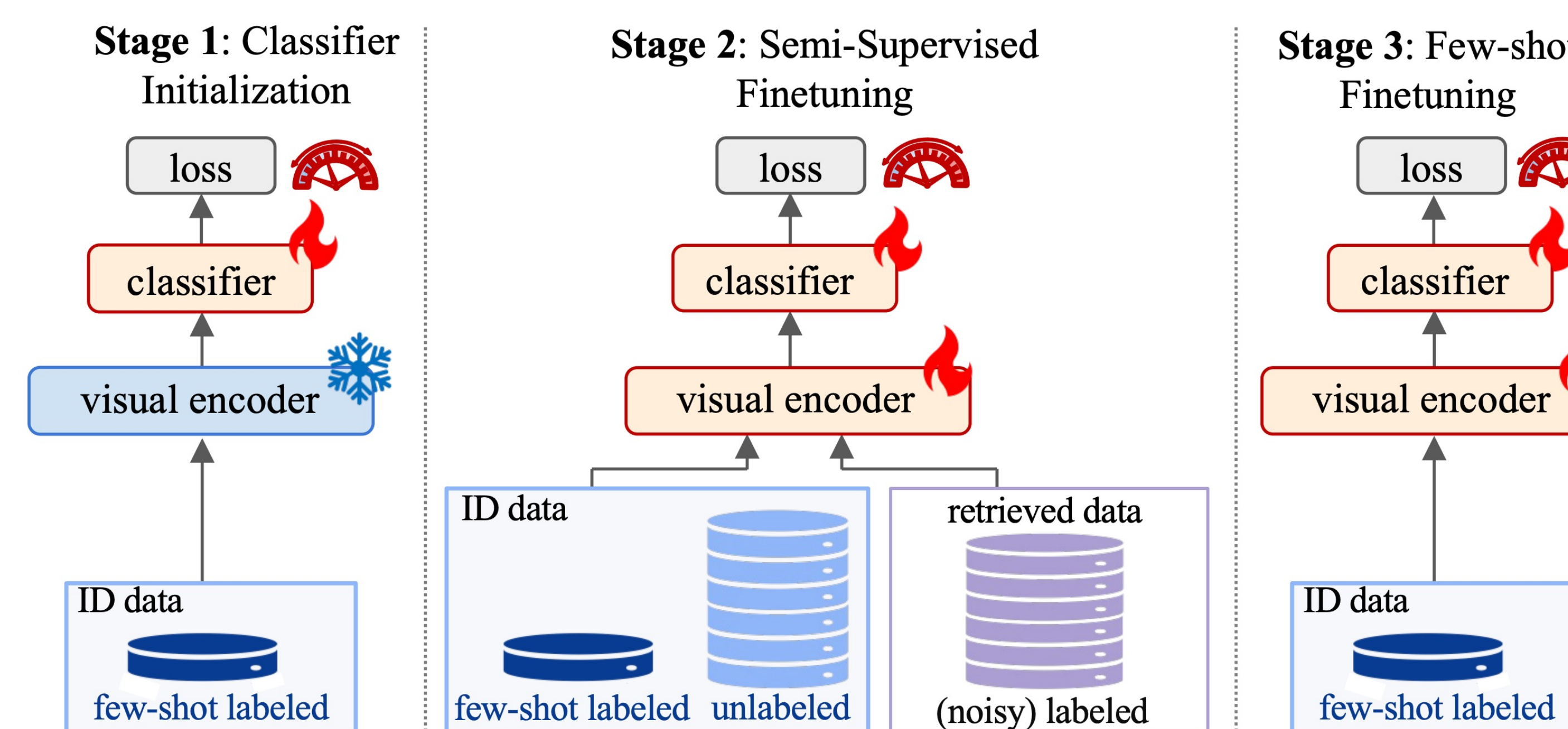
- Root cause: *VLMs produce "flat" distribution of softmax probabilities*, which results in
  - zero utilization of unlabeled data** (w/ default confidence threshold of 0.8) →  $T_{conf}$  sharpens softmax confidences.
  - weak training supervision** that prevents effective finetuning the VLM →  $T_{loss}$  strengthens training supervision.



## Method

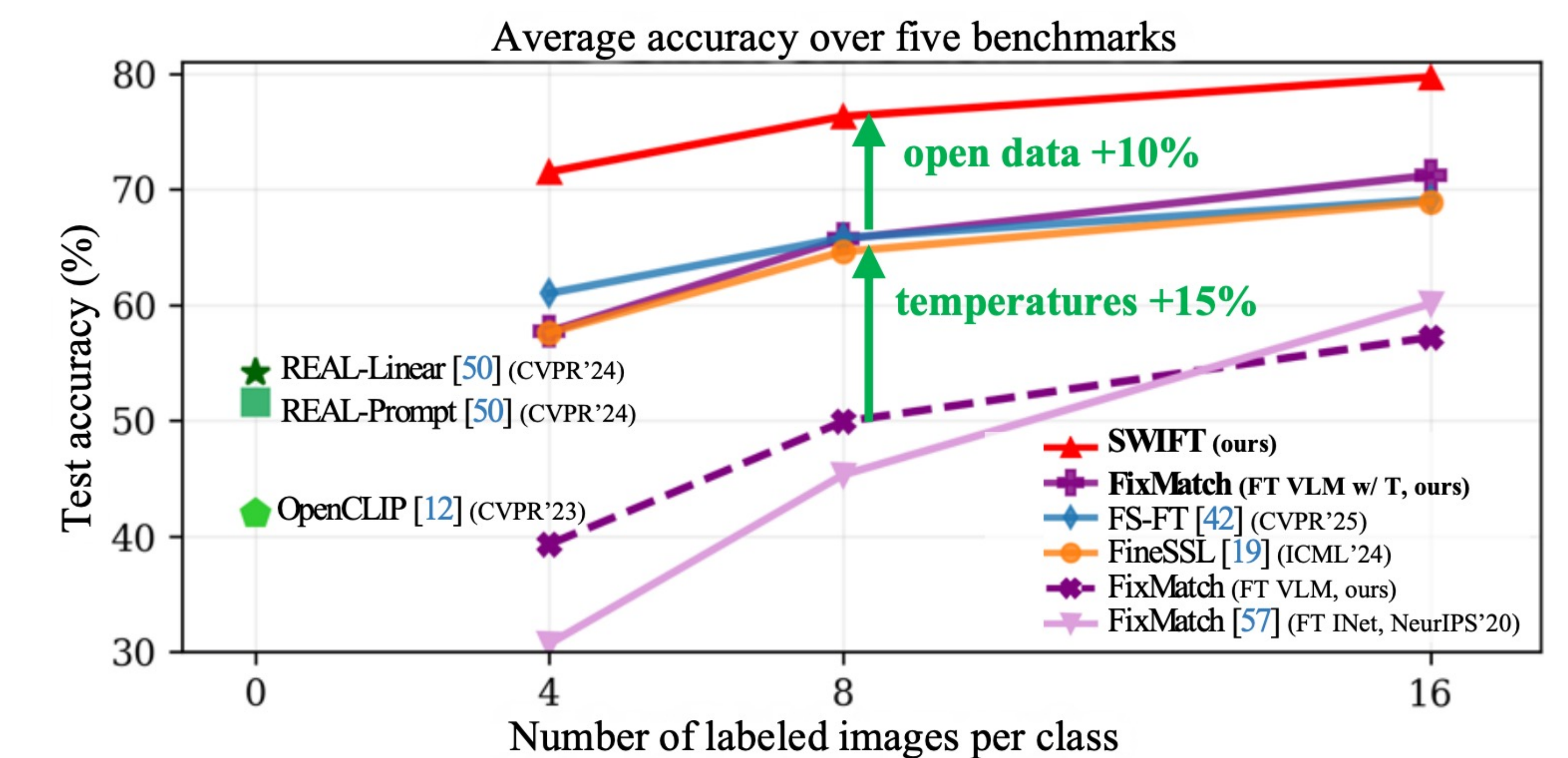
- Our **SWIFT** (stage-wise finetuning with temperatures) exploits retrieved data via a stage-wise training to mitigate domain gaps.

| Dataset                           | Few-shot data | Unlabeled ID data | Retrieved data |
|-----------------------------------|---------------|-------------------|----------------|
| Semi-Aves<br>Phoenicopterus ruber |               |                   |                |
| Aircraft<br>MD-90                 |               |                   |                |
| Cars<br>Suzuki SX4 Sedan 2012     |               |                   |                |



## SOTA Performance

- SWIFT** surpasses prior methods by >5% accuracy.



| method                      | training data | mean accuracy over five datasets |                      |                      |                       |                       |                      |
|-----------------------------|---------------|----------------------------------|----------------------|----------------------|-----------------------|-----------------------|----------------------|
|                             |               | OpenCLIP ViT-B/32                |                      |                      | DINOv2 ViT-B/14       |                       |                      |
|                             |               | 4-shot                           | 8-shot               | 16-shot              | 4-shot                | 8-shot                | 16-shot              |
| FS-FT [42] CVPR'25          | L             | 61.0                             | 65.8                 | 69.1                 | 56.5                  | 71.1                  | 79.6                 |
| FixMatch [57] direct        | L+U           | 39.3                             | 49.9                 | 57.2                 | 50.2                  | 66.3                  | 77.1                 |
| + stage-1: Classifier Init. | L+U           | 56.3 <sup>+17.0</sup>            | 59.8 <sup>+9.9</sup> | 62.2 <sup>+5.0</sup> | 56.7 <sup>+6.5</sup>  | 71.5 <sup>+5.2</sup>  | 80.3 <sup>+3.1</sup> |
| + stage-2: Semi-Sup. FT     | L+U+R         | 68.8 <sup>+11.1</sup>            | 73.3 <sup>+7.6</sup> | 77.3 <sup>+6.1</sup> | 76.3 <sup>+19.6</sup> | 82.0 <sup>+10.5</sup> | 86.3 <sup>+6.0</sup> |
| + stage-3: FS-FT (SWIFT)    | L+U+R         | 71.5 <sup>+2.7</sup>             | 76.3 <sup>+3.0</sup> | 79.7 <sup>+2.4</sup> | 78.2 <sup>+1.9</sup>  | 84.4 <sup>+2.4</sup>  | 87.8 <sup>+1.5</sup> |

## References

- "The Neglected Tails in Vision Language Model". CVPR, 2024.
- "Few-shot Recognition via Stage-wise Retrieval-augmented Finetuning". CVPR 2025.
- "FixMatch: Simplifying semi-supervised learning with consistency and confidence". NeurIPS 2020.
- "Debiased Learning from naturally imbalanced pseudo-labels". CVPR 2022.
- "Erasing the bias: Fine-tuning foundation models for semi-supervised learning". ICML 2024.
- "Revisiting semi-supervised learning in the era of foundation models". NeurIPS 2025.