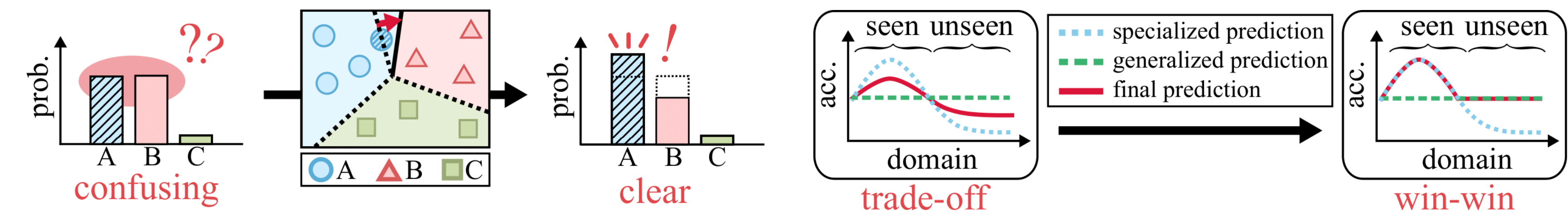


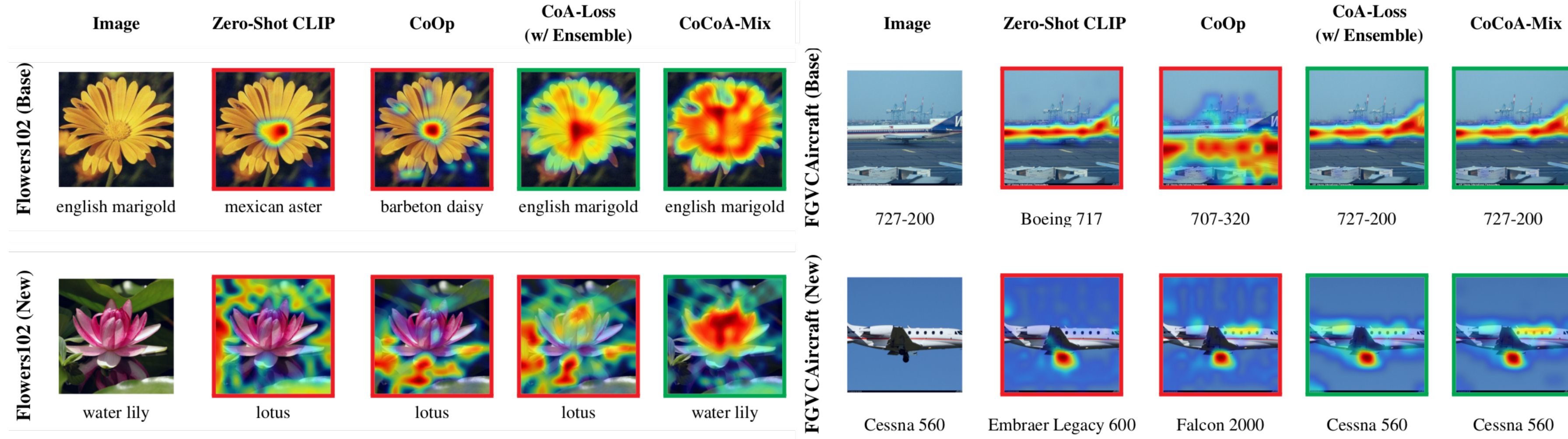
Motivation & Contribution

TL;DR

- Existing **prompt tuning** methods struggle with **class confusion** due to frozen encoders and often **sacrifice task specialization** for better generalization across unseen domains.



- We propose **CoCoA-Mix**, a novel framework combining a **Confusion-Aware Loss (CoA-loss)** to refine decision boundaries for confusing classes, and **Confidence-Aware Weights (CoA-weights)** to balance specialization and generalization in a mixture model.

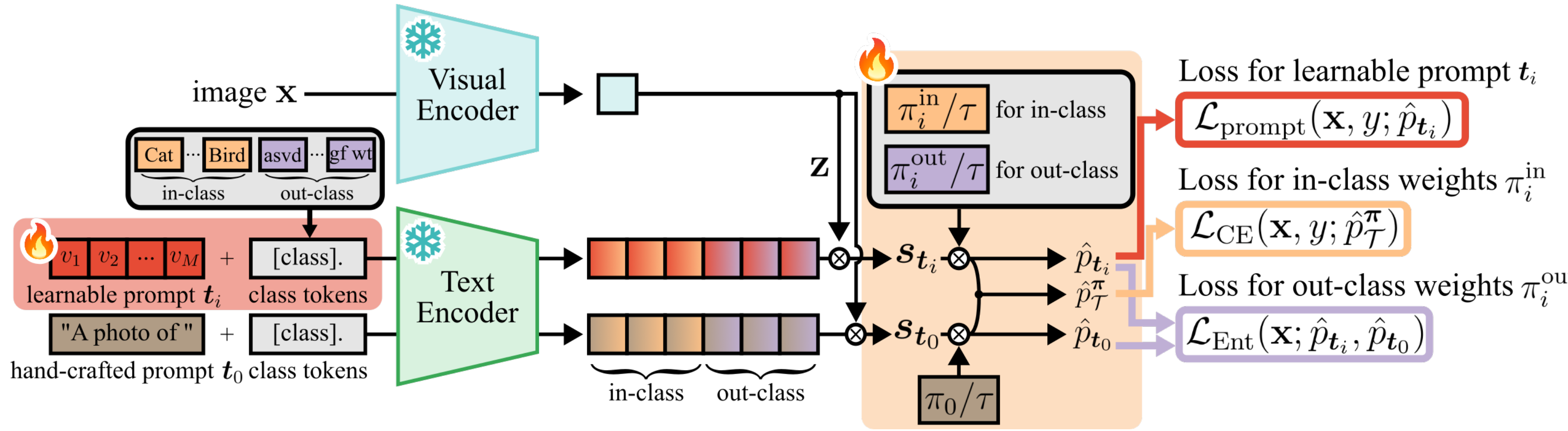


Contribution

- We provide a **mathematical framework** demonstrating that **specialization** and **generalization** can be improved simultaneously.
- We propose a **CoCoA-Mix framework**, consisting of CoA-loss and CoA-weights. **CoA-loss** boosts specialization by improving classification for **confusing cases**, while **CoA-weights** improve generalization by adjusting the weights of individual prompts in the mixture model based on their **confidence over class domains**.
- The proposed method achieves average harmonic mean improvements of 15.28% and 3.28% over zero-shot CLIP in **base-to-new generalization** and **cross-dataset transfer**, respectively; it also improves the average accuracy in **few-shot class-incremental learning** by 5.6%p.

Methodology

CoCoA-Mix (Confusion-and-Confidence-Aware Mixture Model)



- Mixture Model Framework:** A theoretical mixture model that **aggregates predictions from multiple prompts**, enabling a unified analysis of both specialization (task-specific accuracy) and generalization (performance on unseen domains).

The expected error of the mixture model can be bounded as follows:

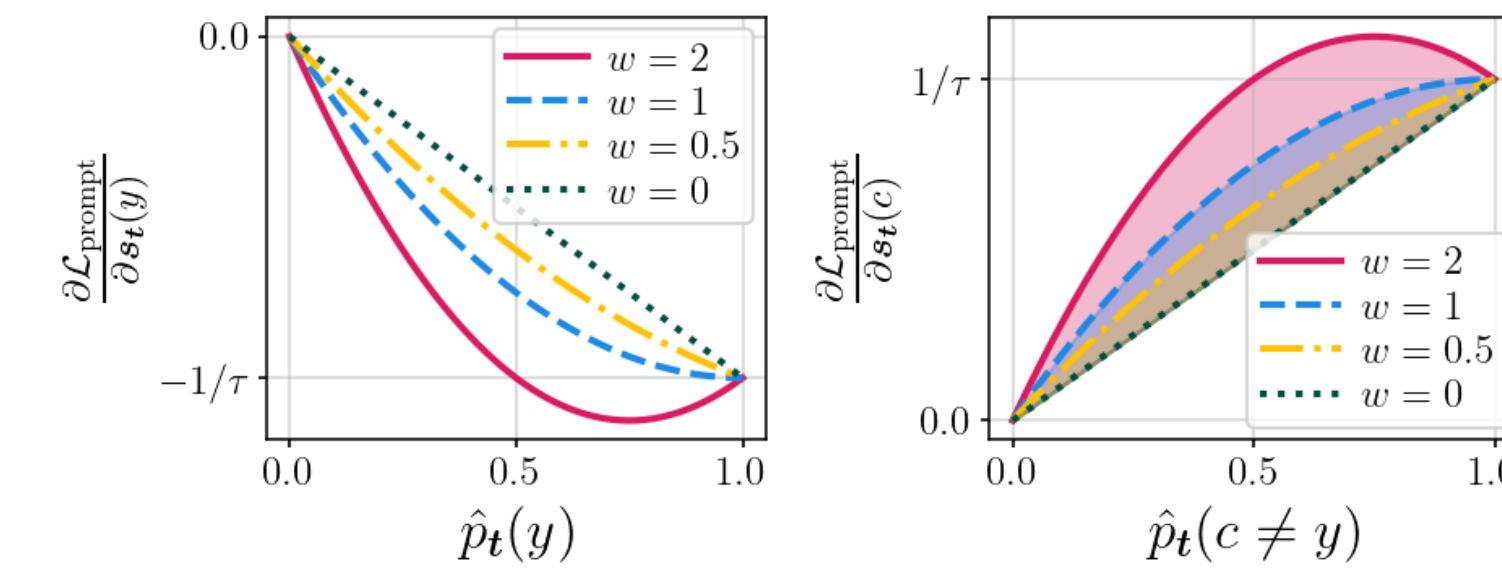
$$\epsilon_T(\hat{p}_T^\pi) \leq \sum_{i=0}^K \lambda_i \left(\underbrace{\pi_i^{\text{in}} \epsilon_{T_i}(\hat{p}_{t_i})}_{\text{specialization error}} + \underbrace{\sum_{\substack{j=0 \\ j \neq i}}^K \pi_j^{\text{out}} \epsilon_{T_i}(\hat{p}_{t_j})}_{\text{generalization error}} \right)$$

- $\hat{p}_t$: prediction with a text prompt t
- $\mathcal{T} = \{t_0, t_1, \dots, t_K\}$: K+1 different prompts
- $\pi = \{(\pi_0^{\text{in}}, \pi_0^{\text{out}}), \dots, (\pi_K^{\text{in}}, \pi_K^{\text{out}})\}$: a set of non-negative weights satisfying $\pi_i^{\text{in}} + \sum_{j \neq i} \pi_j^{\text{out}} = 1$
- The class set of the target domain $\mathcal{D}_T$ be partitioned into K+1 disjoint subsets, with corresponding sub-domains $\mathcal{D}_{T_0}, \mathcal{D}_{T_1}, \dots, \mathcal{D}_{T_K}$, such that $\mathcal{D}_T = \bigsqcup_{i=0}^K \mathcal{D}_{T_i}$
- $\lambda_i = \Pr_{(\mathbf{x}, y) \sim \mathcal{D}_T}[(\mathbf{x}, y) \in \mathcal{D}_{T_i}]$: the probability that a sample from $\mathcal{D}_T$ belongs to $\mathcal{D}_{T_i}$

- Confusion-Aware Loss (CoA-loss):** To enhance **specialization**, CoA-loss assigns **higher gradients to confusing examples**—samples with similar probabilities across multiple classes—thereby refining decision boundaries in prompt tuning.

$$\mathcal{L}_{\text{prompt}} = \mathcal{L}_{\text{CE}} + w \mathcal{L}_{\text{CoA}}$$

$$\mathcal{L}_{\text{CoA}}(\mathbf{x}, y; \hat{p}_t) = 1 - \hat{p}_t(y)$$



- Confidence-Aware Weights (CoA-weights):** To improve **generalization** without compromising specialization, the model **adjusts mixture weights based on each prompt's confidence in the class domain**, assigning more weight to confident prompts and less to uncertain prompts.

In-Class Weights

$$\pi_i^{\text{in}} = \arg \min_{\pi_i^{\text{in}}} \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}_{S_i}} [\mathcal{L}_{\text{CE}}(\mathbf{x}, y; \hat{p}_T^t)]$$

In-Class Weights

$$\pi_i^{\text{out}} = \arg \min_{\pi_i^{\text{out}}} \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}_{S_i}} [\mathcal{L}_{\text{Ent}}(\mathbf{x}; \hat{p}_{t_i}, \hat{p}_{t_0})]$$

$$\mathcal{L}_{\text{Ent}} = \max(0, H(\hat{p}_{t_0}) - H(\hat{p}_{t_i}) + d)$$

Experiments & Analysis

Comparative Study

- CoCoA-Mix is evaluated on 11 datasets across three tasks—**base-to-new generalization**, **few-shot class-incremental learning (FSCIL)**, and **cross-dataset transfer**—demonstrating its broad applicability.

Base-to-New Generalization

METHOD	AVERAGE			IMAGENET			CALTECH101			FOOD101		
	BASE	NEW	H	BASE	NEW	H	BASE	NEW	H	BASE	NEW	H
CLIP	65.14	68.78	66.82	64.43	60.04	62.16	90.64	91.16	90.90	83.58	84.95	84.26
CoOp	77.23	68.56	71.33	73.72 ± 0.29	64.94 ± 0.87	69.05	97.16 ± 0.16	93.92 ± 0.80	95.51	89.19 ± 0.19	88.45 ± 0.89	88.81
ProGrad	78.74	72.19	75.06	74.81 ± 0.29	66.68 ± 0.26	70.51	97.50 ± 0.08	95.49 ± 0.27	96.48	89.93 ± 0.08	89.93 ± 0.58	89.63
KgCoOp	78.67	74.62	76.38	75.44 ± 0.08	69.43 ± 0.29	72.31	97.61 ± 0.33	94.80 ± 0.45	96.18	90.26 ± 0.11	91.25 ± 0.15	90.75
MAPLE	77.14	72.91	74.69	75.40 ± 0.29	70.43 ± 0.12	72.83	97.47 ± 0.31	93.77 ± 1.11	95.57	89.37 ± 0.54	90.77 ± 0.54	90.06
DePT	79.20	66.36	71.78	73.50 ± 0.22	70.00 ± 0.16	71.71	97.83 ± 0.05	95.83 ± 0.25	96.82	89.80 ± 0.05	88.10 ± 0.16	88.94
CoA-Loss	79.12	73.66	76.15	75.68 ± 0.00	67.98 ± 0.31	71.62	97.94 ± 0.14	94.54 ± 0.24	96.21	90.11 ± 0.18	90.87 ± 0.42	90.49
CoCoA-MIX	79.31	75.10	77.03	75.47 ± 0.09	68.92 ± 0.10	72.04	98.02 ± 0.03	94.39 ± 0.10	96.17	90.09 ± 0.16	90.93 ± 0.09	90.50

METHOD	OXFORDPETS			STANFORDCARS			FLOWERS102			FGVCAIRCRAFT		
	BASE	NEW	H	BASE	NEW	H	BASE	NEW	H	BASE	NEW	H
CLIP	90.01	94.24	92.07	55.37	66.65	60.49	69.23	73.90	71.49	19.51	24.60	21.76
CoOp	94.10 ± 0.73	94.42 ± 4.17	94.16	69.54 ± 0.75	71.39 ± 1.28	70.44	90.60 ± 1.50	67.00 ± 1.04	77.01	26.17 ± 7.89	19.50 ± 11.94	11.46
ProGrad	95.00 ± 0.31	97.36 ± 0.42	96.16	71.45 ± 0.39	73.16 ± 0.58	72.29	91.36 ± 0.63	74.92 ± 0.90	82.32	34.21 ± 1.99	28.53 ± 2.08	30.97
KgCoOp	94.65 ± 0.15	97.59 ± 0.08	96.10	68.64 ± 0.35	74.96 ± 0.53	71.66	90.09 ± 0.63	76.31 ± 0.42	82.63	33.43 ± 0.56	32.27 ± 1.19	32.81
MAPLE	94.80 ± 0.94	97.67 ± 0.21	96.21	67.97 ± 0.29	74.40 ± 0.45	71.04	88.03 ± 1.62	73.43 ± 0.49	80.06	31.67 ± 0.66	33.13 ± 2.38	32.29
DePT	94.00 ± 0.29	88.63 ± 0.78	91.23	71.83 ± 0.52	59.27 ± 0.76	64.94	94.53 ± 0.53	66.30 ± 1.42	77.92	35.93 ± 0.93	24.33 ± 0.09	29.01
CoA-Loss	94.90 ± 0.49	97.93 ± 0.08	96.39	72.70 ± 0.11	73.07 ± 1.27	72.87	88.89 ± 1.75	75.58 ± 1.31	81.67	33.91 ± 0.68	32.47 ± 0.37	33.17
CoCoA-MIX	95.16 ± 0.38	97.60 ± 0.09	96.36	73.08 ± 0.25	74.97 ± 0.08	74.01	91.04 ± 1.79	77.37 ± 0.38	83.64	33.51 ± 0.28	34.15 ± 0.14	33.83

FSCIL

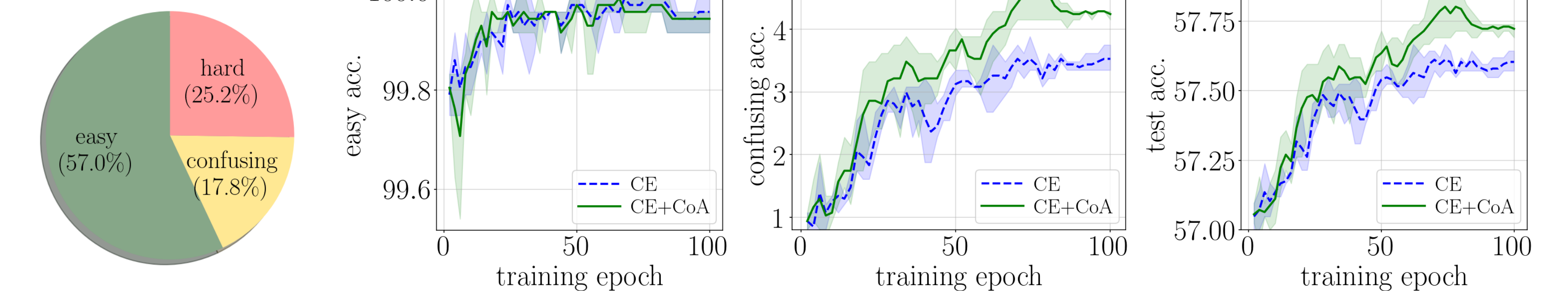
(Few-Shot Class-Incremental Learning)

METHOD	0	1	2	3	4	5	6	7	8	MEAN↑	PD↓
L2P	89.9	86.0	81.8	80.3	80.0	74.6	73.2	72.6	65.0	78.2	24.9
CLIP-ZSL	—	—	—	—	—	—	—	—	—	77.9	—
CoOp-FSCIL	88.6	78.9	77.5	76.0	76.8	78.3	79.2	79.8	79.3	79.4	9.3
FACT w/ CLIP	87.8	84.0	81.4	78.0	77.8	76.3	75.0	72.5	71.9	78.3	15.9
FSPF-FSCIL	86.9	83.1	81.9	80.7	80.4	79.9	80.1	79.9	79.4	81.4	7.5
CoCoA-MIX (OURS)	88.2	85.6	84.6	82.7	82.8	82.5	82.3	81.8	80.8	83.5	7.4

Cross-Dataset Transfer

METHOD	SOURCE	TARGET	H
CLIP	66.73	64.89	63.97
CoOp	69.06 ± 0.43	59.88	61.52
ProGrad	70.21 ± 0.16	62.36	63.58
KgCoOp	70.52 ± 0.05	64.45	65.17
MAPLE	69.53 ± 0.39	65.24	65.26
DePT	68.03 ± 0.05	65.06	64.42
CoCoA-MIX	70.85 ± 0.09	65.27	66.07

Analysis & Ablation Study



Backbone	Method	Loss	Ensemble	Base	New	H
ViT-B/16	CLIP	-	-	65.1	68.8	66.8
	CLIP (w/ Ensemble)	-	Uniform Ensemble	70.6	74.3	72.3
	CoA-loss (w/o Ensemble)	CoA-loss	-	78.6	68.5	72.9
	CoA-loss (w/ Ensemble)	CoA-loss	Uniform Ensemble	79.1	73.7	76.2
	CoCoA-Mix	CoA-loss	CoA-weights	79.3	75.1	77.0
RN50 + RN101 + ViT-B/16 + ViT-B/32	T _{En} + KgCoOp	CoCoOp	Sample-Aware Weight Generator	84.1	75.5	79.2
	T _{En} + CoA-Loss	CoA-loss	Sample-Aware Weight Generator	85.3	75.2	79.5
	CoCoA-Mix	CoA-loss	CoA-weights	85.4	76.3	80.3

- Analyses show that CoA-loss significantly boosts performance on **confusing samples**, while CoA-weights improve generalization across domains without sacrificing specialization.

