

Co-HSF: Resource-Efficient One-Shot Semi-Supervised Adaptation of Histopathology Foundation Models

Luca Cerny Oliveira, M.S.^{1#}, Kartik Patwari, M.S.¹, Xiaoguang Zhu, M.S.¹, Sen-Ching Cheung, Ph.D.², Brittany N. Dugger, Ph.D.³

and Chen-Nee Chuah, Ph.D.¹

¹Department of Electrical and Computer Engineering, University of California Davis,

²Department of Electrical and Computer Engineering, University of Kentucky, ³Department of Pathology and Laboratory Medicine, University of California Davis.

Indicates the corresponding author

UCDAVIS
HEALTH

Alzheimer's Disease
Research Center

Histopathology foundation models

demonstrate promising zero-shot capabilities and achieve state-of-the-art (SOTA) performance after fine-tuning (Tab. 1). Current fine-tuning methods face key limitations:

- Overfitting on one-shot datasets
 - Need for detailed image captions for fine-tuning
 - High computational cost of training
 - Inability to use unlabeled data effectively
- Semi-supervised few-shot learning (SSFSL) offers an alternative, leveraging minimal labeled data alongside unlabeled samples. We study its use for foundation model fine-tuning and propose **Co-filtered Histopathology Semi-Supervised Few-Shot (Co-HSF)**:
- Dual-SSFSL training (teacher-student setup)
 - Novel co-filtering pseudo-labeling technique
 - Effectively exploits unlabeled data while reducing inference times

Co-filtering Pseudo-labeling

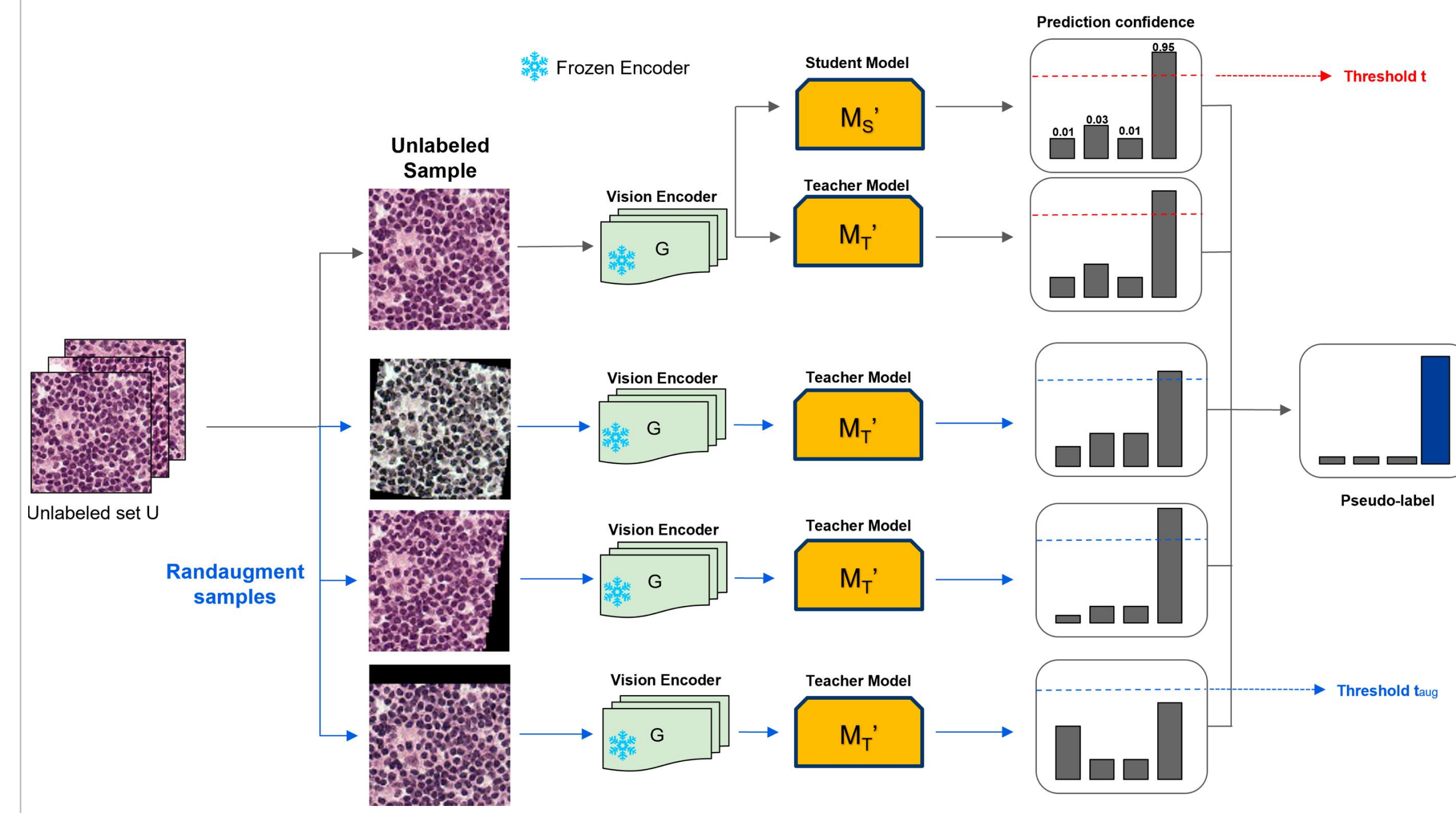


Fig. 2. Overview of the proposed Co-filtering pseudo-labeling strategy: the unlabeled set U is augmented by Randaugment⁶ and featurized by CLIP's pre-trained vision encoder $G(\cdot)$. Both student and teacher models evaluate on U and U_{aug} respectively. The prediction confidences from the evaluations are analyzed, if the prediction confidences for the evaluations of a single sample and their HFT-augmented samples are all above the thresholds t and t_{aug} (both calculated from hyperparameter T), we will add the sample to the pseudo-labeled set.

Results

- Tab. 1 shows Co-HSF consistently outperforms both zero-shot and competing fine-tuning approaches in one-shot settings.
- Tab. 3 and Tab. 4 shows Co-filtering outperforms competing SSFSL, generating a more accurate pseudo-labeled set
- Co-HSF demonstrates lowest memory usage and faster inference times than most compared baselines, while leading accuracy performance (see Tab. 4)

Learning Type	Algorithm	PCam		NCT	
		Accuracy (%)	F-1 Score (%)	Accuracy (%)	F-1 Score (%)
Zero-shot	CLIP	56.57	56.56	28.08	23.64
	CLIP (ViT-B-L16)	57.02	53.15	45.46	39.91
	Zhang et al. 2023	53.35	40.39	50.83	46.82
	PLIP (Huang et al. 2023)	68.91	68.76	47.06	46.70
	CONCH (Lu et al. 2024)	82.66	82.57	51.70	47.83
One-shot	SGD (Huang et al. 2023)	80.55 $\pm$ 3.61	80.41 $\pm$ 4.03	82.26 $\pm$ 5.30	81.54 $\pm$ 1.46
Linear Probing	Logistic (Lu et al. 2024)	80.75 $\pm$ 3.51	80.53 $\pm$ 4.25	84.10 $\pm$ 1.38	83.49 $\pm$ 1.19
One-shot	CoOp (Zhou et al. 2022)	57.80 $\pm$ 0.01	47.04 $\pm$ 0.65	24.30 $\pm$ 0.02	31.58 $\pm$ 6.45
Fine-tuning	CLIPath (Lai et al. 2023)	72.18 $\pm$ 4.04	70.00 $\pm$ 8.44	69.24 $\pm$ 1.46	65.65 $\pm$ 4.76
One-shot FSL	Hu et al. 2021	80.15 $\pm$ 0.78	80.15 $\pm$ 0.79	68.22 $\pm$ 2.05	68.12 $\pm$ 2.10
	Snell et al. 2017	80.75 $\pm$ 3.51	80.53 $\pm$ 3.66	83.29 $\pm$ 5.18	82.19 $\pm$ 5.92
	ICI (Wang et al. 2020)	81.69 $\pm$ 2.58	81.69 $\pm$ 2.59	84.40 $\pm$ 4.16	83.20 $\pm$ 4.77
	PLCM (Huang et al. 2021)	81.54 $\pm$ 1.79	81.49 $\pm$ 1.82	87.49 $\pm$ 3.62	86.62 $\pm$ 4.71
	ICI (Wang et al. 2020)	81.06 $\pm$ 2.07	80.99 $\pm$ 2.02	87.11 $\pm$ 4.07	86.71 $\pm$ 4.60
One-shot	PLCM (Huang et al. 2021)	82.02 $\pm$ 3.43	81.93 $\pm$ 3.49	88.76 $\pm$ 2.96	88.42 $\pm$ 3.35
SSFSL	Co-HSF (proposed)	84.41 $\pm$ 1.38	84.40 $\pm$ 1.38	90.10 $\pm$ 0.52	89.82 $\pm$ 1.86

Tab. 1. Quantitative comparison on PCam and NCT-CRC-HE (NCT). CLIP refers to OpenAI CLIP⁹

Methods

- We use CONCH's vision encoder as $G(\cdot)$ (see Fig.1).
- Randaugment⁶ augments labeled set X_{aug} for teacher model training and unlabeled set U_{aug} for co-filtering pseudo-labeling (see Fig. 2).
- Hyperparameters: T (for pseudo-label confidence threshold), α (for class imbalance mitigation), and step (number of added samples to X each iteration).

Datasets and Evaluation Setup

- Compare average inference time and GPU memory consumption on NVIDIA Tesla T4 and Intel Xeon 4120
- Datasets: (a) PCam⁷ is a binary class (control vs. tumor) breast tissue dataset, and (b) NCT-CRC-HE⁸ is a nine-class colon tissue dataset containing two different tumor-positive classes.

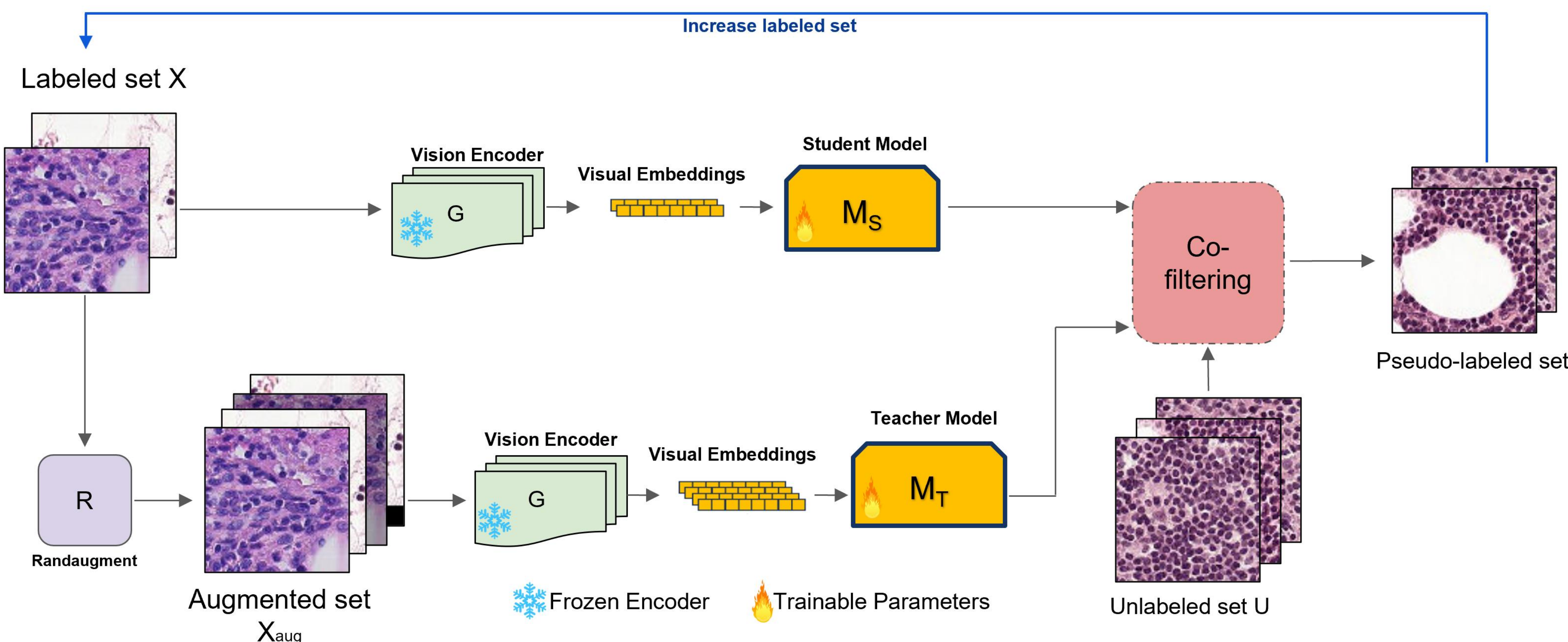


Fig 1. Overview of the proposed framework: the labeled set is augmented by Randaugment⁶ and featurized by any CLIP-based pre-trained vision encoder $G(\cdot)$. Both the student and teacher models are trained using the CONCH⁵ visual embeddings of X and X_{aug} respectively. The trained models and the unlabeled set are inputted to the Co-filtering algorithm (see Fig. 2), which selects samples and pseudo-labels to be added to X for the next iteration.

Co-filtering Pseudo-labeling Performance

Algorithm	Pseudo-label performance	
	Accuracy (%)	F-1 Score (%)
PLCM	88.04 $\pm$ 1.53	84.46 $\pm$ 1.74
ICI	89.33 $\pm$ 0.89	86.03 $\pm$ 1.69
Co-HSF (proposed)	99.96$\pm$0.10	99.20$\pm$1.60

Tab. 2. Quantitative comparison on pseudo-label performance for NCT-CRC-HE8

Algorithm	Pseudo-label performance	
	Accuracy (%)	F-1 Score (%)
PLCM	80.17 $\pm$ 2.56	80.01 $\pm$ 2.58
ICI	96.03 $\pm$ 7.94	96.04$\pm$7.92
Co-HSF (proposed)	99.05$\pm$0.54	92.60 $\pm$ 10.80

Tab. 3. Quantitative comparison on pseudo-label performance for PCam⁷

Resource Utilization

Algorithm	Inference Performance		
	Min	Acc (%)	VRAM(MiB)
CONCH	5.20 $\pm$ 0.01	51.70	2222
CLIPath	9.22 $\pm$ 0.07	69.24 $\pm$ 1.46	2222
Linear Probing	2.35$\pm$0.01	84.10 $\pm$ 1.38	952
ICI	11.86 $\pm$ 0.14	87.11 $\pm$ 4.07	952
PLCM	2.60 $\pm$ 0.01	88.76 $\pm$ 2.96	952
Co-HSF	2.44 $\pm$ 0.01	90.10$\pm$0.52	952

Tab. 4. Quantitative comparison on test set inference time in minutes (Min), accuracy (Acc), and model VRAM utilization on NCT-CRC-HE dataset

Conclusion

- SSFSL can fine-tune foundation models for histopathology classification with lower GPU utilization and higher accuracy.
- Co-HSF can further improve accuracy and reduce inference time when compared to linear probing, adapter-based, prompt-based, and SSFSL methods.
- Co-HSF is better suited for deployability due to its lower GPU utilization and inference times (see Tab. 4)
- In the future, we aim to test the proposed method over additional, diverse histopathology datasets, different evaluation tasks (e.g. detection), and discuss scenarios such as class-imbalanced unlabeled sets and selection criteria for one-shot labeled dataset.

References

- Lai et al. 2023. International Conference on Computer Vision Workshops.
- Zhou et al. 2022. Conference on Computer Vision and Pattern Recognition.
- Mirza et al. 2024. Advances in Neural Information Processing Systems.
- Javed et al. 2024. Conference on Computer Vision and Pattern Recognition.
- Lu et al. 2024. Nature Medicine.
- Cubuk et al. 2020. Conference on Computer Vision and Pattern Recognition Workshops.
- Veeling et al. 2018. Medical Image Computing and Computer Assisted Intervention.
- Kather et al. 2018. Zenodo.
- Radford et al. 2021. International Conference on Machine Learning.

Acknowledgments: This project was made possible by a grant from the National Institute on Aging (NIA) of the National Institutes of Health (NIH) under Award Number R01AG062517 and U24NS133949, Noyce Initiative UC Partnerships in Computational Transformation Program, NIH-National Institute On Aging awards #2R01-AG0652517, a 2023-2024 UC Davis Chancellor's Fellowship and Child Family Endowed Professorship Funds.