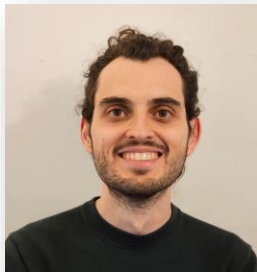


Semi-Supervised Few-Shot Adaptation of Vision-Language Models

Julio
Silva-Rodríguez



Ender
Konukoglu



Outline

A. Vision-Language Models (VLMs)

- Contrastive Language-Image Pre-training (CLIP).
- Zero-shot and few-shot inference.
- Medical vision-language models.

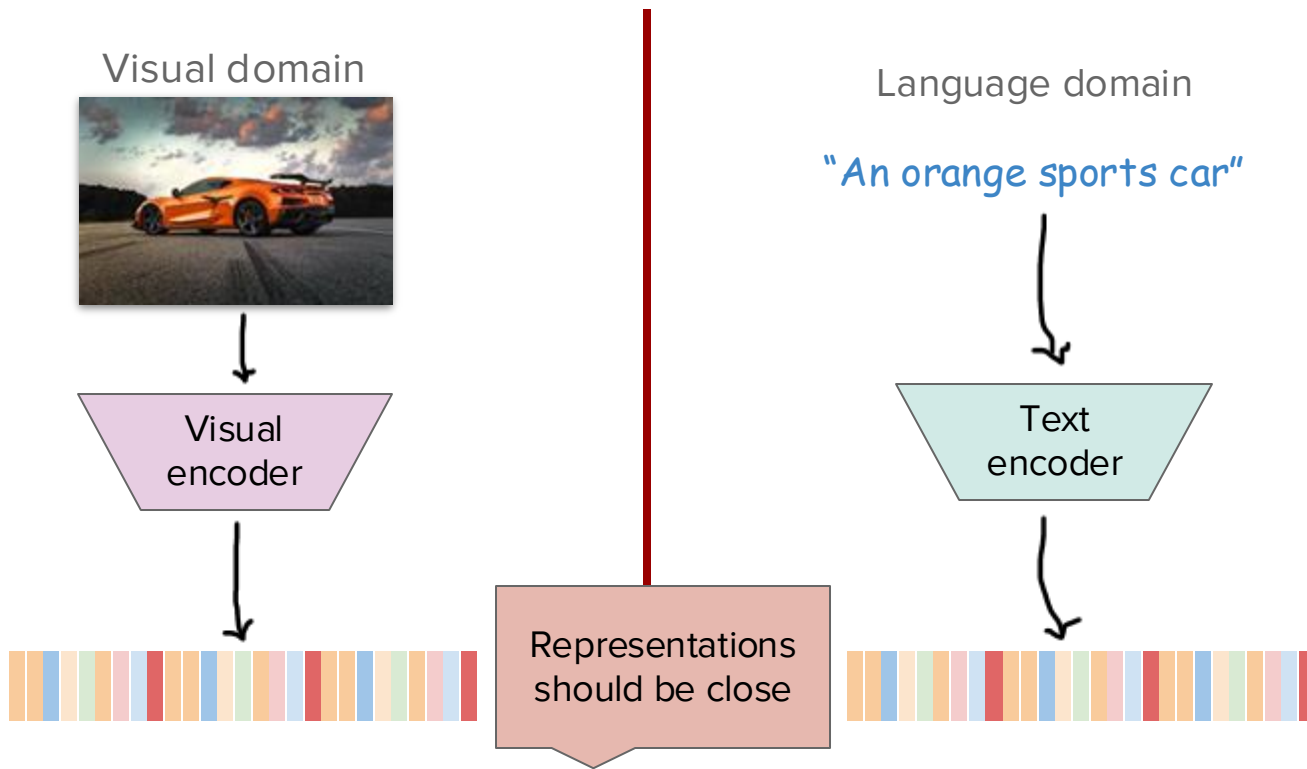
B. Few-shot adaptation

- Text-informed linear probes.

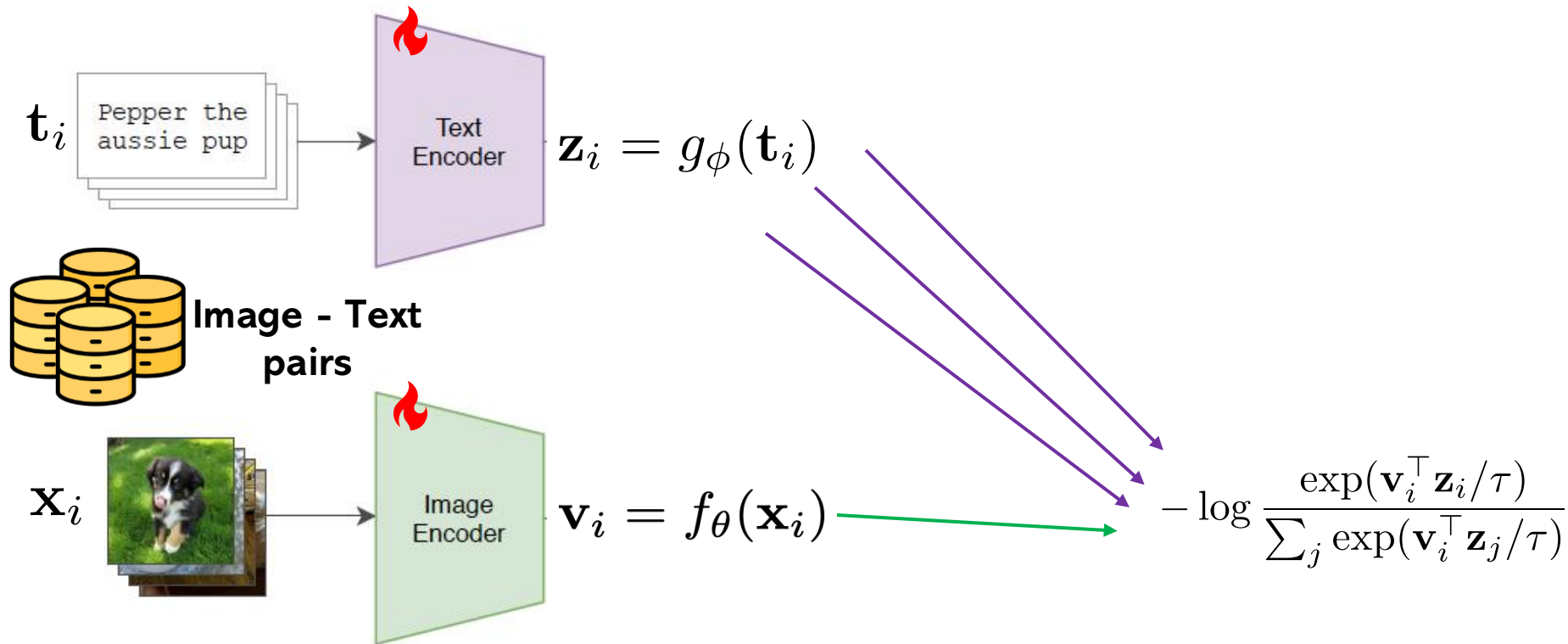
C. Semi-supervised few-shot adaptation

- Motivation.
- Learning from unlabeled data.
- SS-Text-U: objective and block-wise optimization.
- Results & exploratory analysis.

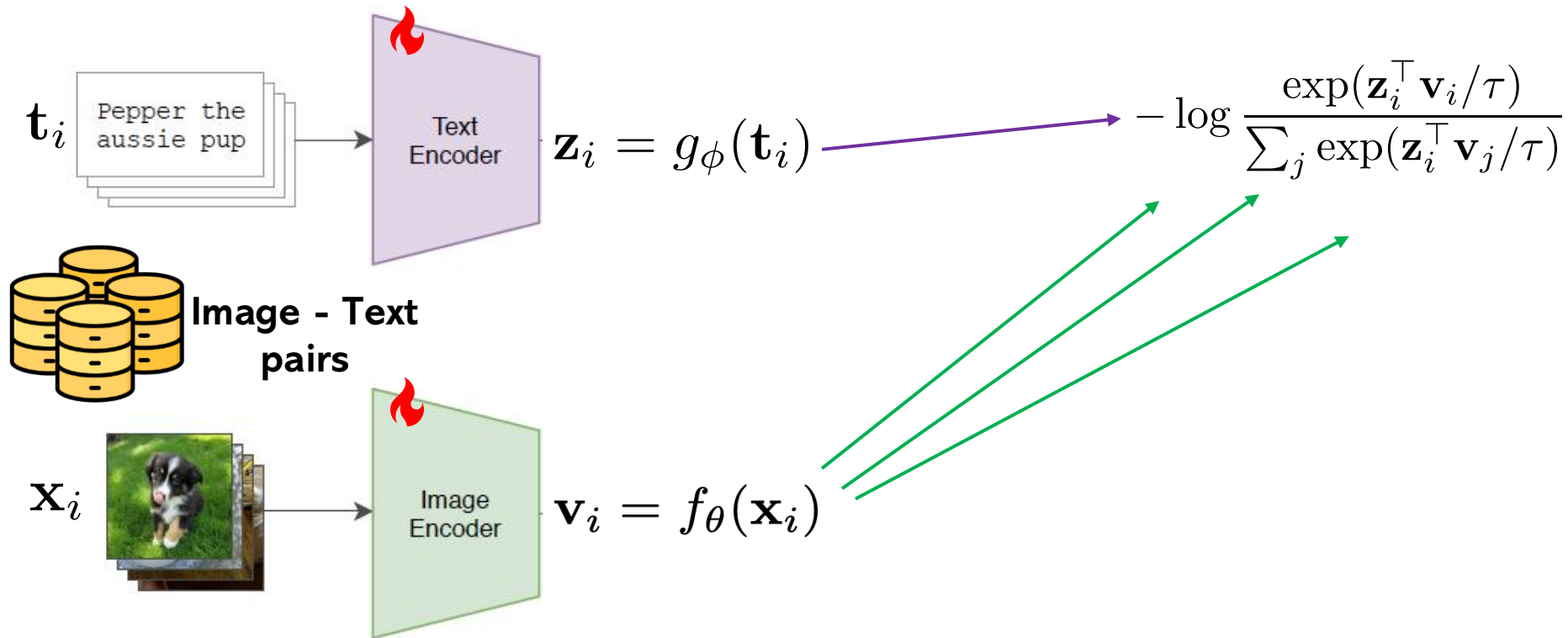
Embedding semantic knowledge at larger scale trough text supervision



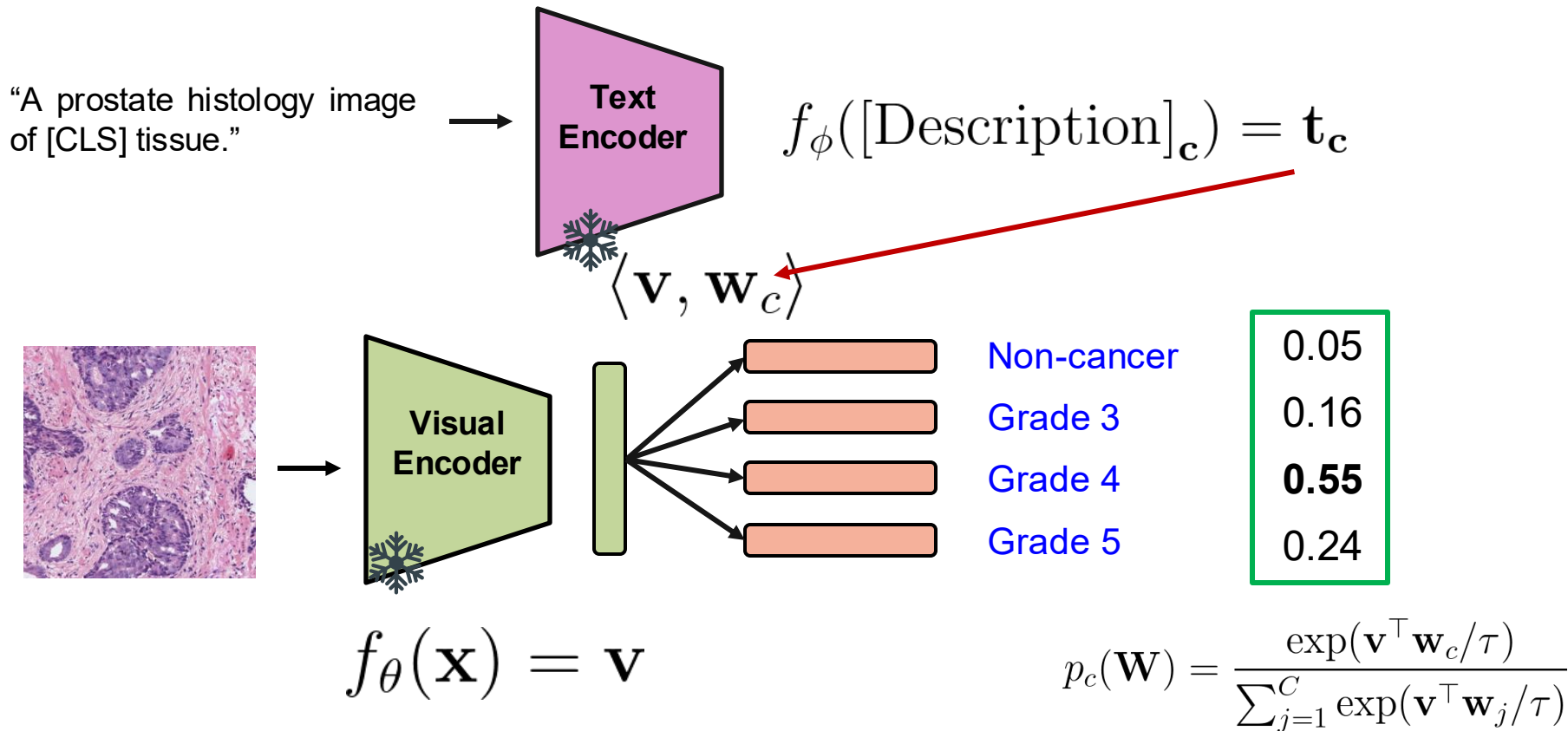
Contrastive Language-Image Pre-training (CLIP)



Contrastive Language-Image Pre-training (CLIP)



“Zero-shot” text-driven prediction for new tasks



Few-shot adaptation via linear probing

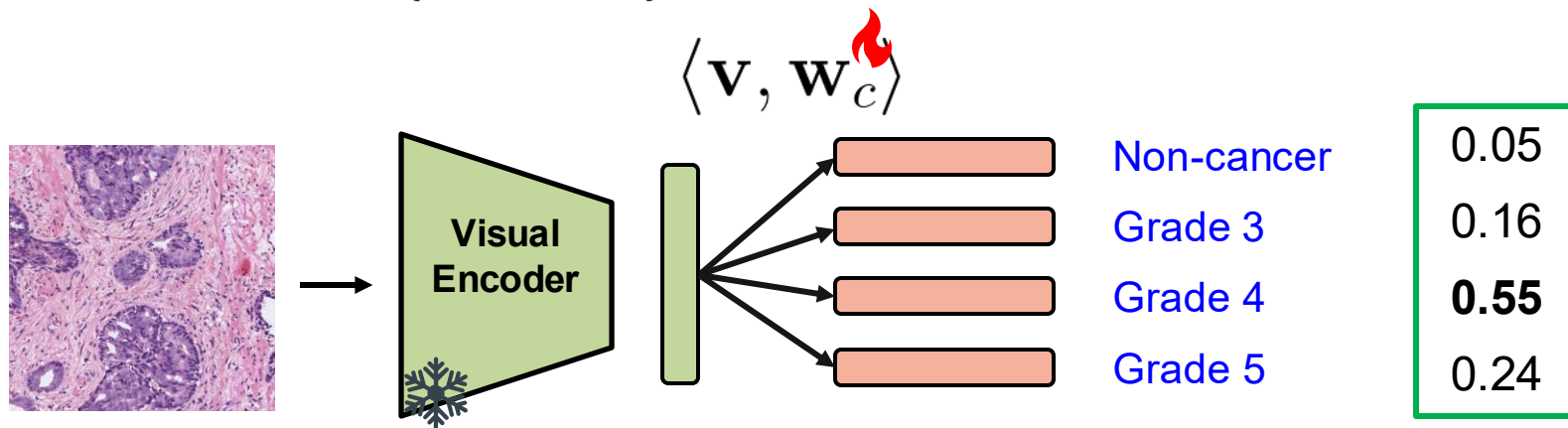
Support set

$$\{(\mathbf{x}_i, y_i)\}_{i=1}^N$$



$$N = C \times K; K \in \{1, 2, 4, 8, 16\}$$

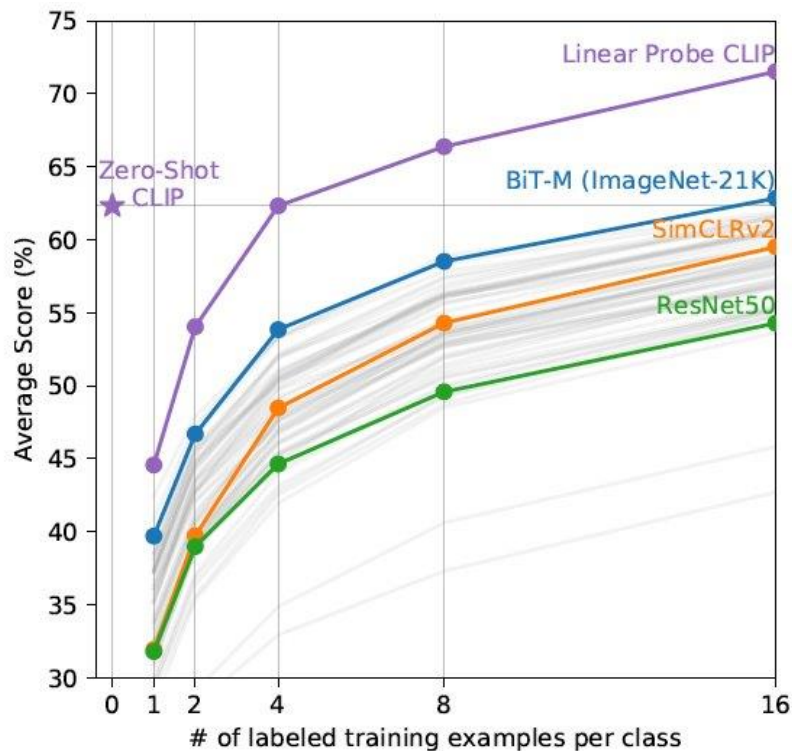
$$\min_{\mathbf{W}} \mathcal{L}(\mathbf{W}) = -\frac{1}{N} \sum_{i=1}^N \sum_{c=1}^C y_{ic} \ln(p_{ic}(\mathbf{W}))$$



$$f_{\theta}(\mathbf{x}) = \mathbf{v}$$

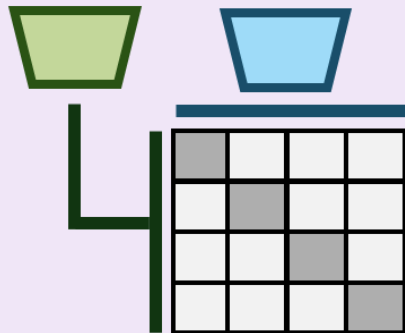
$$p_c(\mathbf{W}) = \frac{\exp(\mathbf{v}^{\top} \mathbf{w}_c / \tau)}{\sum_{j=1}^C \exp(\mathbf{v}^{\top} \mathbf{w}_j / \tau)}$$

Few-shot adaptation via linear probing



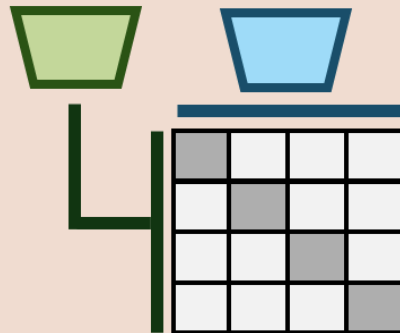
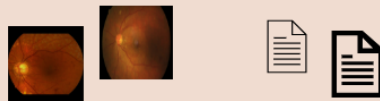
Medical Vision-Language Models

CONCH



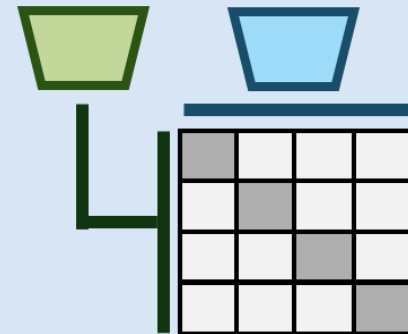
Histology

FLAIR



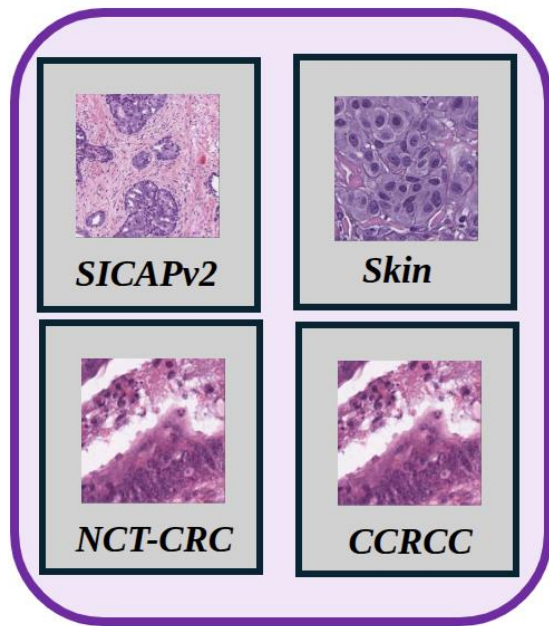
Ophtalmology

CONVIRT

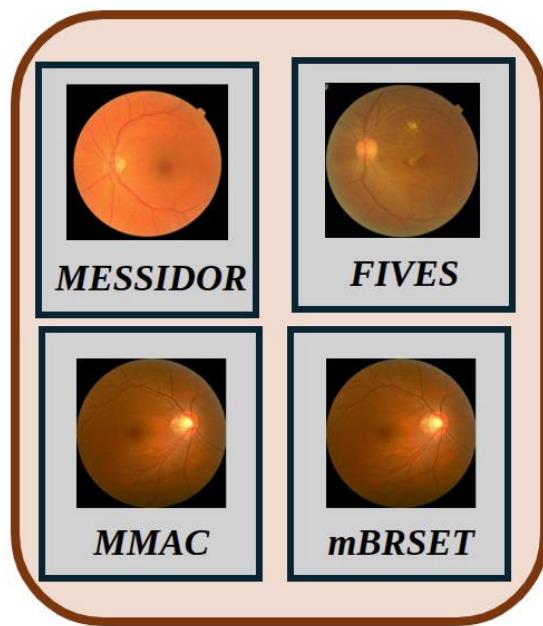


Chest X-ray

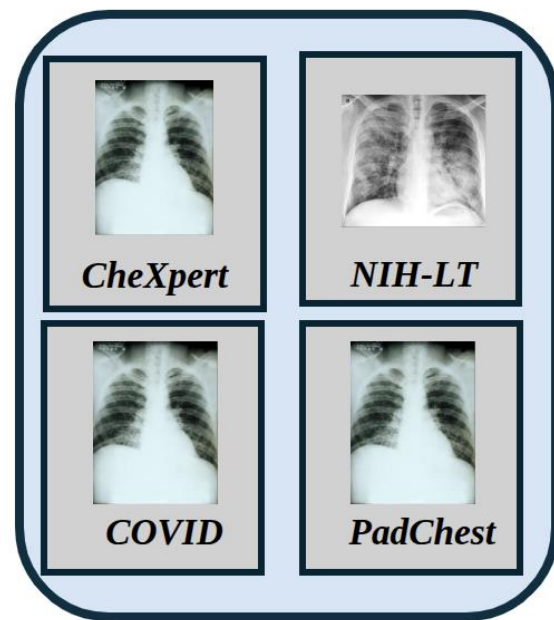
Medical Vision-Language Models



Histology



Ophthalmology



Chest X-ray

Text-informed linear probes.

Few-shot learning objective:

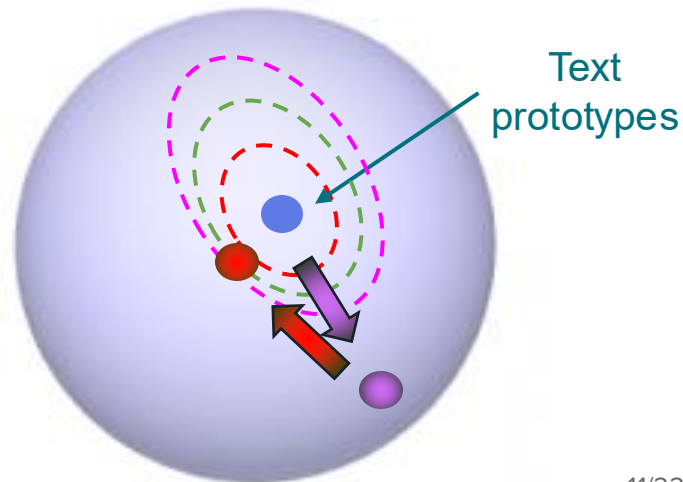
$$\min_{\mathbf{W}} \mathcal{L}(\mathbf{W}) = \frac{1}{|\mathcal{S}|} \sum_{i \in \mathcal{S}} \mathcal{H}(\mathbf{y}_i, \mathbf{p}_i(\mathbf{W})) + \sum_{c \in \mathcal{Y}} \lambda_c^T \|\mathbf{w}_c - \mathbf{t}_c\|_2^2$$

Cross-entropy loss:

$$\mathcal{H}(\mathbf{y}_i, \mathbf{p}_i(\mathbf{W})) = - \sum_{c \in \mathcal{Y}} y_{i,c} \ln(p_{i,c}(\mathbf{W}))$$

SoftMax similarity:

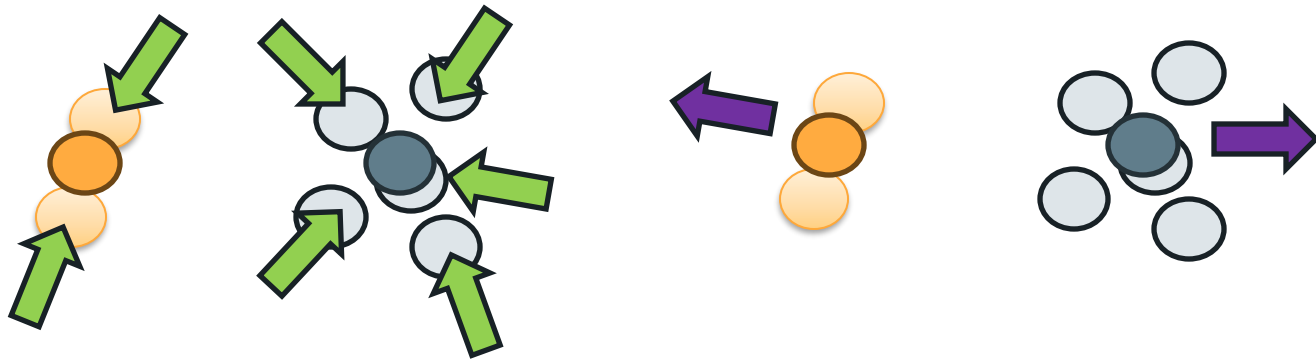
$$p_{i,c}(\mathbf{W}) = \frac{\exp(\mathbf{v}_i^\top \mathbf{w}_c / \tau)}{\sum_{j \in \mathcal{Y}} \exp(\mathbf{v}_i^\top \mathbf{w}_j / \tau)}$$



Text-informed linear probes.

A closer look to cross-entropy loss:

$$\mathcal{H}(\mathbf{y}_i, \mathbf{p}_i(\mathbf{W})) = \underbrace{- \sum_{c \in \mathcal{Y}} y_{i,c} (\mathbf{v}_i^\top \mathbf{w}_c / \tau)}_{\mathcal{H}^t(\mathbf{y}_i, \mathbf{W})} + \underbrace{\sum_{c \in \mathcal{Y}} y_{i,c} \ln \left(\sum_{j \in \mathcal{Y}} \exp(\mathbf{v}_i^\top \mathbf{w}_j / \tau) \right)}_{\mathcal{H}^{\text{cont}}(\mathbf{W})}$$



Text-informed linear probes.

Constrained linear probe approximation with closed-form solution:

$$\min_{\mathbf{W}} \mathcal{L}_{\text{FEW-SHOT}}(\mathbf{W}) = \frac{1}{|\mathcal{S}|} \sum_{i \in \mathcal{S}} \mathcal{H}^t(\mathbf{y}_i, \mathbf{W}) + \sum_{c \in \mathcal{Y}} \lambda_c^T \|\mathbf{w}_c - \mathbf{t}_c\|_2^2$$

Text-informed linear probes.

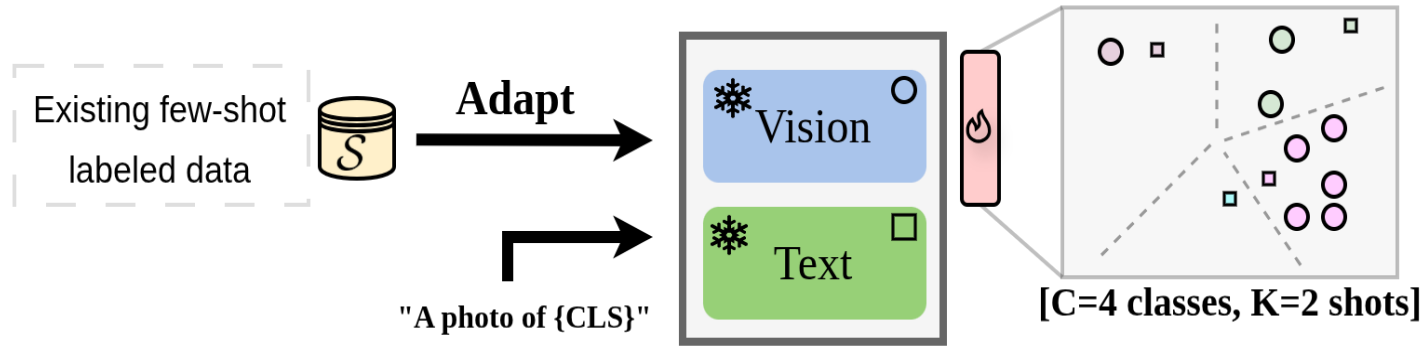
Constrained linear probe approximation with closed-form solution:

$$\min_{\mathbf{W}} \mathcal{L}_{\text{FEW-SHOT}}(\mathbf{W}) = -\frac{1}{|\mathcal{S}|} \sum_{c \in \mathcal{Y}} y_{i,c} (\mathbf{v}_i^\top \mathbf{w}_c / \tau) + \sum_{c \in \mathcal{Y}} \lambda_c^\top \|\mathbf{w}_c - \mathbf{t}_c\|_2^2$$

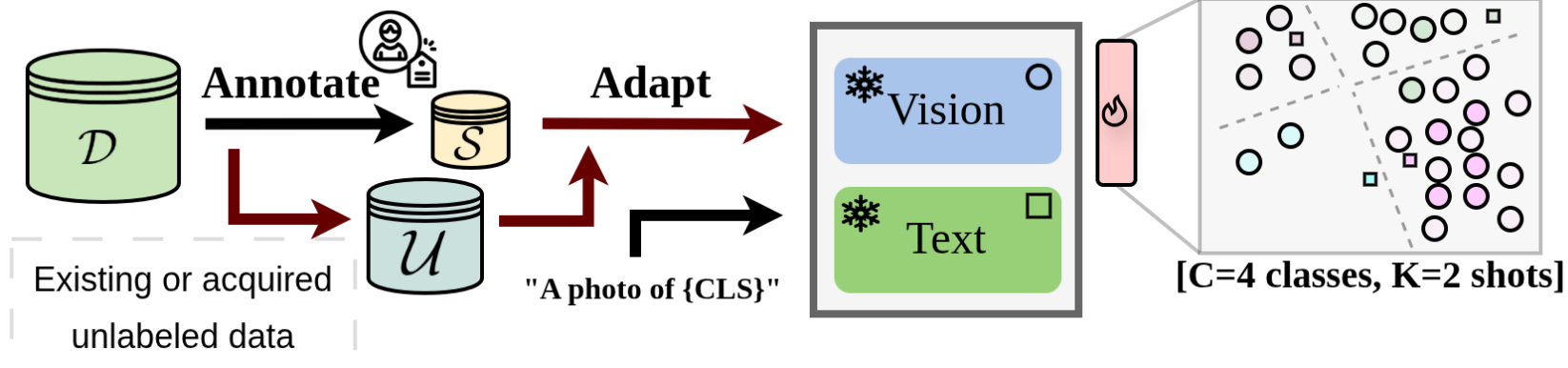
$$\mathbf{w}_c^t = \arg \min_{\mathbf{w}_c} \frac{\partial \mathcal{L}_{\text{FEW-SHOT}}(\mathbf{W})}{\partial \mathbf{w}_c} = \frac{1}{2\lambda_c^\top |\mathcal{S}| \tau} \sum_{i \in \mathcal{S}} y_{i,c} \mathbf{v}_i + \mathbf{t}_c$$


Visual Textual

Semi-supervision: making the most of unlabelled data



Semi-supervision: making the most of unlabelled data



Learning from unlabelled data

Combining a few-shot support set with unlabeled data:

$$\mathcal{D} = \{i \in \mathbb{N} : 1 \leq i \leq N + M\} = \mathcal{S} \cup \mathcal{U}$$


$\{(\mathbf{v}_i, \mathbf{y}_i)\}_{i \in \mathcal{S}} \quad \{(\mathbf{v}_i)\}_{i \in \mathcal{U}}$

Learning from unlabelled data

Incorporating supervisory signals through unlabeled data:

$$\mathcal{L}_U(\mathbf{W}, \mathbf{z}) = \frac{1}{|\mathcal{U}|} \sum_{i \in \mathcal{U}} \mathcal{H}^t(\mathbf{z}_i, \mathbf{W}) \text{ s.t. } \hat{\mathbf{m}} = \mathbf{m}$$


pseudo-labels
(treated as learnable variables)

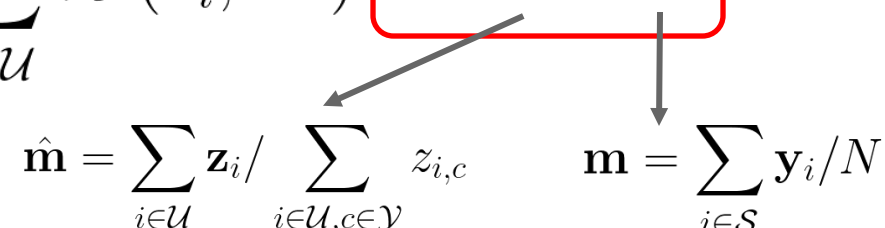
$$\mathbf{z}_i = (z_{i,k})_{(1 \leq k \leq C)} \in \Delta_C, i \in \mathcal{U}$$

Learning from unlabelled data

Incorporating supervisory signals through unlabeled data:

$$\mathcal{L}_U(\mathbf{W}, \mathbf{z}) = \frac{1}{|\mathcal{U}|} \sum_{i \in \mathcal{U}} \mathcal{H}^t(\mathbf{z}_i, \mathbf{W}) \quad \text{s.t.} \quad \hat{\mathbf{m}} = \mathbf{m}$$

$\hat{\mathbf{m}} = \sum_{i \in \mathcal{U}} \mathbf{z}_i / \sum_{i \in \mathcal{U}, c \in \mathcal{Y}} z_{i,c} \qquad \mathbf{m} = \sum_{i \in \mathcal{S}} \mathbf{y}_i / N$



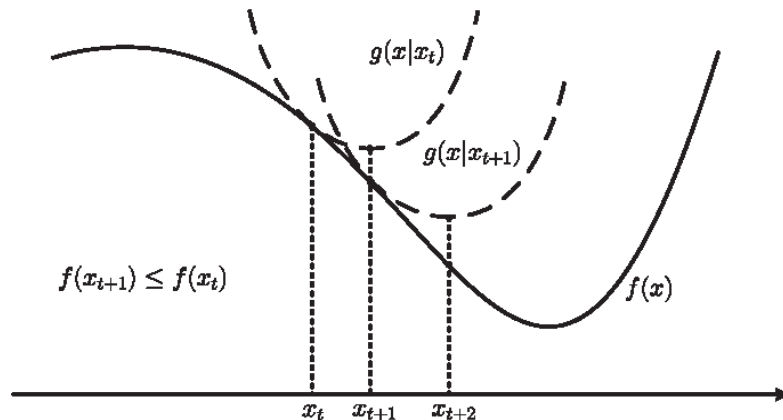
SS-Text-U: Objective & block-wise optimization

$$\min_{\mathbf{W}, \mathbf{z}} \mathcal{L}_{\text{SEMI}}(\mathbf{W}, \mathbf{z}) = \mathcal{L}_{\text{FEW-SHOT}}(\mathbf{W}) + \lambda^U \mathcal{L}_U(\mathbf{W}, \mathbf{z})$$

Block-wise optimizer.

For t in T :

1. Freeze $\mathbf{W}$, argmin $\mathbf{Z}$ wrt $\mathbf{Z}$ -block.
2. Freeze $\mathbf{Z}$, argmin $\mathbf{W}$ wrt $\mathbf{W}$ -block.



Lange et al. Optimization transfer using surrogate objective functions (2000).

Razaviyayn et al. A unified convergence analysis of block successive minimization methods for non-smooth optimization (2013).

Wright et al. Coordinate descent algorithms (2015).

Figure from Sun et al. Majorization-Minimization Algorithms in Signal Processing, Communications, and Machine Learning (2017).

SS-Text-U: Objective & block-wise optimization

1. Z-block updates.

$$\min_{\mathbf{z}} \mathcal{L}_{\mathcal{U}}(\mathbf{W}^{t-1}, \mathbf{z}) = -\frac{1}{|\mathcal{U}|} \sum_{i \in \mathcal{U}} \sum_{c \in \mathcal{Y}} z_{i,c} (\mathbf{v}_i^{\top} \mathbf{w}_c^{t-1} / \tau) \text{ s.t. } \hat{\mathbf{m}} = \mathbf{m}$$

$(\mathbf{w}_c^0 = \mathbf{t}_c)$



SS-Text-U: Objective & block-wise optimization

1. Z-block updates.

$$\min_{\mathbf{z}} \mathcal{L}_{\mathcal{U}}(\mathbf{W}^{t-1}, \mathbf{z}) = -\frac{1}{|\mathcal{U}|} \sum_{i \in \mathcal{U}} \sum_{c \in \mathcal{Y}} z_{i,c} (\mathbf{v}_i^{\top} \mathbf{w}_c^{t-1} / \tau) \text{ s.t. } \hat{\mathbf{m}} = \mathbf{m}$$

(in the matrix form)

$$(\mathbf{w}_c^0 = \mathbf{t}_c)$$

$$\mathbf{S} = ((\mathbf{v}_i^{\top} \mathbf{w}_c^{t-1} / \tau)_{c \in \mathcal{Y}})_{i \in \mathcal{U}} \quad \mathbf{Q} = (\mathbf{z}_i)_{i \in \mathcal{U}}$$

SS-Text-U: Objective & block-wise optimization

1. Z-block updates.

$$\min_{\mathbf{z}} \mathcal{L}_{\mathcal{U}}(\mathbf{W}^{t-1}, \mathbf{z}) = -\frac{1}{|\mathcal{U}|} \sum_{i \in \mathcal{U}} \sum_{c \in \mathcal{Y}} z_{i,c} (\mathbf{v}_i^{\top} \mathbf{w}_c^{t-1} / \tau) \text{ s.t. } \hat{\mathbf{m}} = \mathbf{m}$$

(in the matrix form)

$$(\mathbf{w}_c^0 = \mathbf{t}_c)$$

$$\mathbf{S} = ((\mathbf{v}_i^{\top} \mathbf{w}_c^{t-1} / \tau)_{c \in \mathcal{Y}})_{i \in \mathcal{U}} \quad \mathbf{Q} = (\mathbf{z}_i)_{i \in \mathcal{U}}$$

The optimization can be addressed from an *Optimal Transport* perspective:

$$\max_{\mathbf{Q} \in \mathcal{Q}} \text{tr}(\mathbf{Q}^{\top} \mathbf{S}), \text{ s.t. } \hat{\mathbf{m}} = \mathbf{m}$$

SS-Text-U: Objective & block-wise optimization

Iterative solution via the Sinkhorn-Knopp algorithm.

$$\mathbf{Q}^* = \text{Diag}(\mathbf{r}^{(T_{\text{OT}})}) \mathbf{Q}^{(0)} \text{Diag}(\mathbf{c}^{(T_{\text{OT}})})$$

$$\mathbf{r}^{(t_{\text{OT}})} = \mathbf{m} / (\mathbf{Q}^{(0)} \mathbf{c}^{(t_{\text{OT}}-1)}) \quad \mathbf{c}^{(t_{\text{OT}})} = \mathbf{u}_{(|\mathcal{U}|)} / (\mathbf{Q}^{(0)} \mathbf{r}^{(t_{\text{OT}})})$$

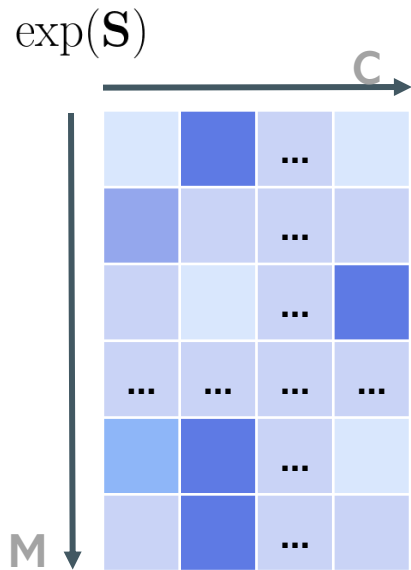
$$\mathbf{Q}^{(0)} = \exp(\mathbf{S}) / \sum (\exp(\mathbf{S})) \quad (\text{batch-level normalized logits for initialization})$$

$$\mathbf{c}^{(0)} = \mathbf{1}_{(|\mathcal{U}|)} \quad (\text{normalization initialization – initially, all samples weight the same})$$

$$\mathbf{u}_{(|\mathcal{U}|)} \quad (\text{target uniform sample-marginal distribution})$$

SS-Text-U: Objective & block-wise optimization

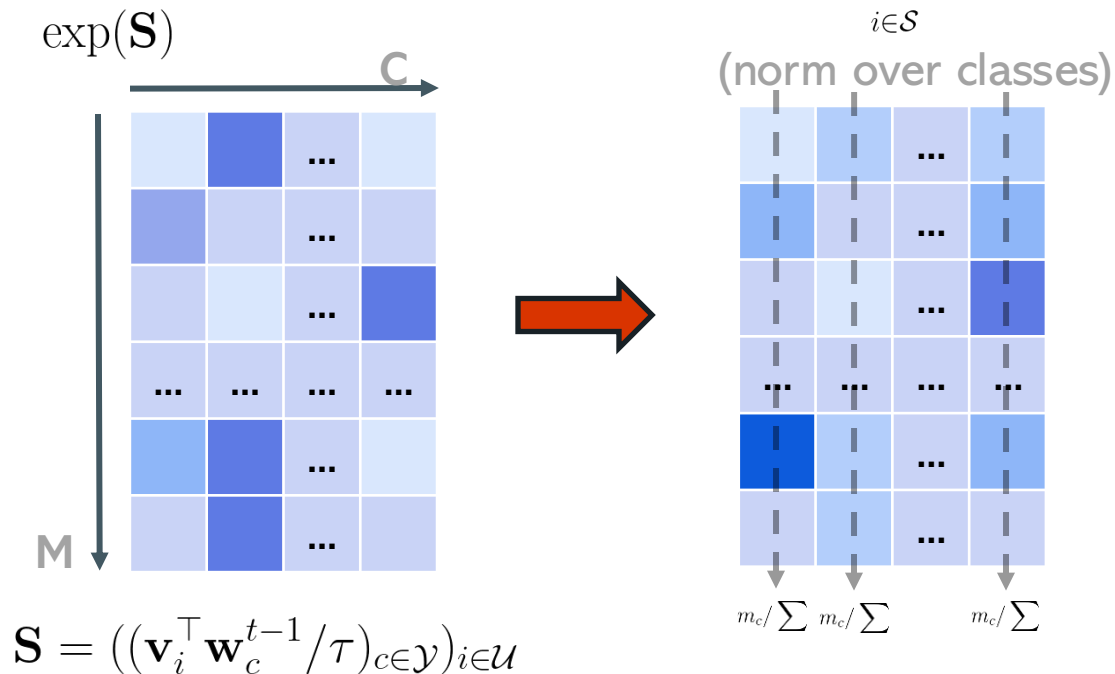
Sinkhorn-Knopp algorithm.



$$\mathbf{S} = ((\mathbf{v}_i^\top \mathbf{w}_c^{t-1} / \tau)_{c \in \mathcal{Y}})_{i \in \mathcal{U}}$$

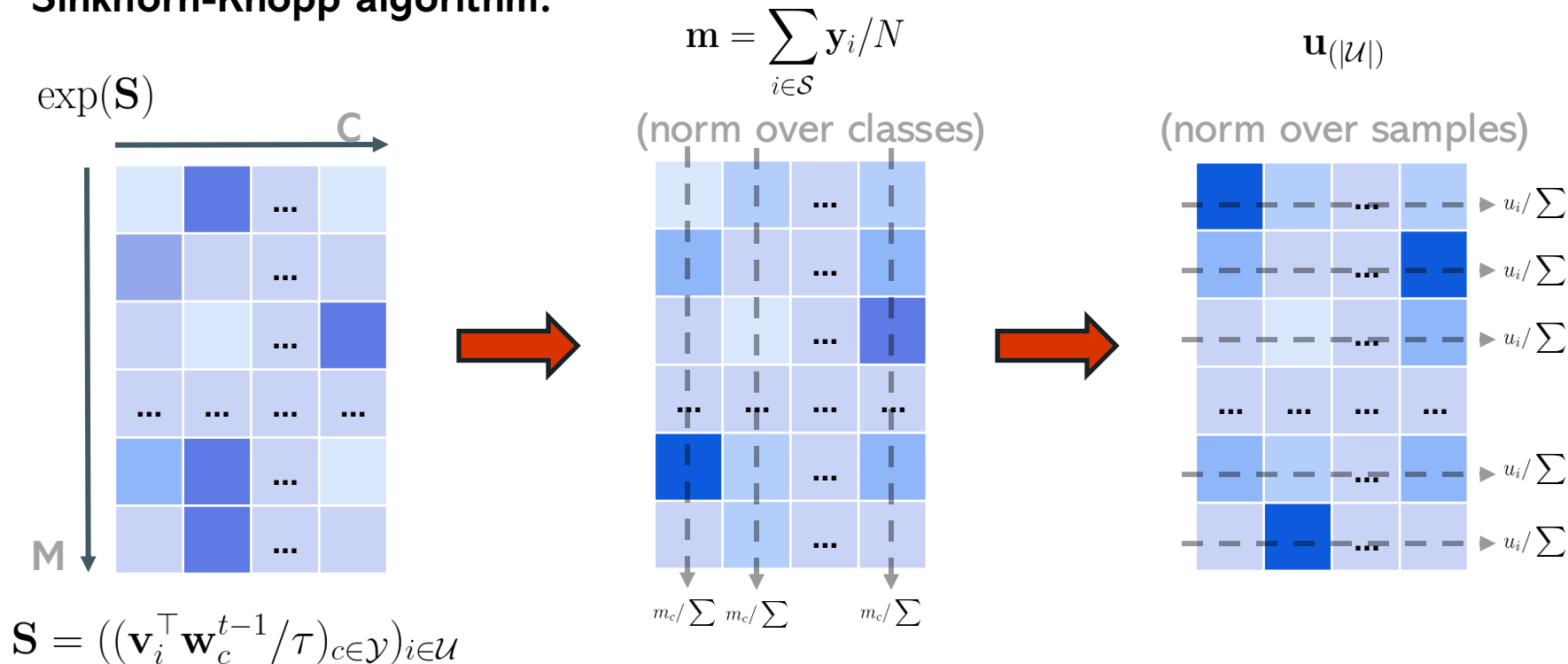
SS-Text-U: Objective & block-wise optimization

Sinkhorn-Knopp algorithm.



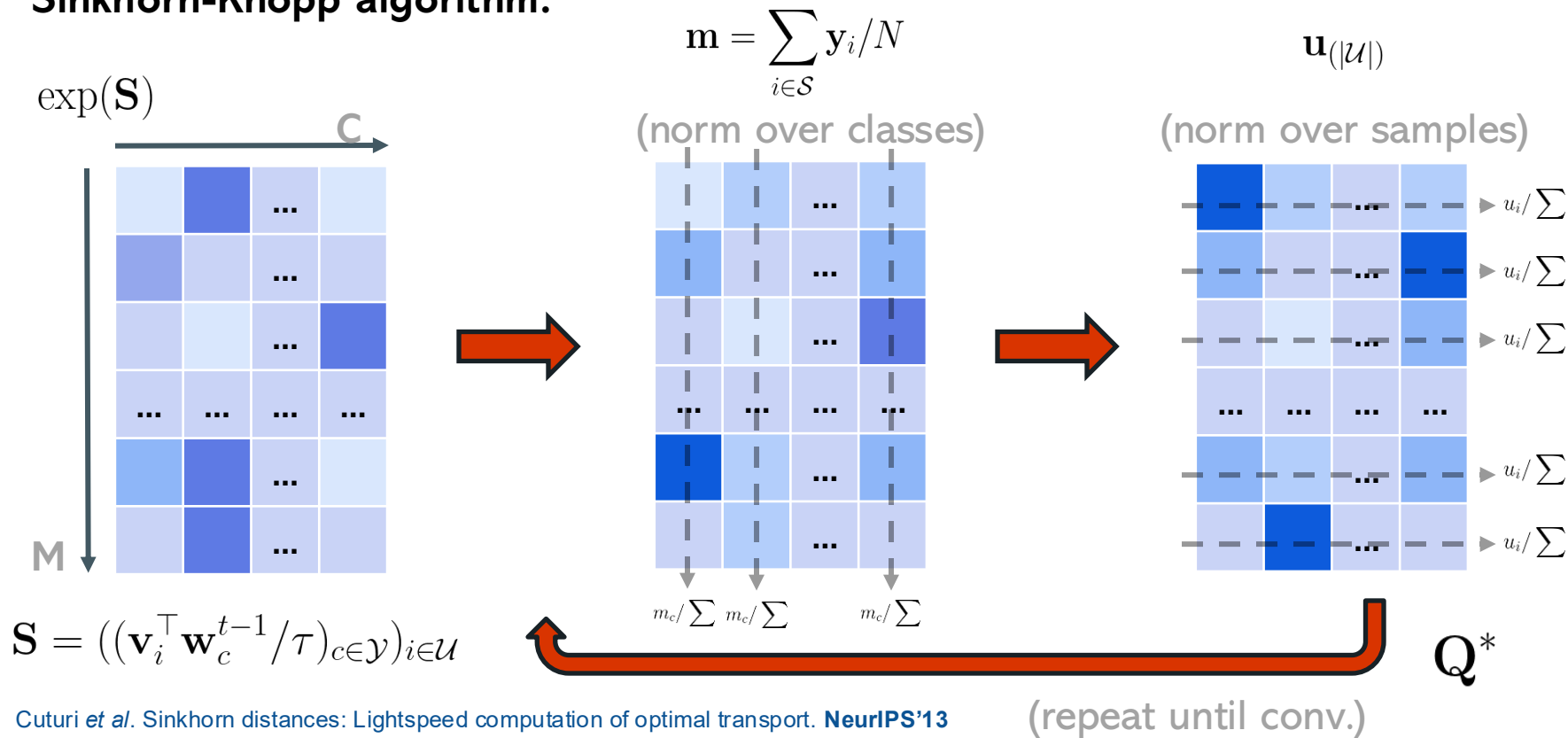
SS-Text-U: Objective & block-wise optimization

Sinkhorn-Knopp algorithm.



SS-Text-U: Objective & block-wise optimization

Sinkhorn-Knopp algorithm.



SS-Text-U: Objective & block-wise optimization

2. Closed-form solution for the W-block updates.

$$\mathbf{w}_c^t = \arg \min_{\mathbf{w}_c} \frac{\partial \mathcal{L}_{\text{SEMI}}(\mathbf{W}, \mathbf{z}^t)}{\partial \mathbf{w}_c} = \frac{1}{2\lambda_c^T |\mathcal{S}| \tau} \sum_{i \in \mathcal{S}} y_{i,c} \mathbf{v}_i + \frac{\lambda_c^U}{2\lambda_c^T |\mathcal{U}| \tau} \sum_{i \in \mathcal{U}} z_{i,c}^t \mathbf{v}_i + \mathbf{t}_c$$

SS-Text-U: Objective & block-wise optimization

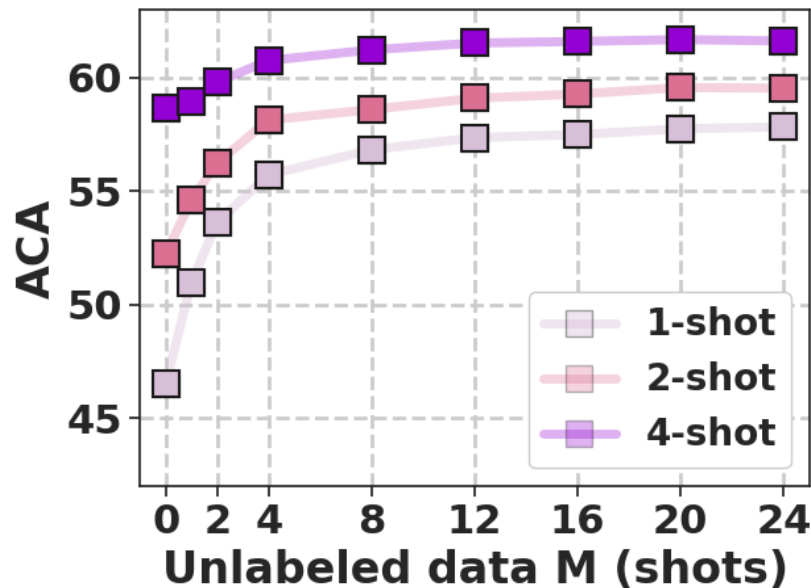
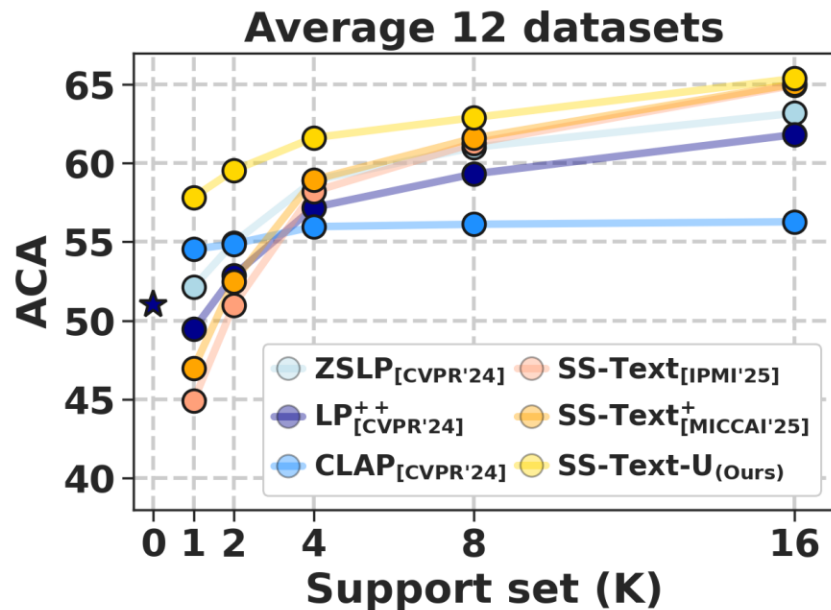
2. Closed-form solution for the W-block updates.

$$\mathbf{w}_c^t = \arg \min_{\mathbf{w}_c} \frac{\partial \mathcal{L}_{\text{SEMI}}(\mathbf{W}, \mathbf{z}^t)}{\partial \mathbf{w}_c} = \frac{1}{2\lambda_c^T |\mathcal{S}| \tau} \sum_{i \in \mathcal{S}} y_{i,c} \mathbf{v}_i + \frac{\cancel{\lambda_c^U}}{2\lambda_c^T |\mathcal{U}| \tau} \sum_{i \in \mathcal{U}} z_{i,c}^t \mathbf{v}_i + \mathbf{t}_c$$

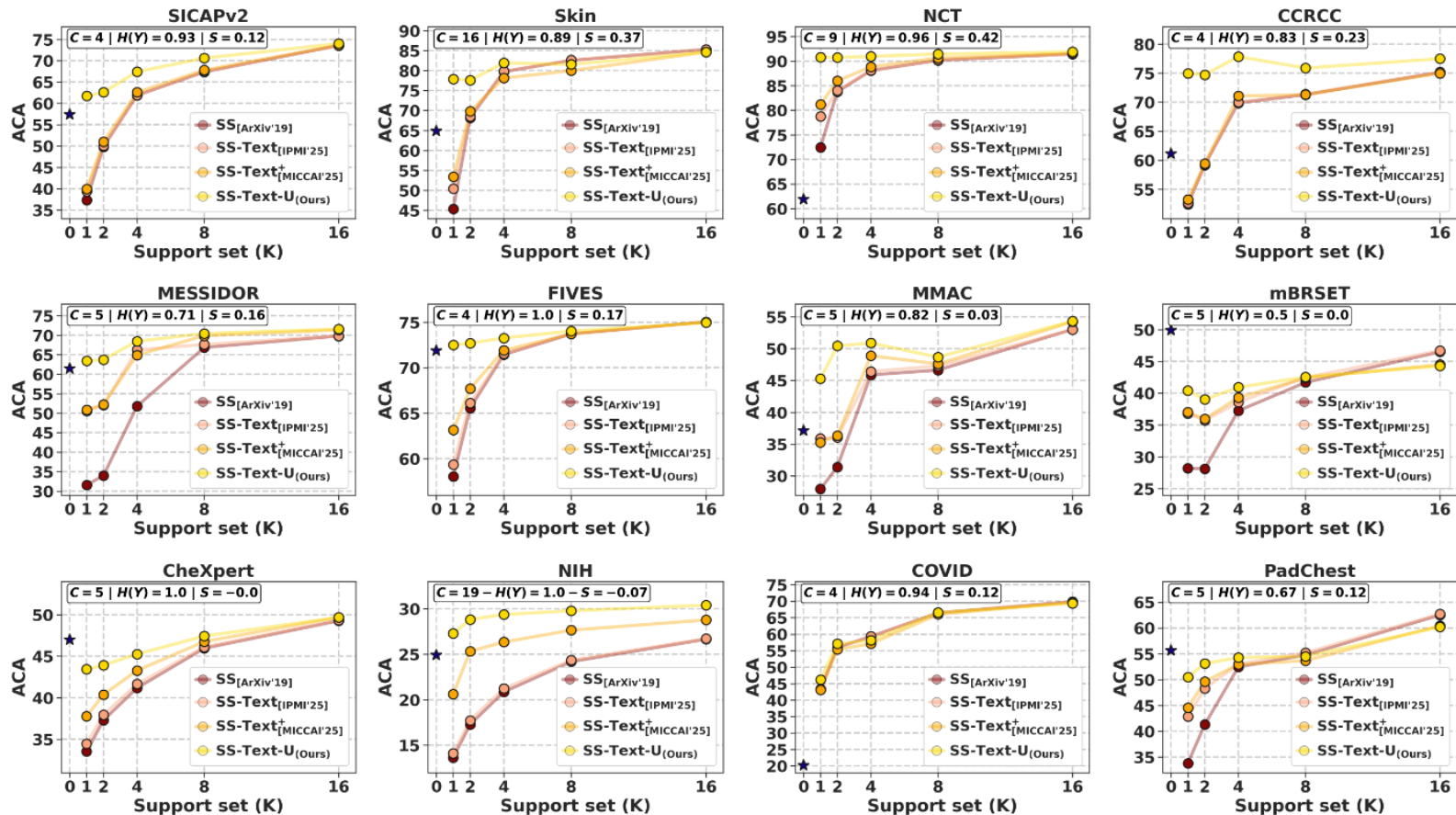
(data-driven hyper-parameters)

$$\lambda_c^T = (1/K_c) = (1/\sum_{i \in \mathcal{S}} y_{i,c}) ; \quad \cancel{\lambda_c^U} = 2\lambda_c^T$$

Results & exploratory analysis



Results & exploratory analysis



Results & exploratory analysis

