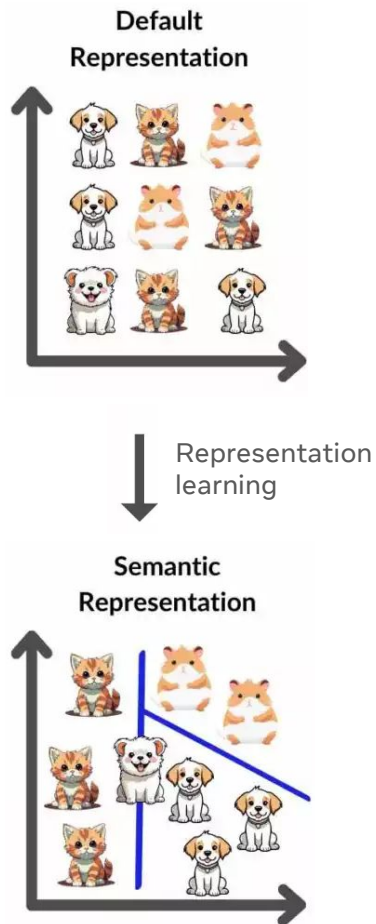


Robust image retrieval with deep learning.

Scientific seminar @ Orasis 2025 – Elias Ramzi
June 13, 2025



le cnam
cedric EA4629

COEXYA
Connect skills, create more

valeo.ai

Summary.

Robust image retrieval:

-  ROADMAP: optimization of rank losses.
-  HAPPIER: Hierarchical Image Retrieval for Robust Ranking.

Few shot robustness in the era of foundation models:

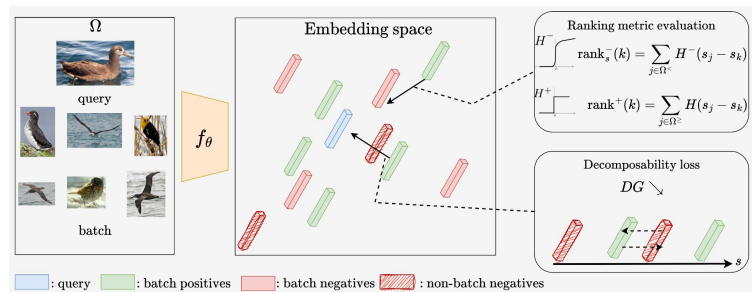
-  HEAT: Post-hoc out-of-distribution detection.
-  GalLoP: Few shot adaptation of vision-language models.

Summary and contributions.



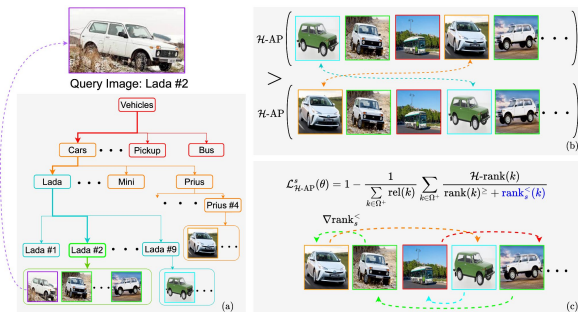
ROADMAP:

Optimization of Ranking Losses for Image retrieval.



HAPPIER:

Hierarchical Image Retrieval for Robust Ranking.



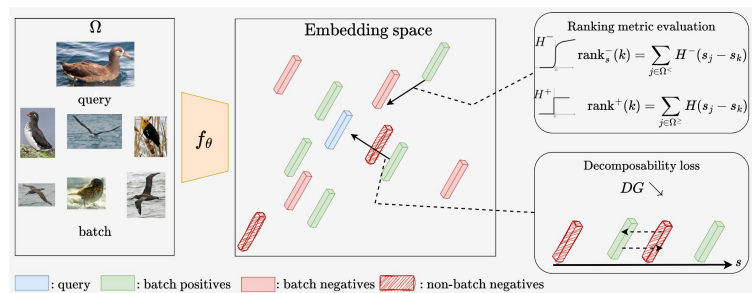
- Ramzi, Elias, et al. "Robust and Decomposable Average Precision for Image Retrieval." *NeurIPS*, 2021.
- Ramzi, Elias, et al. "Hierarchical Average Precision Training for Pertinent Image Retrieval." *ECCV*, 2022.
- Ramzi, Elias, et al. "Optimization of Rank Losses for Image Retrieval." *TPAMI*, 2025.

Summary and contributions.



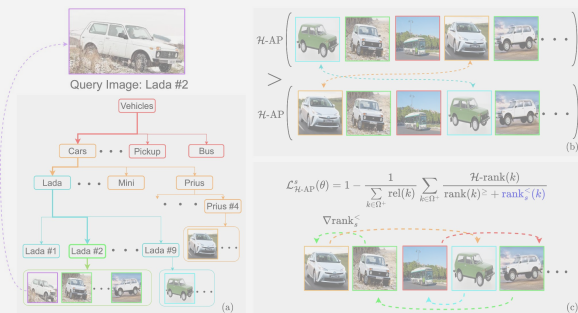
ROADMAP:

Optimization of Ranking Losses for Image retrieval.



HAPPIER:

Hierarchical Image Retrieval for Robust Ranking.



■ Ramzi, Elias, et al. “Robust and Decomposable Average Precision for Image Retrieval.” *NeurIPS*, 2021.

■ Ramzi, Elias, et al. “Optimization of Rank Losses for Image Retrieval.” *TPAMI*, 2025.

Image retrieval.

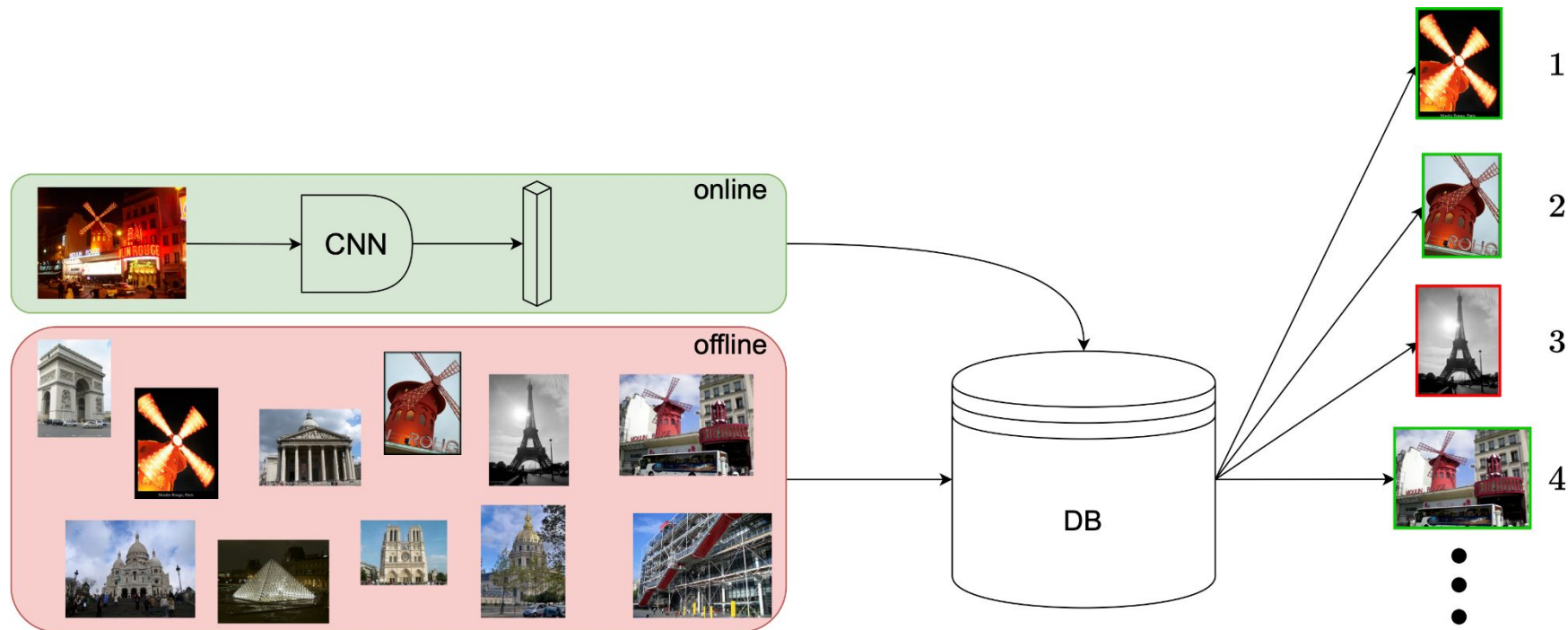
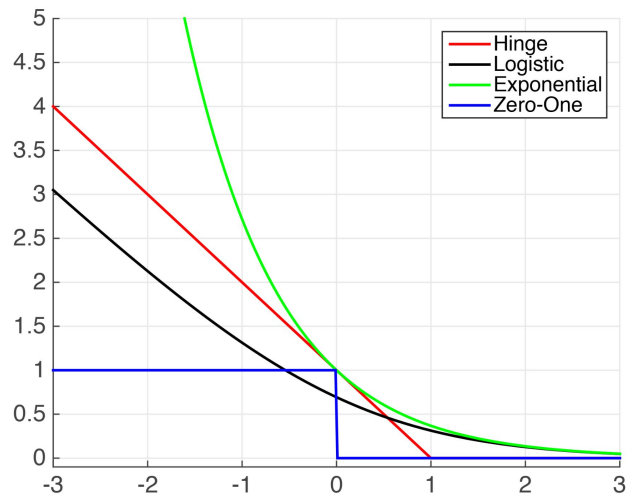
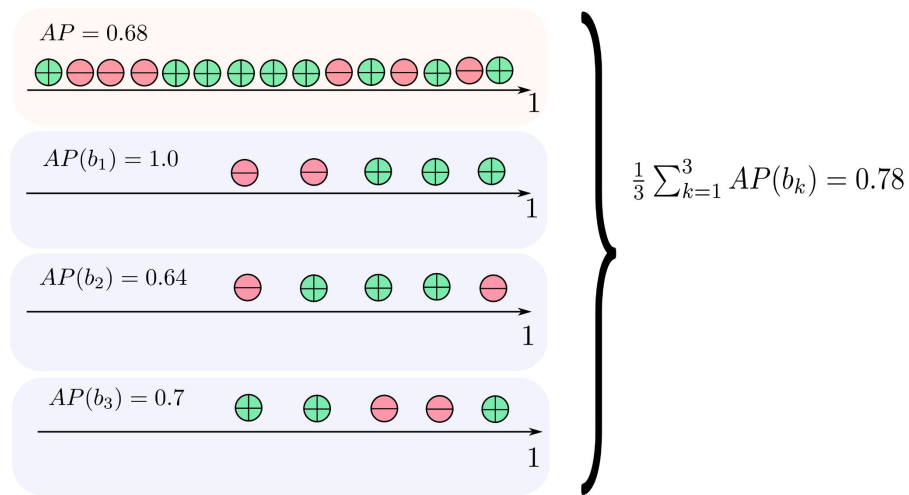


Image retrieval → retrieve similar images to the query image.

Optimizing rank losses.



Evaluation metrics are not differentiable → can not be used directly for SGD.



Evaluation metrics are non-decomposable → shortcomings for mini-batch training.



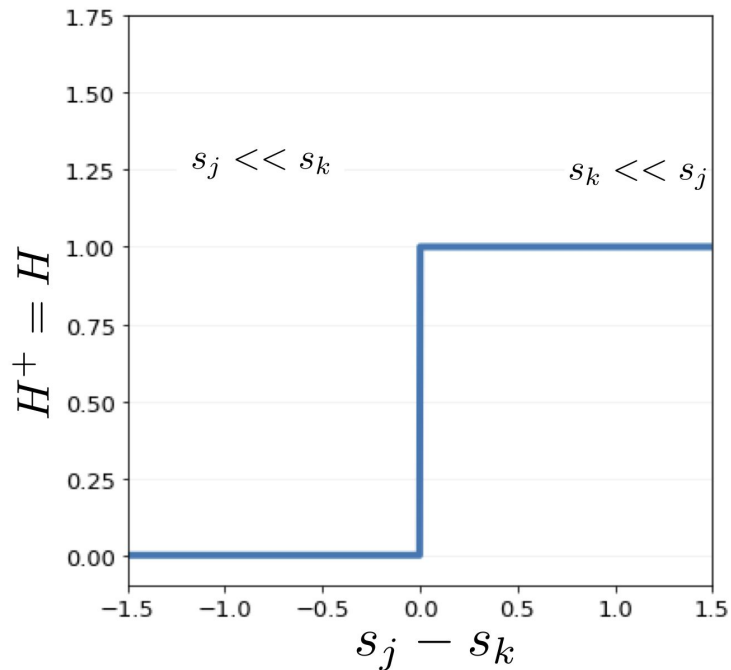
Evaluation metrics in image retrieval → rank-based:

$$\text{AP} = \frac{1}{|\Omega^+|} \sum_{k \in \Omega^+} \frac{\text{rank}^+(k)}{\text{rank}(k)}$$

$$\begin{aligned} \text{Recall@n} &= \frac{\# \text{ number of positive before } n}{|\Omega^+|} \\ &= \frac{1}{\min(|\Omega^+|, n)} \sum_{k \in \Omega^+} H(n - \text{rank}(k)) \end{aligned}$$

Definition of the rank with Heaviside functions:

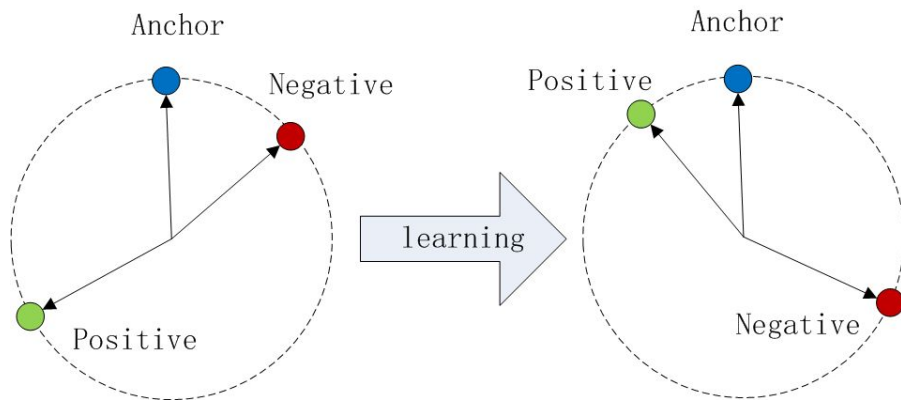
$$\begin{aligned} \text{rank}(k) &= \# \text{ number of similarities greater than } s_k \\ &= 1 + \sum_{j \in \Omega} H(s_j - s_k) \end{aligned}$$



s_k : cosine similarity between image k and the query.

Surrogate losses: approximation of the metric. 🌐 ROADMAP

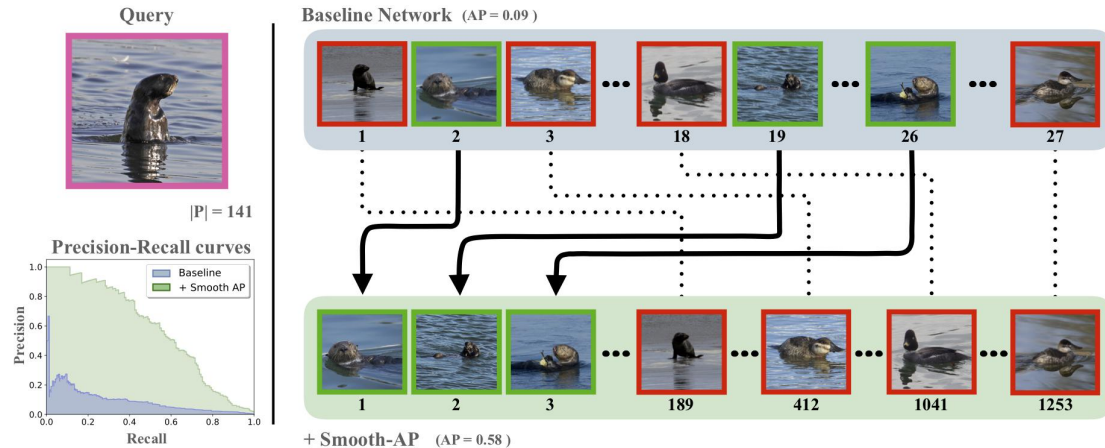
- Coarse upper-bound of the metrics.
- Not well-aligned with evaluation metrics: supports bottom vs. top of the ranking



Surrogate losses:

- Triplet loss – Wu, Chao-Yuan, et al. *ICCV*, 2017.
- NSM: Zhai, Andrew, et al. *BMVC*, 2018.
- Multi-similarity loss – Wang, Xun, et al. *CVPR*, 2019.

Surrogate losses: approximation of the rank.



- No strong theoretical guarantees: not an upper bound.
- Ill-behaved gradient flow.

Long line of work:

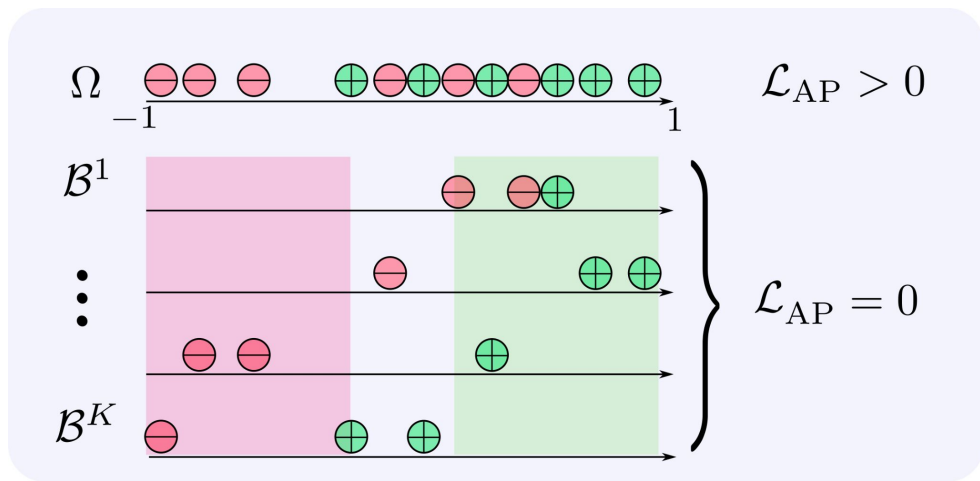
- Learned Ranking – SoDeep: Engilberge, Martin, et al. *CVPR*, 2019.
- Histogram binning – SoftBin: Revaud, Jerome, et al. *ICCV*, 2019.
- Blackbox optimization of AP: Rolínek, Michal, et al. *CVPR*, 2020.
- Smooth approximations of AP: Brown, Andrew, et al. *ECCV*, 2020.

Non-decomposability.

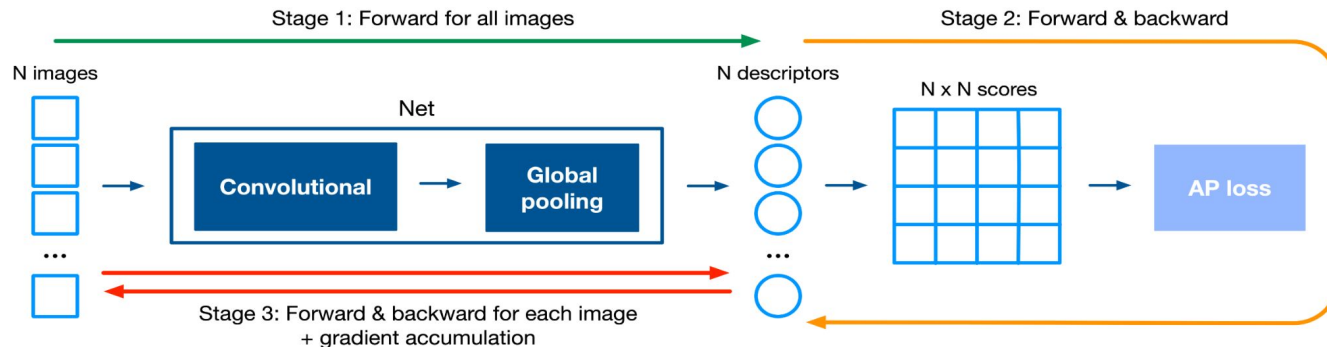
Ranking-based metrics are not decomposable:

$$\mathcal{M}(\Omega) \neq \frac{1}{|\mathcal{B}|} \sum_{b \in \mathcal{B}} \mathcal{M}(b)$$

- Loss on batches are biased.
- Training can stop before global loss is zero.



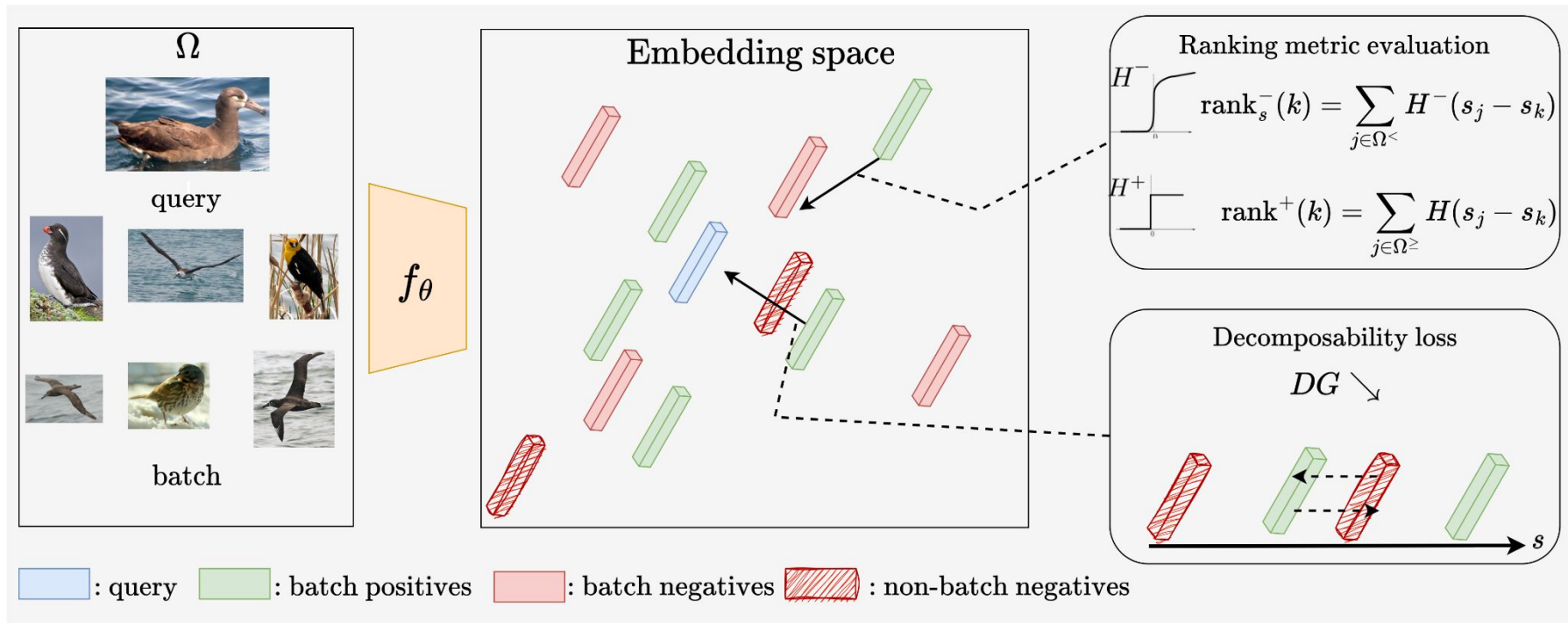
Addressing non-decomposability.



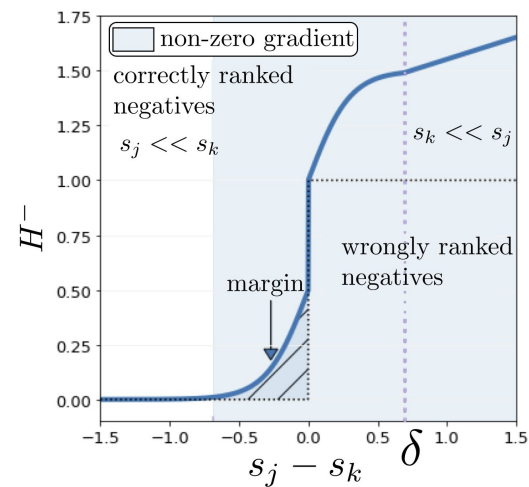
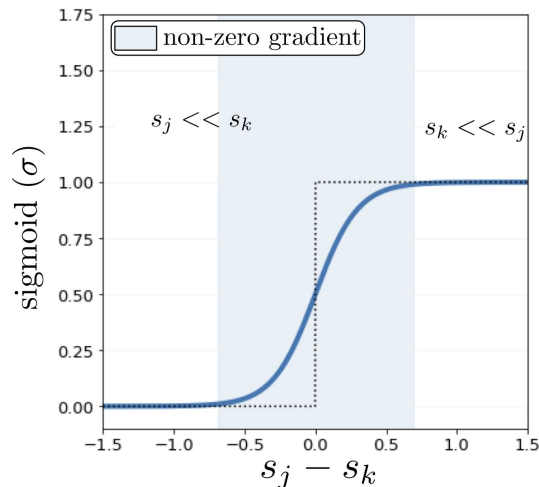
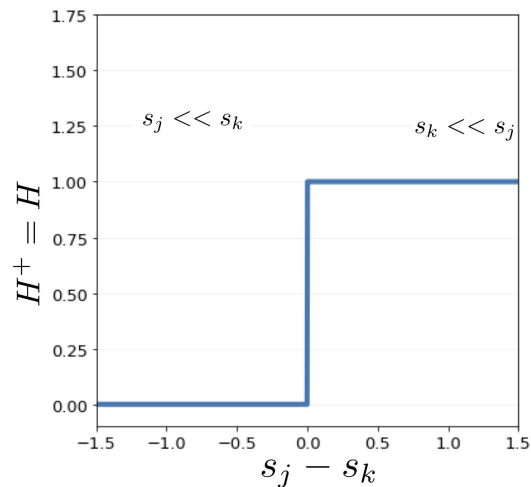
- Fewer works: brute force approaches
 - ◆ Hard negative mining.
 - ◆ Doubling the number of forward passes.
 - ◆ Storing the dataset.
- Overhead in training time.

- Negative mining – Wu, Chao-Yuan, et al. *ICCV*, 2017.
- Large batches – Revaud, Jerome, et al. *ICCV*, 2019.
- cross batch memory – Wang, Xun, et al. *CVPR*, 2020.

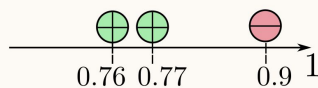
Robust and decomposable rank losses.



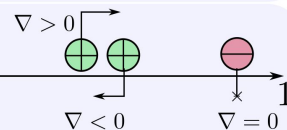
SupRank: smooth approximation of the rank. ROADMAP



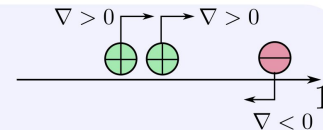
$$\mathcal{L}_{\text{AP}} = 0.41$$



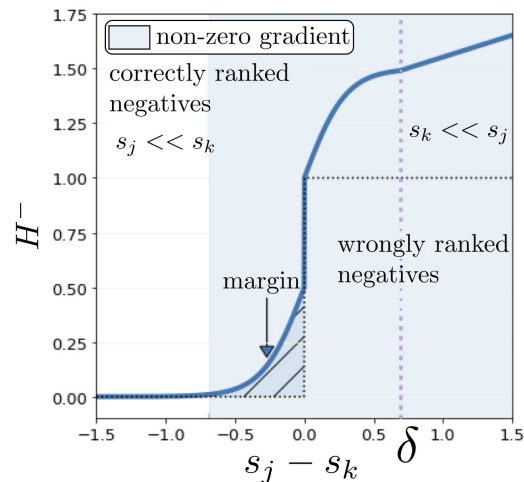
$$\mathcal{L}_{\text{SmoothAP}} = 0.38$$



$$\mathcal{L}_{\text{SupAP}} = 0.44$$



SupRank: smooth approximation of the rank.

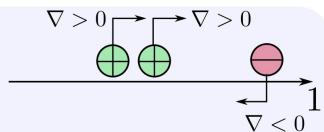


$$\text{AP} = \frac{1}{|\Omega^+|} \sum_{k \in \Omega^+} \frac{\text{rank}^+(k)}{\text{rank}^+(k) + \text{rank}^-(k)}$$

$$\text{rank}_s^-(k) = \sum_{j \in \Omega_k^-} H^-(s_j - s_k)$$

- Optimizing $\text{rank}^- \rightarrow$ smooth approximation + upper bound.
- Not optimizing $\text{rank}^+ \rightarrow$ well-behaved gradients.

$\mathcal{L}_{\text{SupAP}} = 0.44$

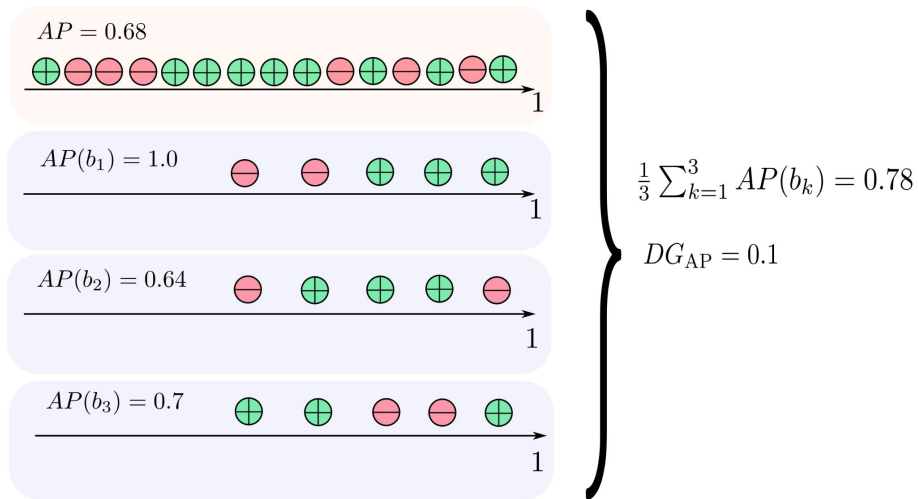


The decomposability gap.

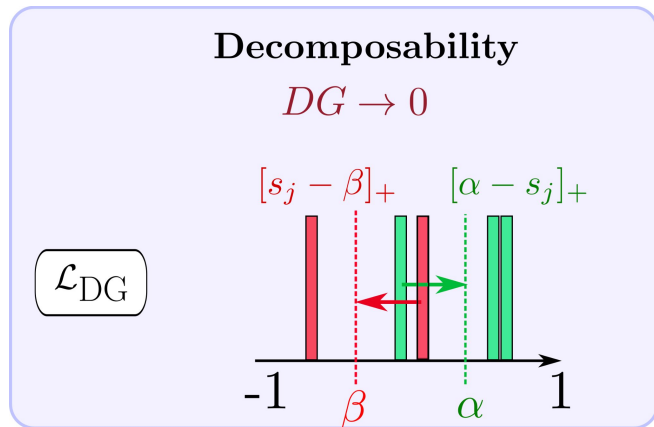


$$DG(\Omega) = \frac{1}{|\mathcal{B}|} \sum_{b \in \mathcal{B}} \mathcal{M}(b) - \mathcal{M}(\Omega)$$

→ Difference between the average ranking-based metrics on batch and its value on the whole dataset.



Decomposable rank losses.



$$\mathcal{L}_{DG}(\theta) = \frac{1}{|\Omega^+|} \sum_{\mathbf{x}_j \in \Omega^+} [\alpha - s_j]_+ + \frac{1}{|\Omega^-|} \sum_{\mathbf{x}_j \in \Omega^-} [s_j - \beta]_+$$

- Calibrates scores across batches:
 - ◆ positive scores greater than α .
 - ◆ negative scores lower than β .
- Explicitly optimize an upper-bound on the decomposability gap.

Instantiation to image retrieval.



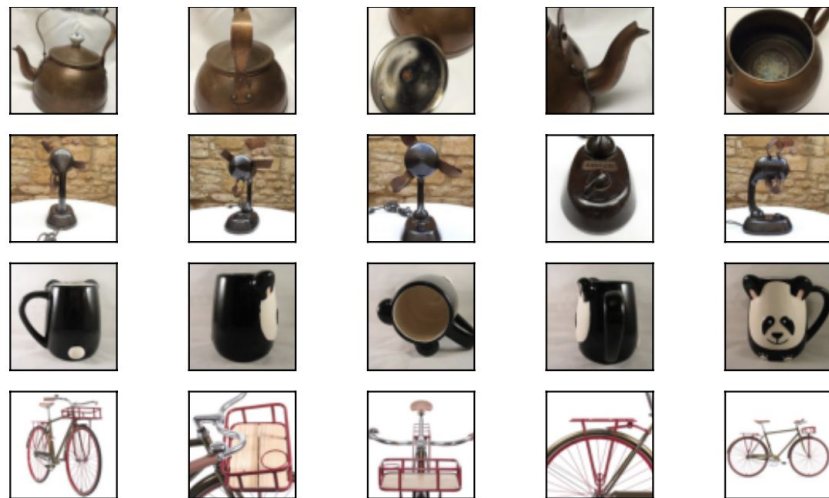
Framework applicable with ranking-based loss: *e.g.* AP, Recall, NDCG.

Application to Average Precision:

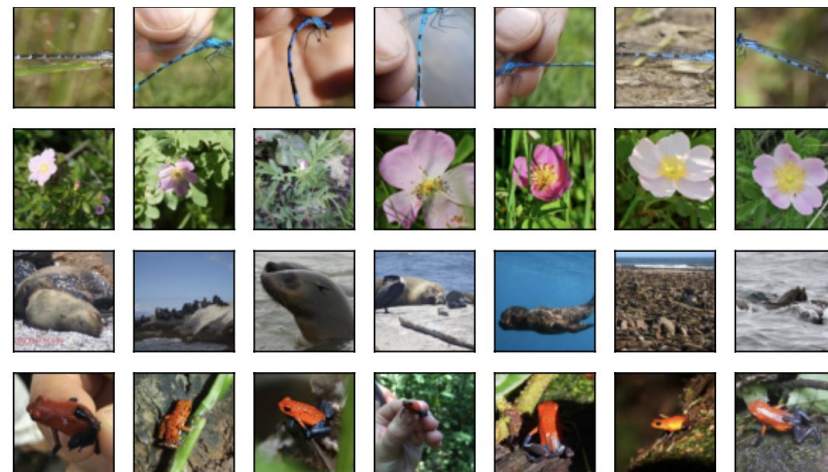
$$\mathcal{L}_{\text{SupAP}} = 1 - \frac{1}{|\Omega^+|} \sum_{k \in \Omega^+} \frac{\text{rank}^+(k)}{\text{rank}^+(k) + \text{rank}_s^-(k)}$$

$$\mathcal{L}_{\text{ROADMAP}}(\boldsymbol{\theta}) = (1 - \lambda) \cdot \mathcal{L}_{\text{SupAP}}(\boldsymbol{\theta}) + \lambda \cdot \mathcal{L}_{\text{DG}}(\boldsymbol{\theta})$$

Experimental validation.



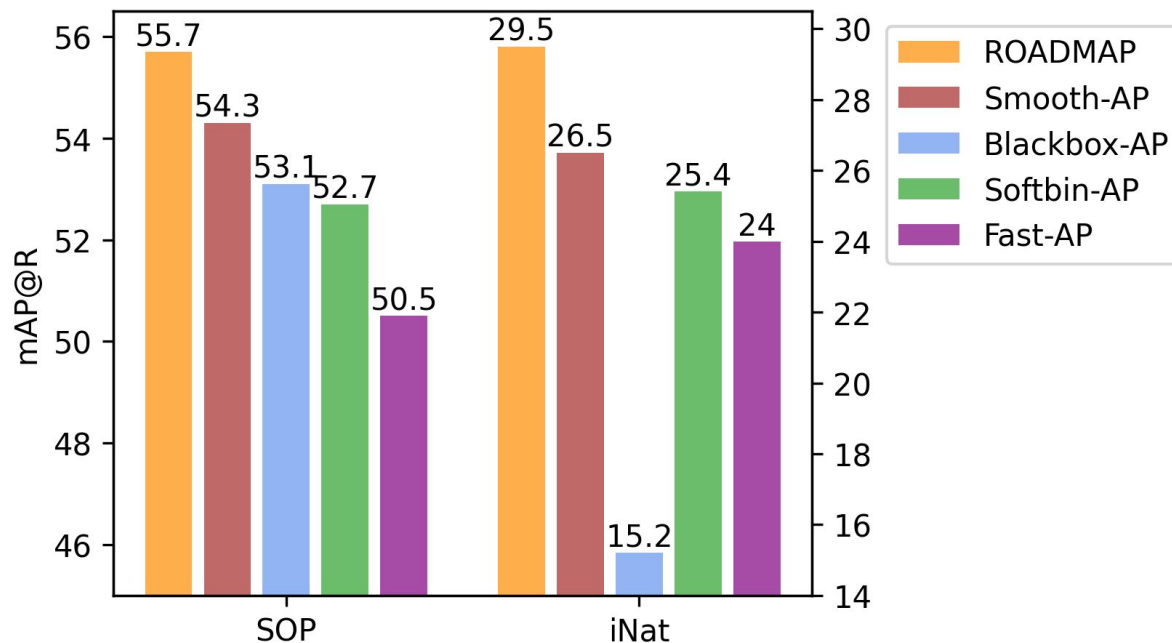
SOP: retail products.



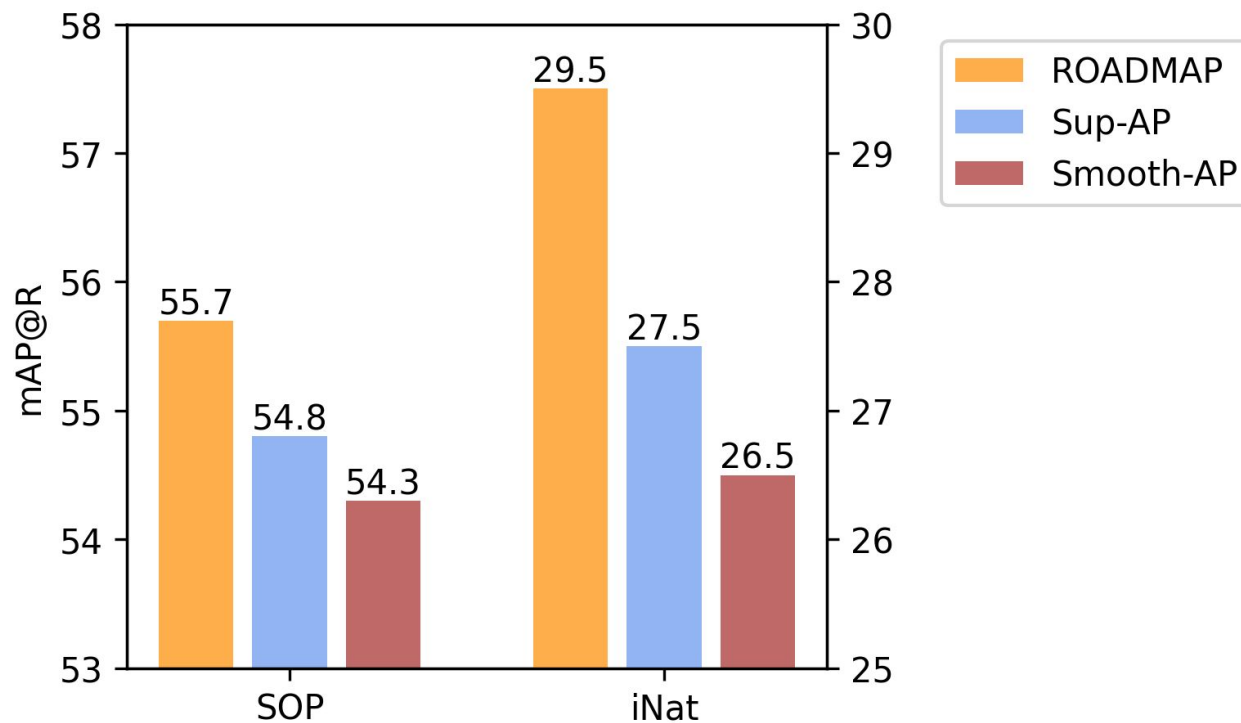
iNaturalist: wildlife images.

- Song, Hyun Oh, et al. “Deep Metric Learning via Lifted Structured Feature Embedding.” *CVPR*, 2016.
- Van Horn, Grant, et al. “The iNaturalist Species Classification and Detection Dataset.” *CVPR*, 2018.

Comparison to AP methods.



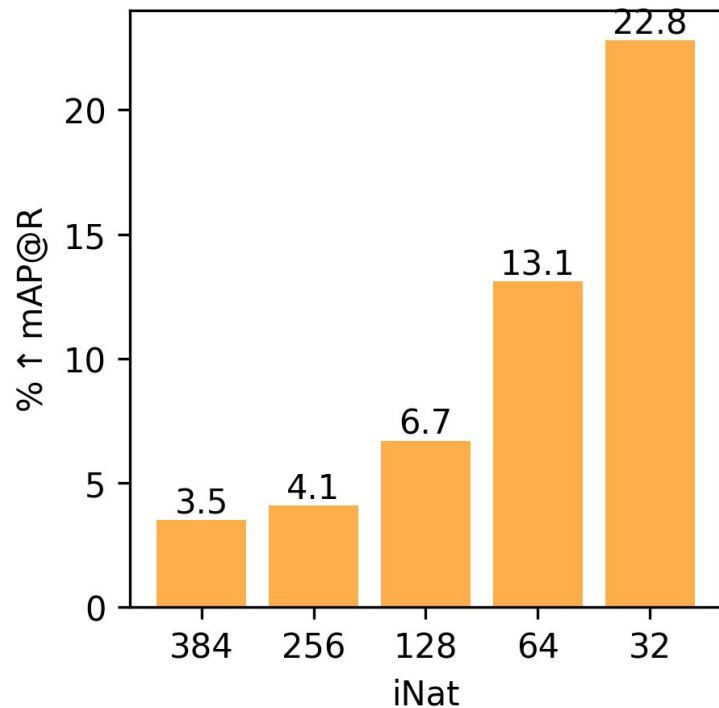
- **Fair comparison** under the same setting (backbone, batch size, data, optimization...)
- Significant gains when **changing the loss**.



→ Significant gain with SupRank.

→ Further boost with decomposability loss.

Impact of the decomposability loss.



→ Biggest **relative** increase on small batches.

→ ROADMAP works with **small batches**.

Conclusion on ROADMAP.



Pros:

- **Smooth** and **decomposable** surrogate for rank losses.
- Consistent and **significant improvements**.

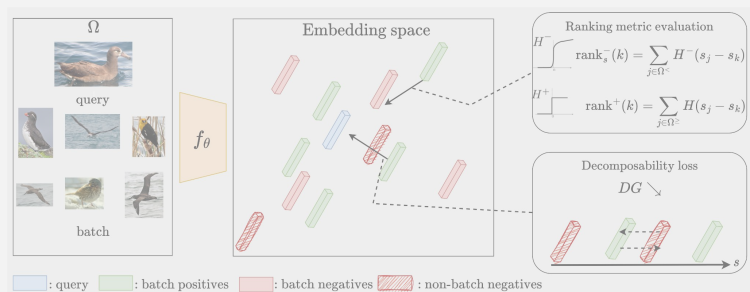
Limitation:

- Only defined for binary labels.

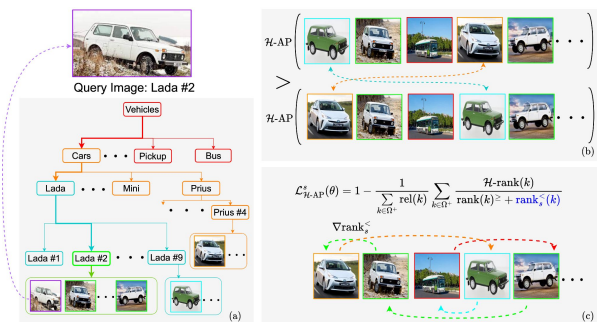


Summary and contributions.

ROADMAP: Optimization of Ranking Losses for Image retrieval.

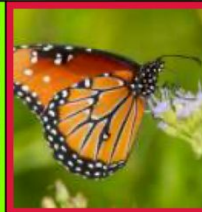


🐱 HAPPIER: Hierarchical Image Retrieval for Robust Ranking.



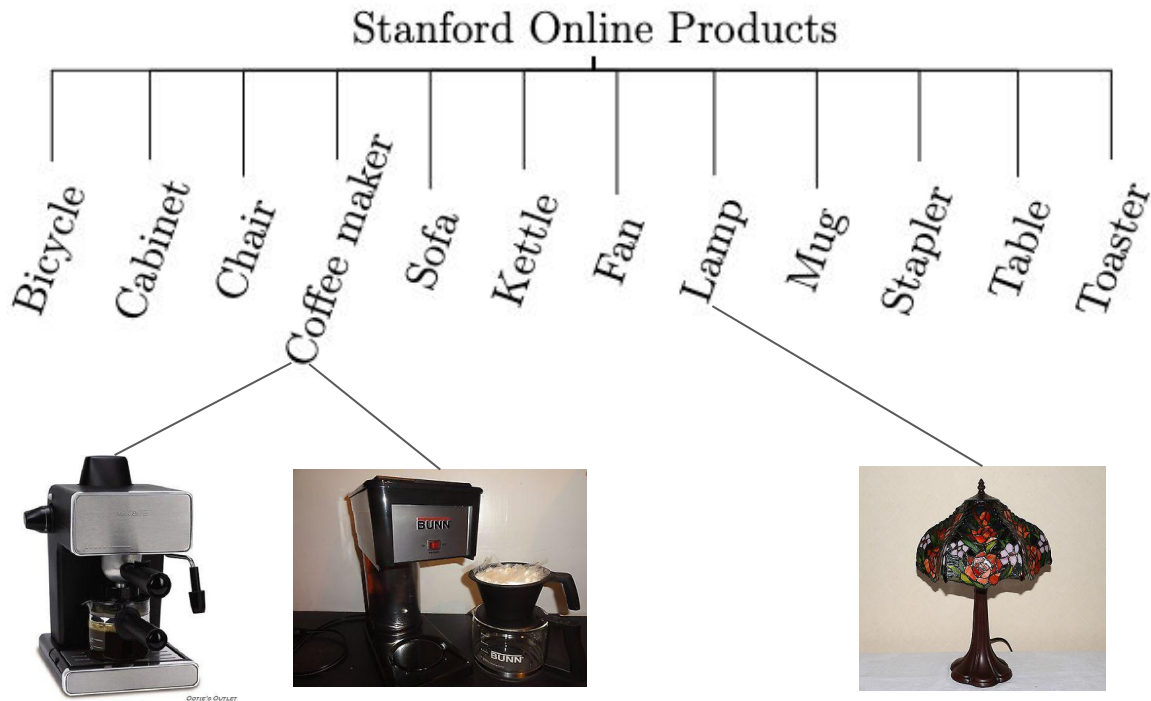
- Ramzi, Elias, et al. "Hierarchical Average Precision Training for Pertinent Image Retrieval." *ECCV*, 2022.
- Ramzi, Elias, et al. "Optimization of Rank Losses for Image Retrieval." *TPAMI*, 2025.

Hierarchical Image Retrieval for Robust Ranking. HAPPIER

		HAPPIER					
		rank 1	rank 2	rank 3	rank 4	rank 5	rank 6
	$\mathcal{H}$ -AP						
	AP						
	0.94						
		Baseline					
	$\mathcal{H}$ -AP						
	AP						
	0.68						

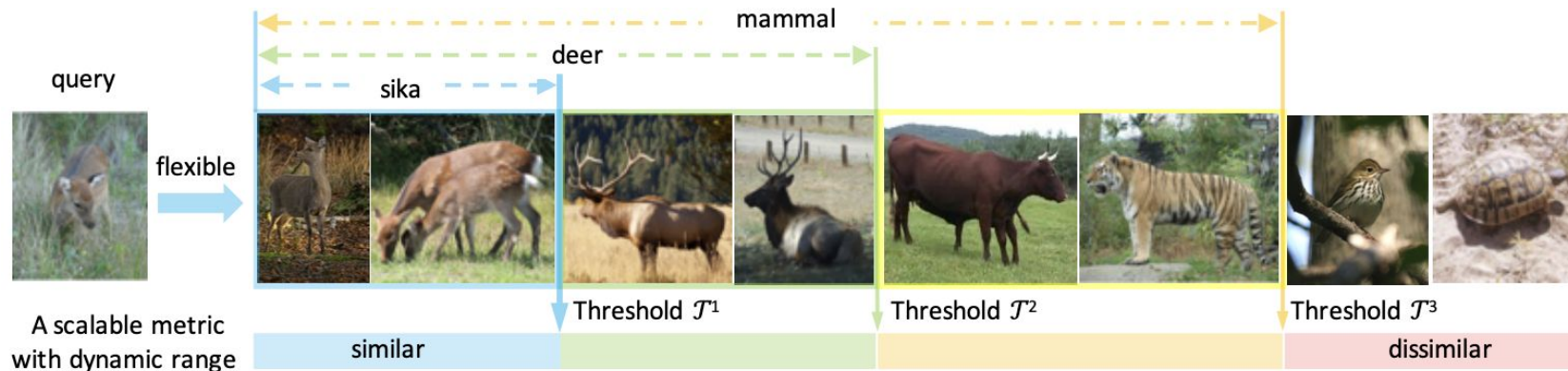
- Image retrieval is binary → do not take into account mistake severity.
- Extend AP to graded setting to take importance of errors into account.

Hierarchical learning.



- Hierarchy is a good proxy for human perception of mistake severity.
- Several public datasets include hierarchical annotations.

Hierarchical image retrieval in the literature.

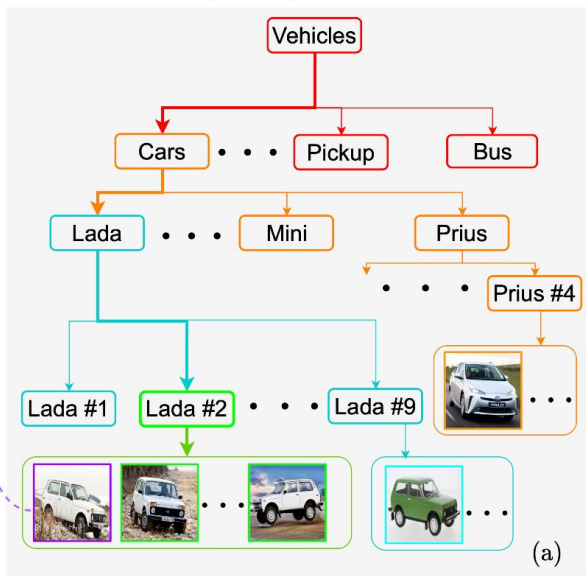


- Introduction of the “Dynamic Metric learning” datasets.
- CSL = triplet + proxy loss for hierarchical setting.
 - ◆ Limitation: **not aligned** with evaluation metrics.

■ Sun, Yifan, et al. "Dynamic metric learning: Towards a scalable metric space to accommodate multiple semantic scales." *CVPR*, 2021.



Query Image: Lada #2





- Leverage a hierarchical tree to design a **graded similarity**, or “relevance” between categories.
- **Decreasing function of the distance** in the hierarchical tree.

$$\text{rel}(k) = \frac{l/L}{|\Omega^{(l)}|}$$

- l : level of the closest ancestor in the tree.
- L total number of levels.

$$\mathcal{H}\text{-rank}(k) = \text{rel}(k) + \sum_{j \in \Omega^+} \min(\text{rel}(k), \text{rel}(j)) \cdot H(s_j - s_k)$$



Query Image:
Lada #2



$$\mathcal{H}\text{-rank}(1) = \frac{1}{3} = 1/3$$

$$\mathcal{H}\text{-rank}(2) = 1 + \frac{1}{3} = 4/3$$

$$\mathcal{H}\text{-rank}(4) = \frac{2}{3} + \left(\frac{1}{3} + \frac{2}{3}\right) = 5/3$$

$$\mathcal{H}\text{-rank}(5) = 1 + \left(\frac{1}{3} + 1 + \frac{2}{3}\right) = 3$$

Lada #2
rel := 1

Lada #9
rel := 2/3

Prius #4
rel := 1/3

Bus
rel := 0

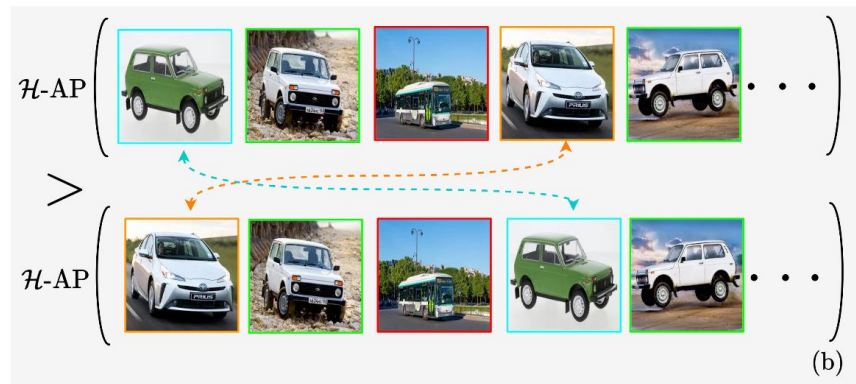
- Errors in ranking → **weighted by relevance.**
- Correct $\mathcal{H}$ -rank → decreasing order of relevance.

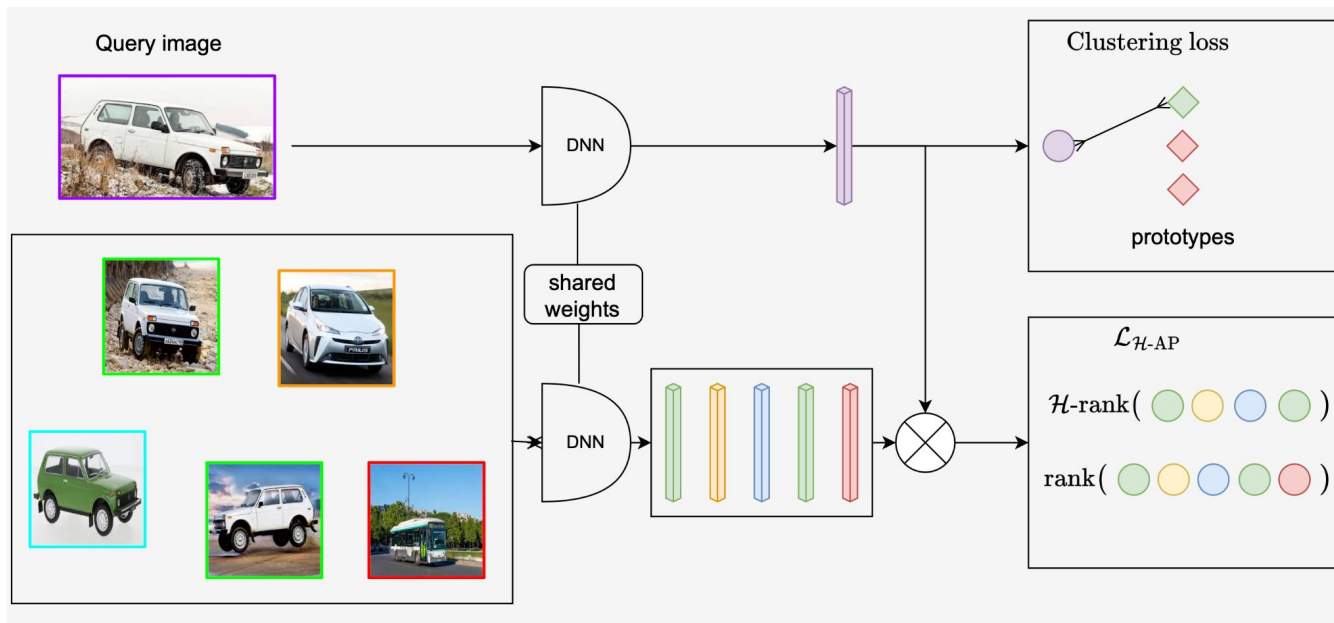
Hierarchical average precision.

$$\mathcal{H}\text{-AP} = \frac{1}{\sum_{k \in \Omega^+} \text{rel}(k)} \sum_{k \in \Omega^+} \frac{\mathcal{H}\text{-rank}(k)}{\text{rank}(k)}$$

$\mathcal{H}\text{-AP}$ properties:

- Consistent generalization of AP.
- Keeps desirable properties of AP.
 - Penalize wrong rankings.
 - Emphasis for the top of the ranking.
 - Relevant in imbalanced settings.
- Is flexible wrt. the relevance.

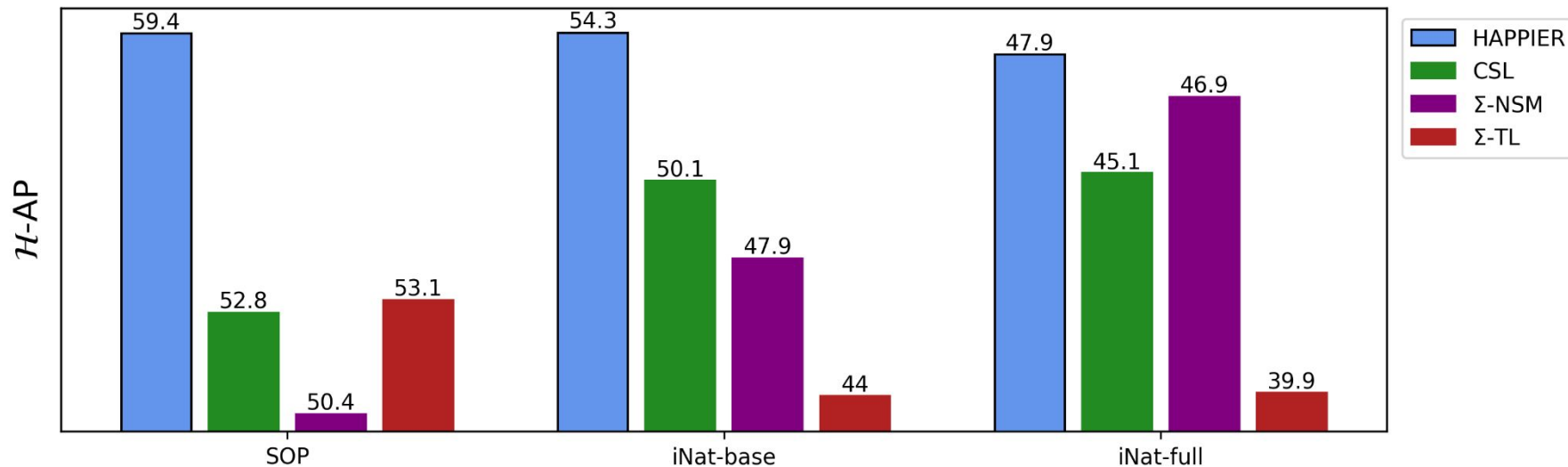




$$\mathcal{L}_{\text{Sup-}\mathcal{H}\text{-AP}}(\theta) = 1 - \frac{1}{\sum_{k \in \Omega^+} \text{rel}(k)} \sum_{k \in \Omega^+} \frac{\mathcal{H}\text{-rank}(k)}{\text{rank}^+(k) + \text{rank}_s^-(k)}$$

$$\mathcal{L}_{\text{HAPPIER}}(\theta) = (1 - \lambda) \cdot \mathcal{L}_{\text{Sup-}\mathcal{H}\text{-AP}}(\theta) + \lambda \cdot \mathcal{L}_{\text{DG}}^*(\theta)$$

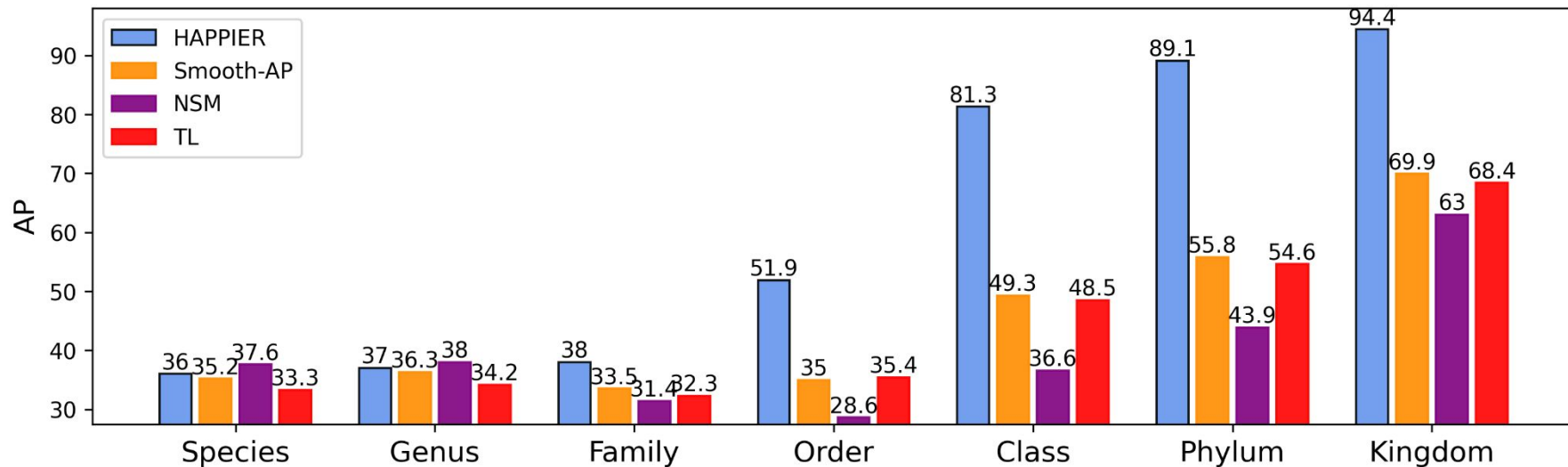
Comparison to hierarchical methods.



→ HAPPIER outperforms the state-of-the-art hierarchical loss CSL on fine-grained and hierarchical retrieval.

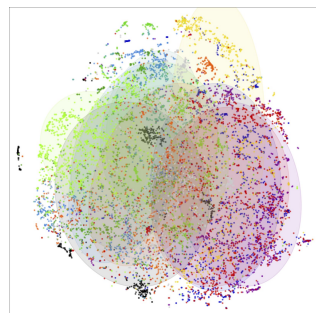
- TL: Wu, Chao-Yuan, et al. *ICCV*, 2017.
- NSM: Zhai, Andrew, et al. *BMVC*, 2018.
- CSL: Sun, Yifan, et al. *CVPR*, 2021.

Comparison to fine-grained methods.

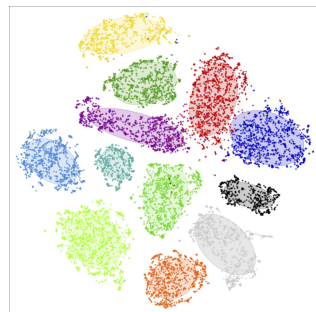


- On par for fine-grained retrieval (“Species”).
- **Large gains** on other hierarchical levels from “Family”.

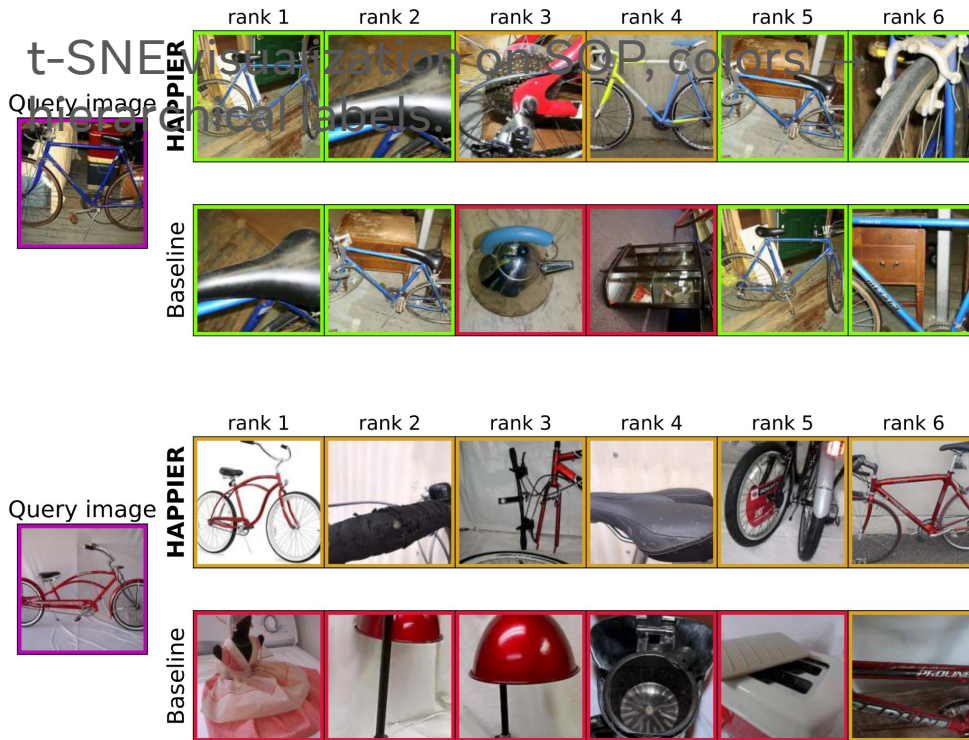
Qualitative results.



Baseline



HAPPIER



HAPPIER:

→ High quality of clusters.

→ Better mistakes, even in failure cases.

H-GLDv2: a hierarchical landmark dataset.

Query:



Relevant index images:



GLDv2 → **large scale** landmarks retrieval dataset.

No hierarchical annotations → how difficult is it to create hierarchical annotations?

Scraping Wikimedia Commons.

St Oswald's Church, Dean [\[Collapse\]](#)

church in Dean, Cumbria, UK



Upload media: [?](#)

[Wikipedia](#)

Instance of [church building](#)

Dedicated to Oswald of Worcester

Made from material calciferous sandstone

Location Dean, Allerdale, Cumbria, North West England, England, UK

Architectural style English Gothic architecture

Norman architecture

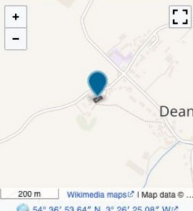
Diocese Diocese of Carlisle

Heritage designation Grade I listed building (1986–)

Inception 12th century

Religion or worldview Anglicanism

[official website](#) [?](#)



Buena Vista Lagoon [\[Collapse\]](#)

lake in United States of America



Upload media: [?](#)

[Wikipedia](#)

Instance of [lake](#)

Location California



300 m Wikimedia maps: [?](#) | Map data © O...

[Authority control](#) [\[Collapse\]](#)

[Q4985455](#)

[Reasonator](#) [?](#) • [Scholia](#) [?](#) • [PetScan](#) [?](#) • [statistics](#) [?](#) • [WikiMap](#) [?](#) • [Locator tool](#) [?](#) • [KML file](#) [?](#) • [WikiShootMe](#) [?](#) • [OpenStreetMap](#) [?](#) • [Search depicted](#)

Shōmyō Falls [\[Collapse\]](#)

waterfall in Toyama Prefecture, Japan



Upload media: [?](#)

[Wikipedia](#)

Instance of [waterfall](#)

Part of Japan's Top 100 Waterfalls (38)

Named after Shōmyō

Located in protected area Chūbu-Sangaku National Park

Location Ashikuraji, Tateyama, Nakanikawa district, Toyama Prefecture, Japan

Located in or next to body of water Shōmyō River

Heritage designation Place of Scenic Beauty of Japan

natural monument

Height 350 m

126 m (Q46868895)

Elevation above sea level 1,260 m

Mosquée de la Divinité [\[Collapse\]](#)

building in Senegal



Upload media: [?](#)

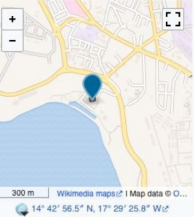
[Wikipedia](#)

Instance of [mosque](#)

Location Dakar, Dakar Department, Dakar, Senegal

Start time 1997

Religion or worldview Islam



300 m Wikimedia maps: [?](#) | Map data © O...

[Authority control](#) [\[Collapse\]](#)

[Q3324921](#)

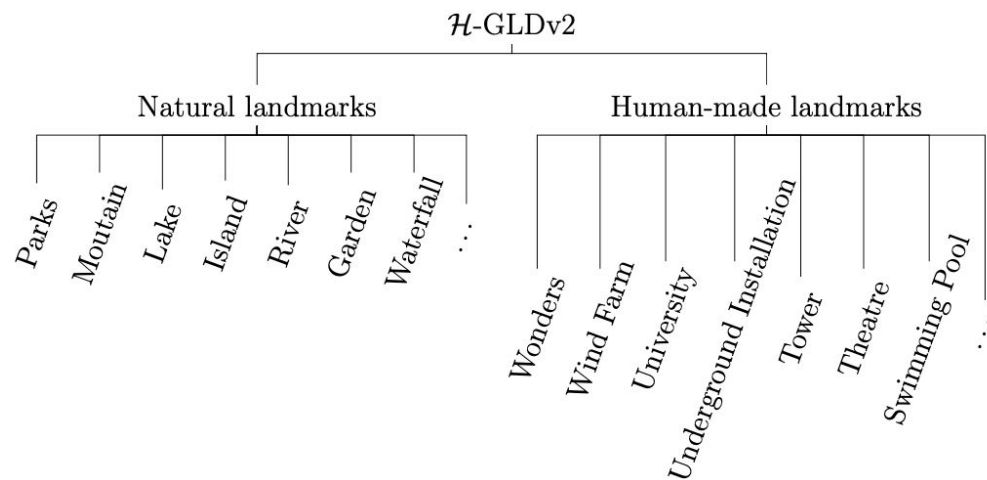
[Reasonator](#) [?](#) • [Scholia](#) [?](#) • [PetScan](#) [?](#) • [statistics](#) [?](#) • [WikiMap](#) [?](#) • [Locator tool](#) [?](#) • [KML file](#) [?](#) • [WikiShootMe](#) [?](#) • [OpenStreetMap](#) [?](#) • [Search depicted](#)

Wikimedia Commons → the largest open database of landmarks.

Scraped labels include:

- Church building.
- Church building (1172–1954)
- Cathedral.
- Castle.
- Corsican nature reserve.
- New Zealand great walks.
- Waterfall.
- Arch-gravity dam.
- Canal.
- Association football venue.
- Astronomical observatory.
- Village.

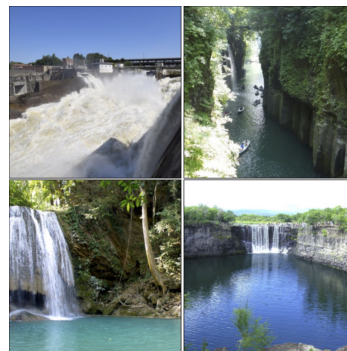
Final super categories.



Bridge.



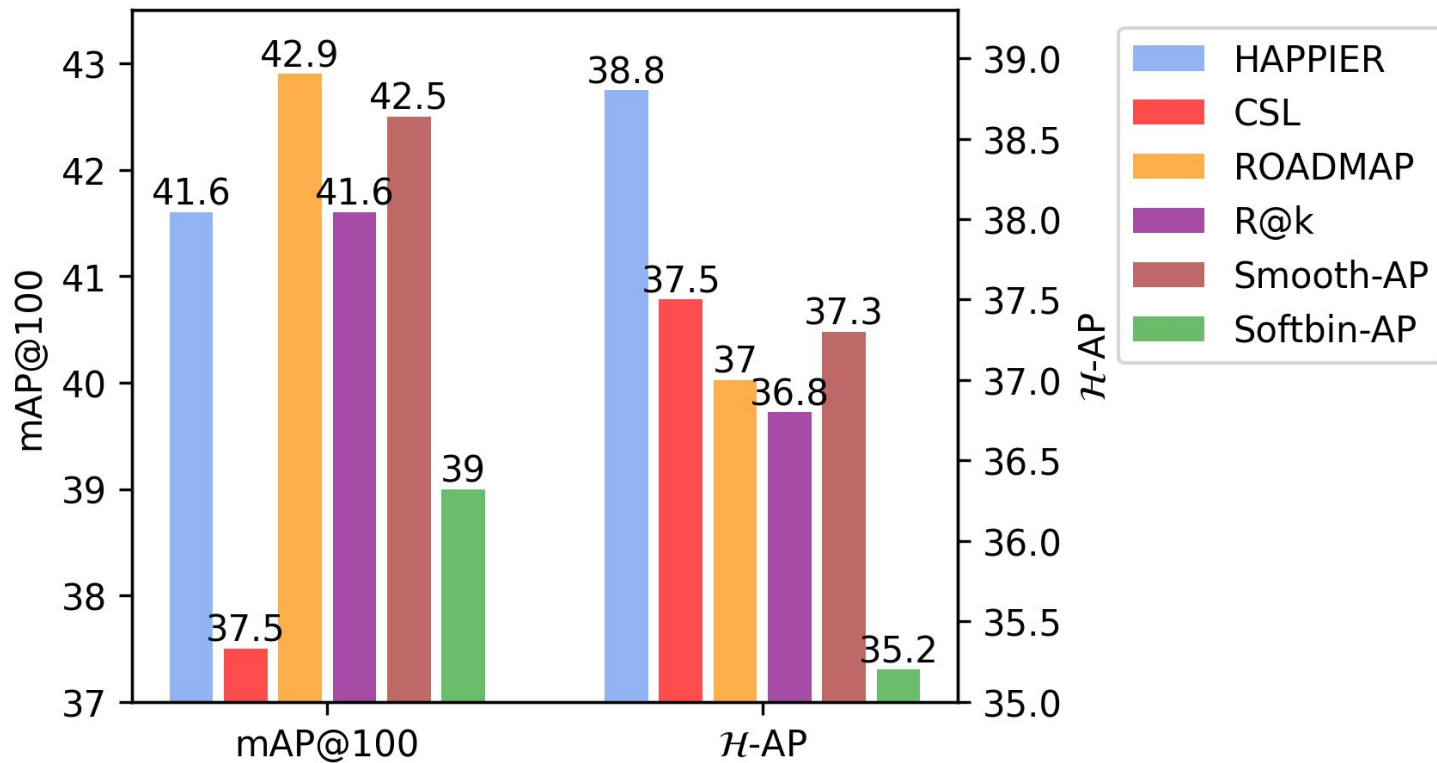
Castle.



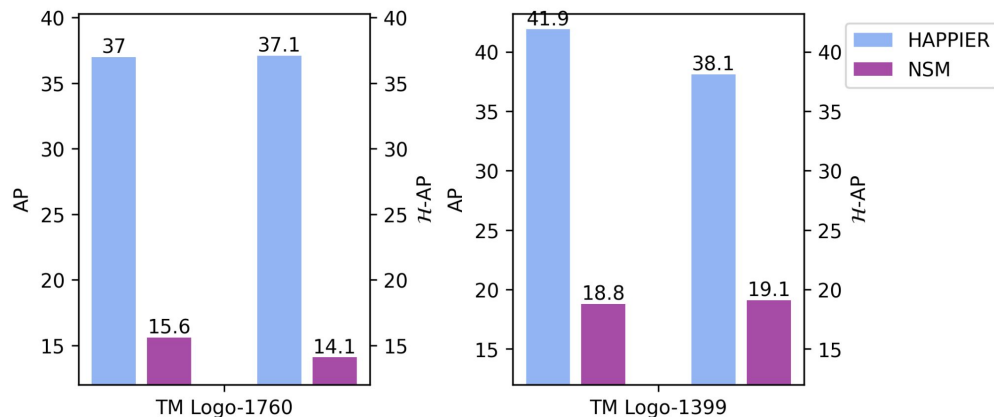
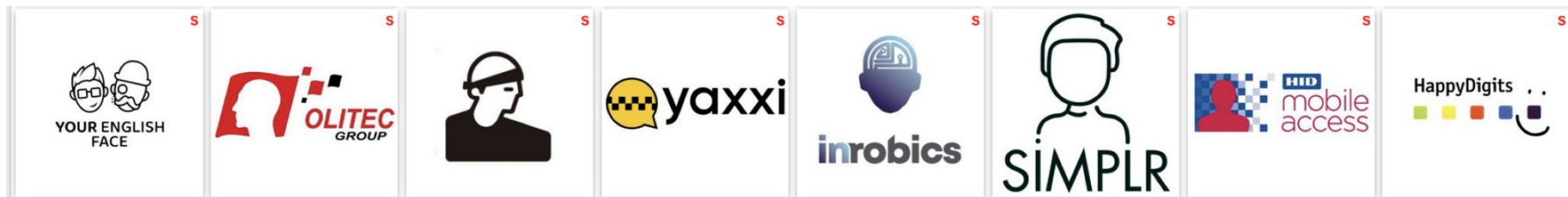
Waterfall.



Volcano.



Results on Coexya data.



→ HAPPIER can be extended to **multi-label setting**.

→ Optimizing HAPPIER works best on both **fine-grained** and **hierarchical metrics**.

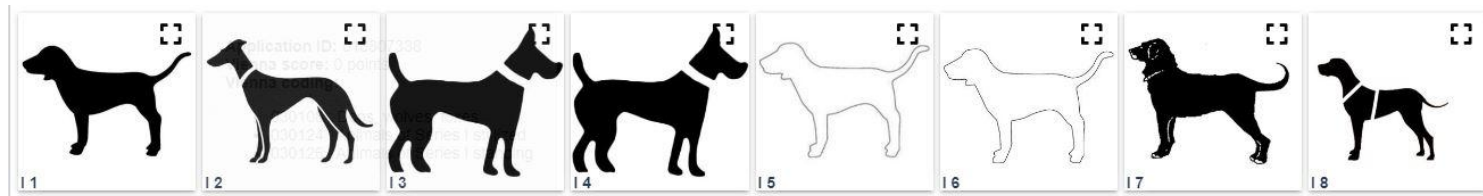
Qualitative results on Coexya data.



HAPPIER



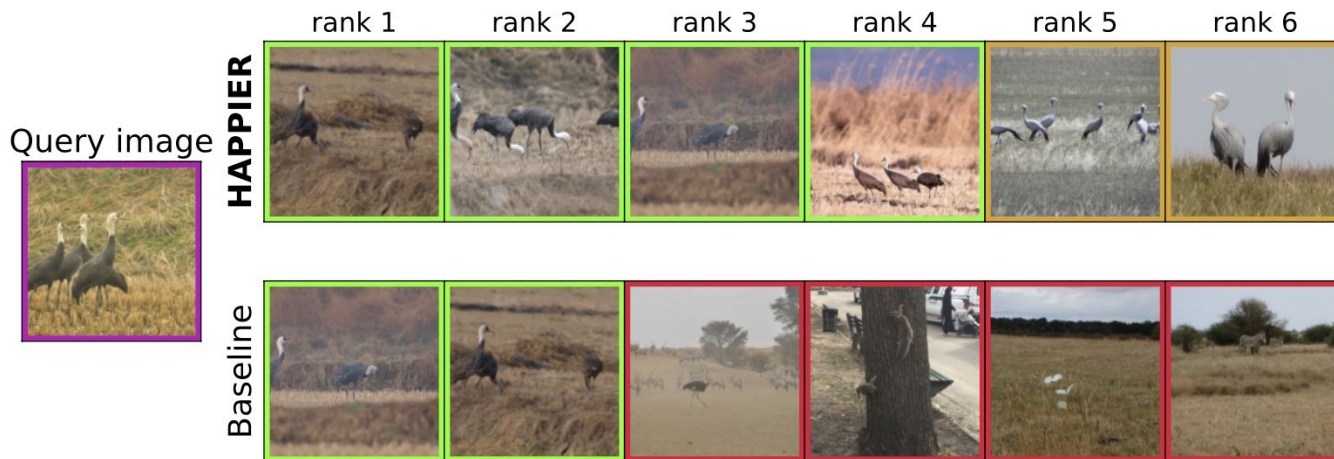
Base



- HAPPIER → retrieves images with **both labels**.

Conclusion on HAPPIER.

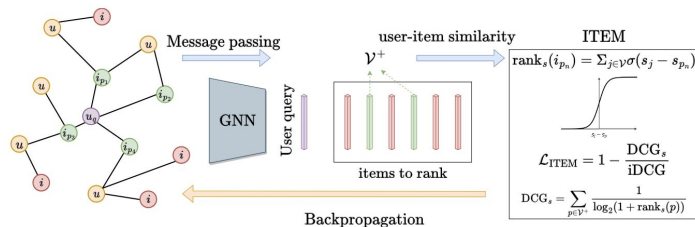
- Optimizing hierarchical metrics:
 - **Better robustness** wrt. to mistake severity.
 - Performances on par on fine-grained metrics.
- Hierarchical annotations → **not too costly** + boost models' robustness.
- HAPPIER → production in Acsepto



Short term extensions.

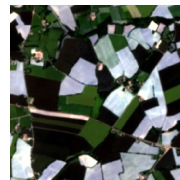
ROADMAP:

- Optimize other metrics: Recall, NDCG.
- Use in GNN recommender systems.

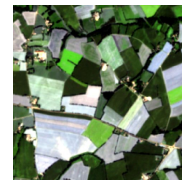


🐱 HAPPIER:

- Use in multi-label setting + business driven insights.
- Aerial image time series ranking.



2019-06-01



2019-08-30



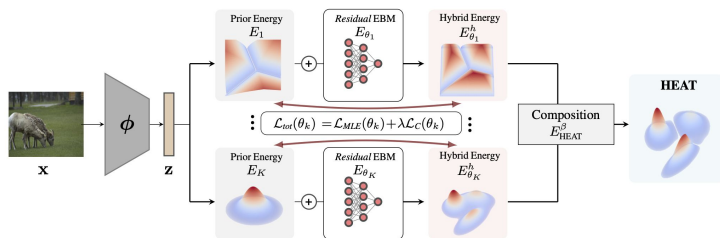
2019-09-19

Few shot
robustness in the
era of foundation
models.

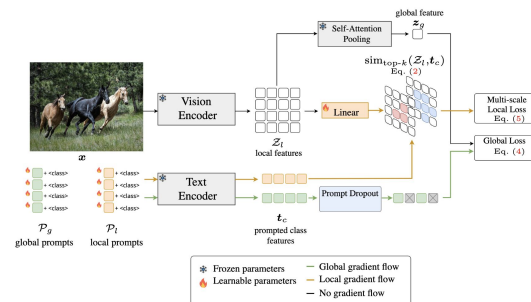
Summary and contributions.



HEAT:
Post-hoc out-of-distribution
detection.



GalLoP:
Few shot adaptation of
vision-language models.

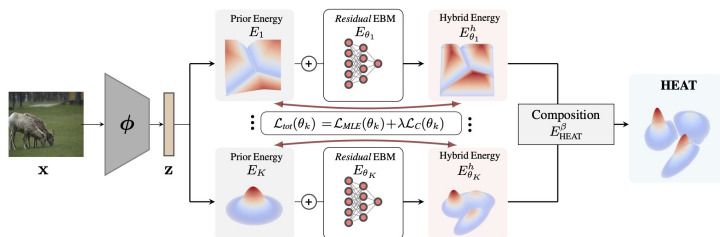


■ Lafon, Marc, et al. "Hybrid Energy Based Model in the Feature Space for Out-of-Distribution Detection." *ICML*, 2023.

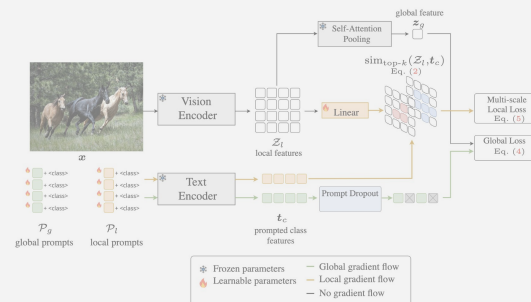
■ Lafon, Marc, et al. "GalLoP: Learning global and local prompts for vision-language models", *ECCV*, 2024

Summary and contributions.

🔥 HEAT: Post-hoc out-of-distribution detection.



🐎 GalLoP: Few shot adaptation of vision-language models.



■ Lafon, Marc, et al. "Hybrid Energy Based Model in the Feature Space for Out-of-Distribution Detection." *ICML*, 2023.

Out-of-distribution detection.



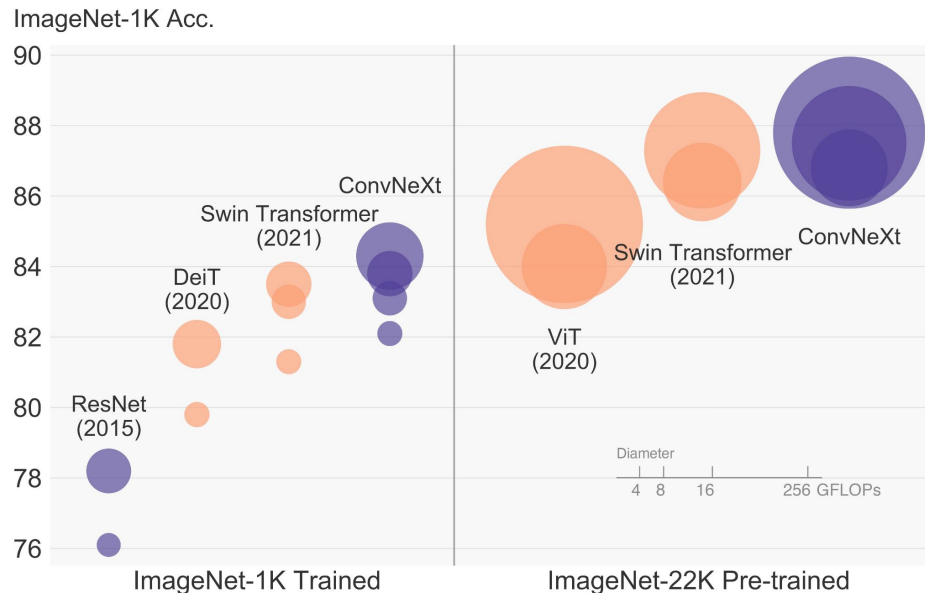
AI confuses a horse-drawn carriage (OOD) for a truck.

- Detection of in-distribution (ID) = similar to train dataset, vs. out-of-distribution (OOD) = dissimilar to train dataset.
- Important task for real world deployment.

OOD $\rightarrow$ binary problem:

$$G_{\lambda}(x) = \begin{cases} \text{ID, if } E(x) \leq \lambda_{95\%} \\ \text{OOD, if } E(x) > \lambda_{95\%} \end{cases}$$

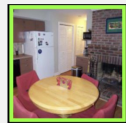
Post-hoc out-of-distribution detection.



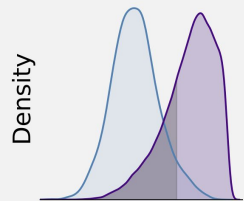
- Enjoy strong predictive performances from off-the-shelf neural networks.
- Use any backbones → ConvNets, transformers.
- Limited compute overhead at inference time.

Prior out-of-distribution scorers.

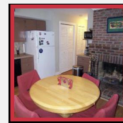
Prior K - (EL)



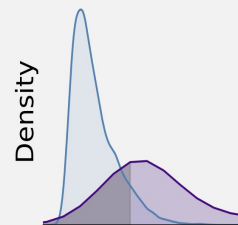
$FPR_{Near} = 45.2\%$
 $FPR_{Far} = 17.8\%$



Prior 1 - (GMM)



$FPR_{Near} = 51.6\%$
 $FPR_{Far} = 7.0\%$



Classification based methods:

- Better at distinguishing near-OOD.
- No control for far-OOD.

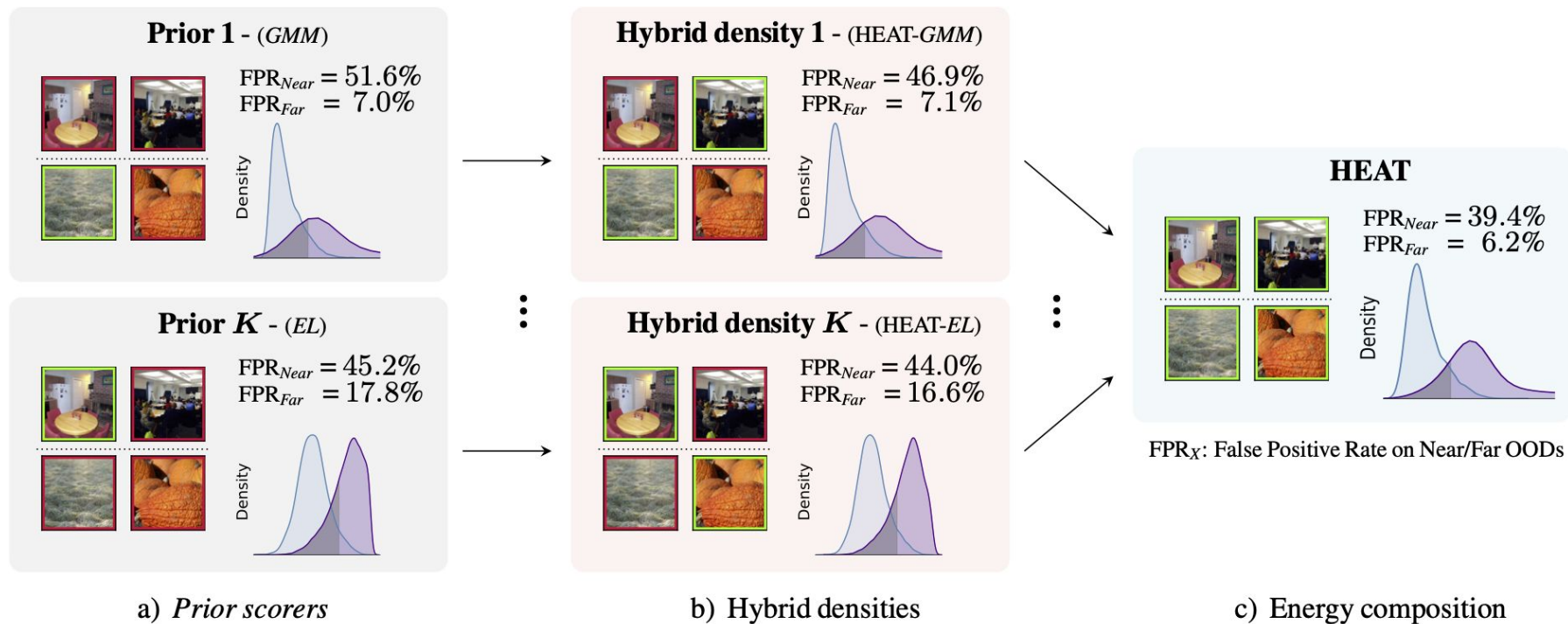
Parametric methods:

- Strong guarantees for far-OOD.
- Lower performances on near-OOD.

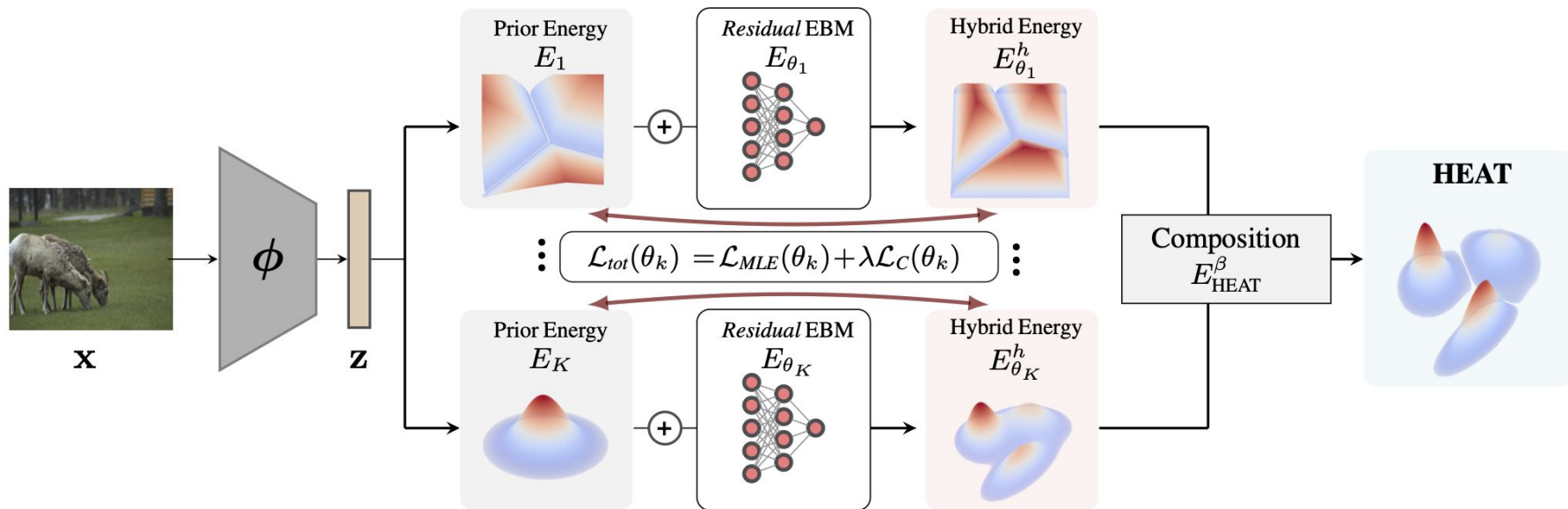
■ Energy logits (EL): Liu, Weitang, et al. NeurIPS, 2020.
■ DICE: Sun, Yiyong, et al. ECCV, 2022.

■ SSD (GMM): Sehwal, Vikash, et al. ICLR, 2021.
■ KNN: Sun, Yiyong, et al. ICML, 2022.

Motivations.



HEAT for out-of-distribution detection.



- EBM $\rightarrow$ data driven residual learning.
- Composition $\rightarrow$ leverage $\neq$ modeling biases.

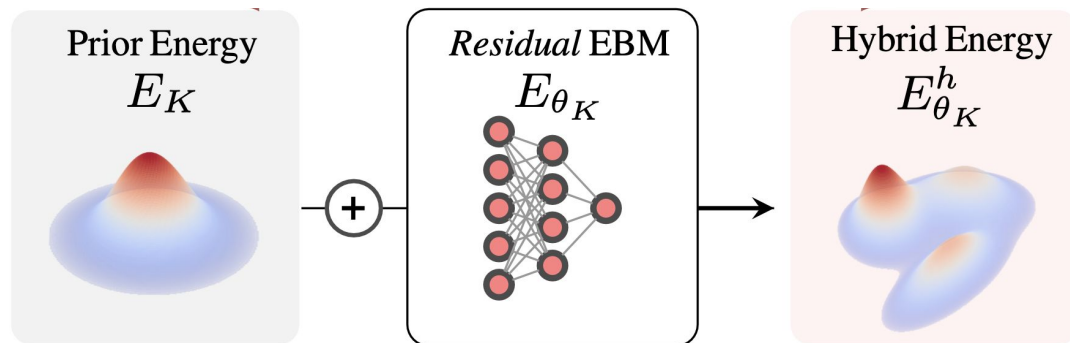
Hybrid energy-based density estimation.

Hybrid density:

$$p_{\theta_k}^h(z) = \frac{1}{Z(\theta_k)} p_{\theta_k}^r(z) q_k(z)$$

Hybrid energy:

$$E_{\theta_k}^h(z) = E_{q_k}(z) + E_{\theta_k}(z)$$



 EBM – LeCun, Yann, et al. *Predicting structured data 1.0* (2006).

 Energy correction – Deng, Yuntian, et al. *ICLR 2020*.

Composition of refined prior density estimators. HEAT

Class prior:

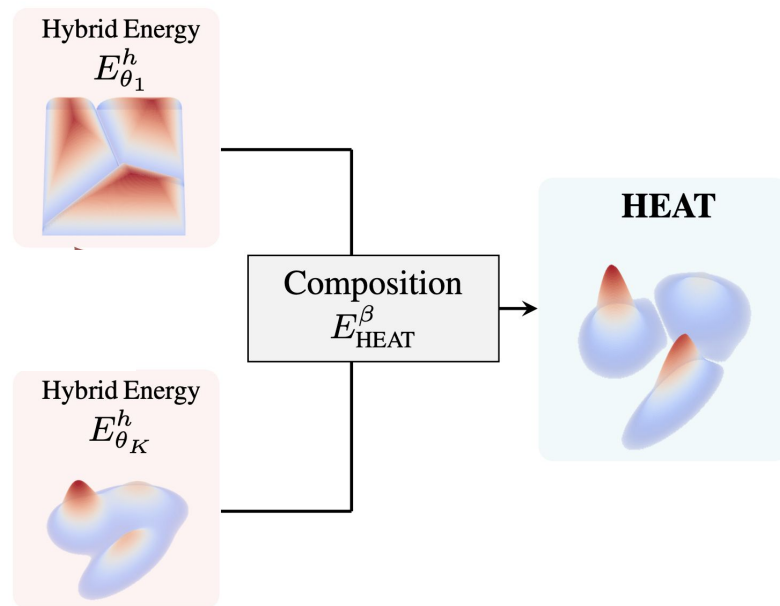
$$E_{\theta_l}^h(z) = -\log \sum_c e^{f(z)[c]} + E_{\theta_l}^r(z)$$

Feature & style prior:

$$E_{\theta_g}^h(z) = -\log \sum_c e^{-\frac{1}{2}(z-\mu_c)^T \Sigma^{-1}(z-\mu_c)} + E_{\theta_g}^r(z)$$

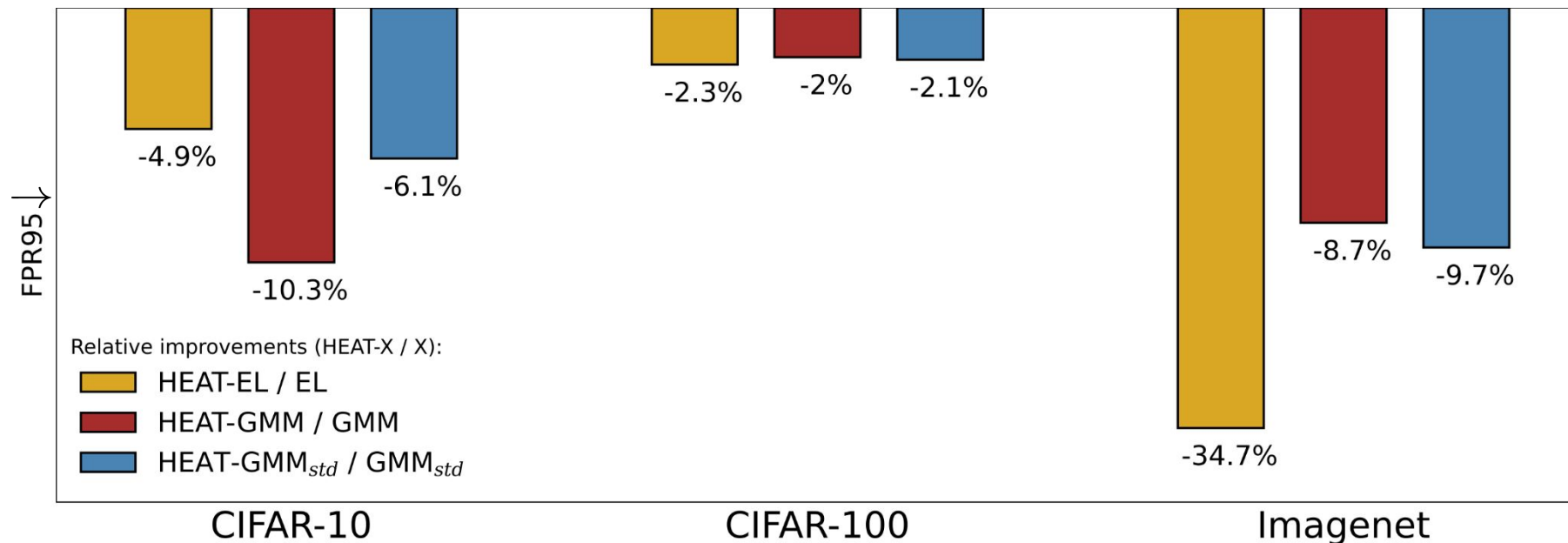
Energy function composition:

$$E_{\text{HEAT}}^\beta = \frac{1}{\beta} \log \sum_{k=1}^K e^{\beta E_{\theta_k}^h}$$

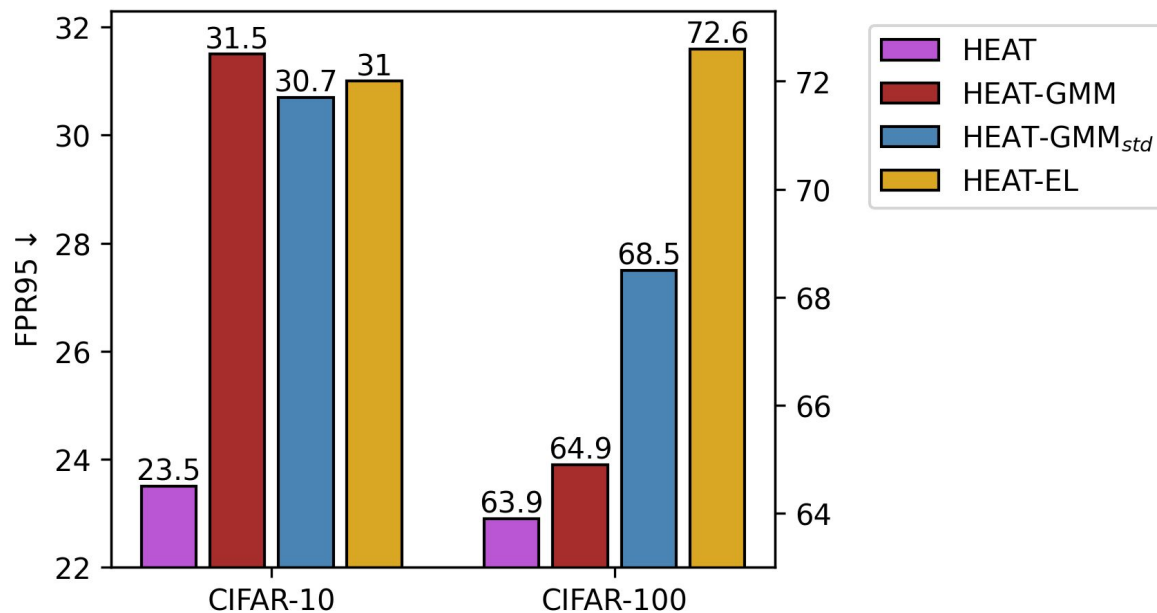


- EL – Liu, Weitang, et al. NeurIPS, 2020.
- GMM – Lee, Kimin, et al. NeurIPS, 2018.
- Gram matrix – Sastry, Chandramouli Shama, et al. ICML, 2020.

Residual learning.

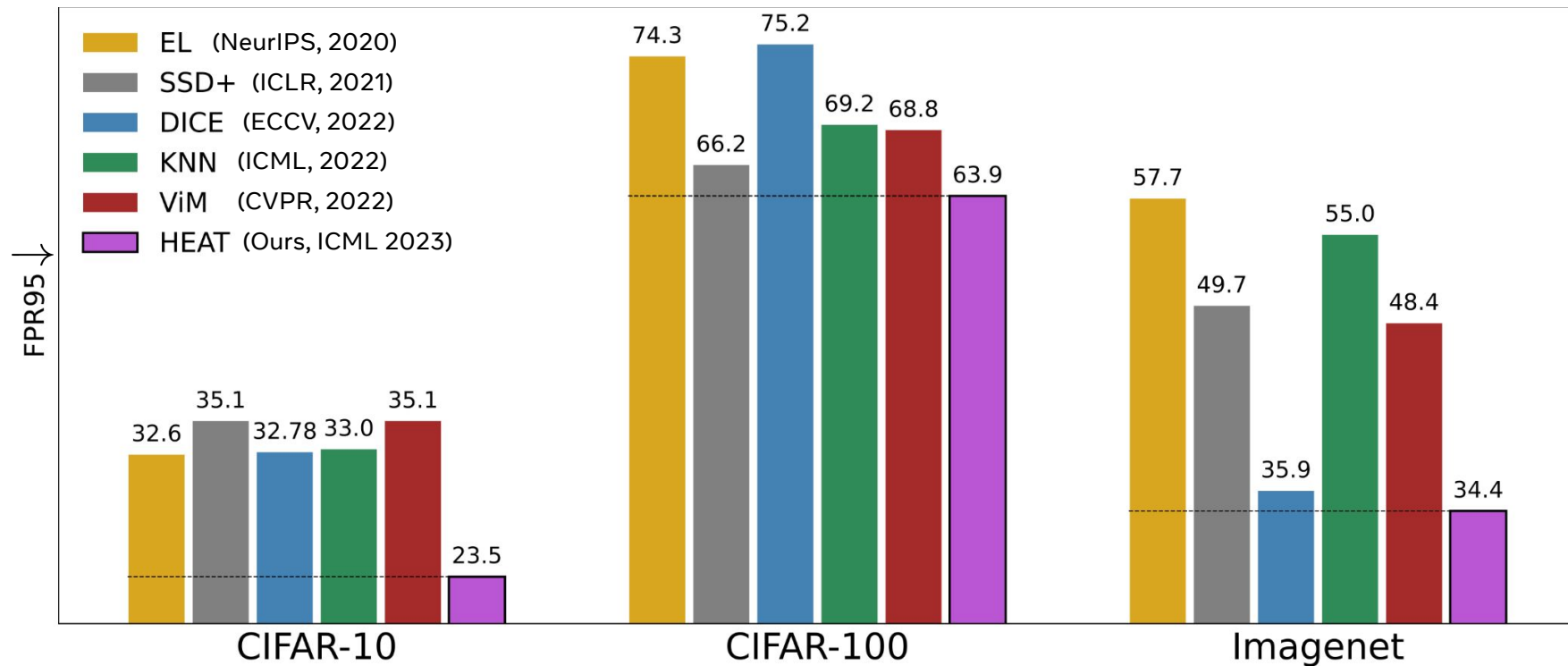


- Residual learning → better performances



- Results better than individual scorers.
- Composition → performance boost.

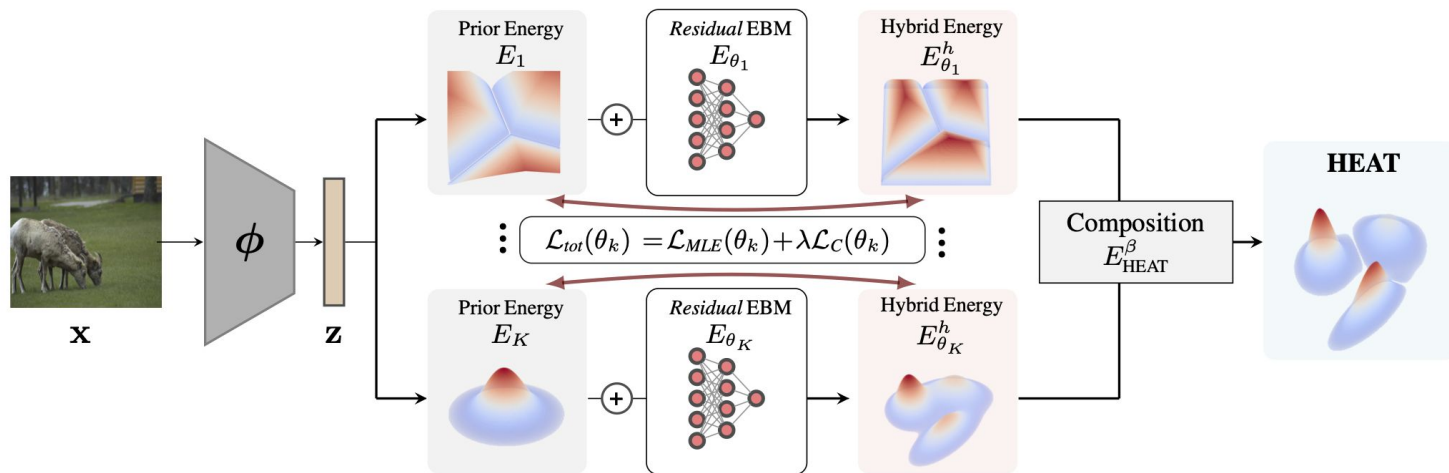
Comparison to state-of-the-art methods.



- HEAT → competitive performances on 3 datasets

Conclusion on HEAT.

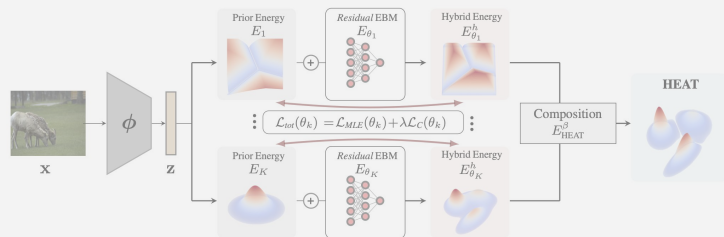
- Post-hoc OOD detection $\rightarrow$ no assumption on models.
- EBM models boost prior OOD scorers.
- Composing refined OOD scorers $\rightarrow$ boost performances.



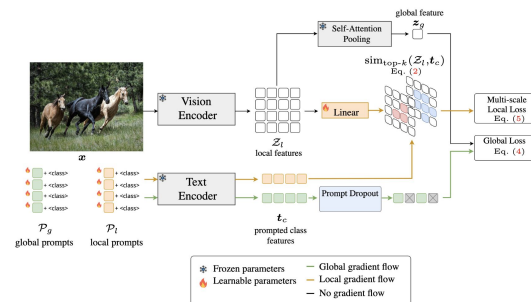
Summary and contributions.



HEAT:
Post-hoc out-of-distribution
detection.



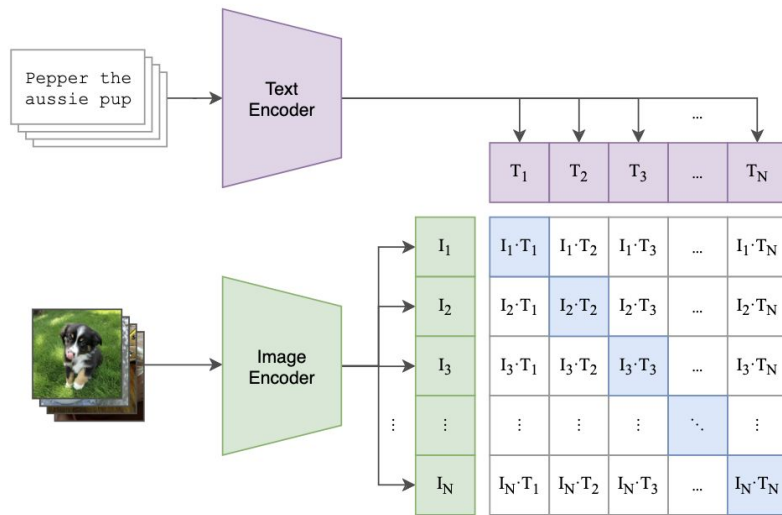
GalLoP:
Few shot adaptation of
vision-language models.



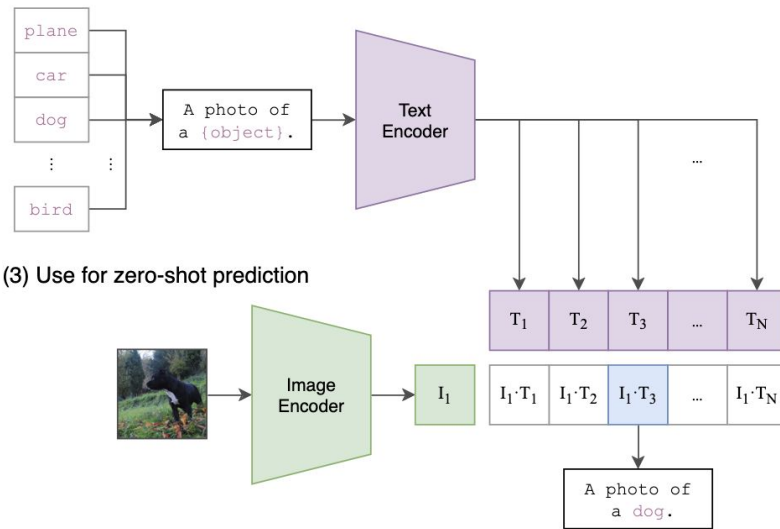
■ Lafon, Marc, et al. “GalLoP: Learning global and local prompts for vision-language models”, *ECCV*, 2024

Vision-language models.

(1) Contrastive pre-training



(2) Create dataset classifier from label text



(3) Use for zero-shot prediction

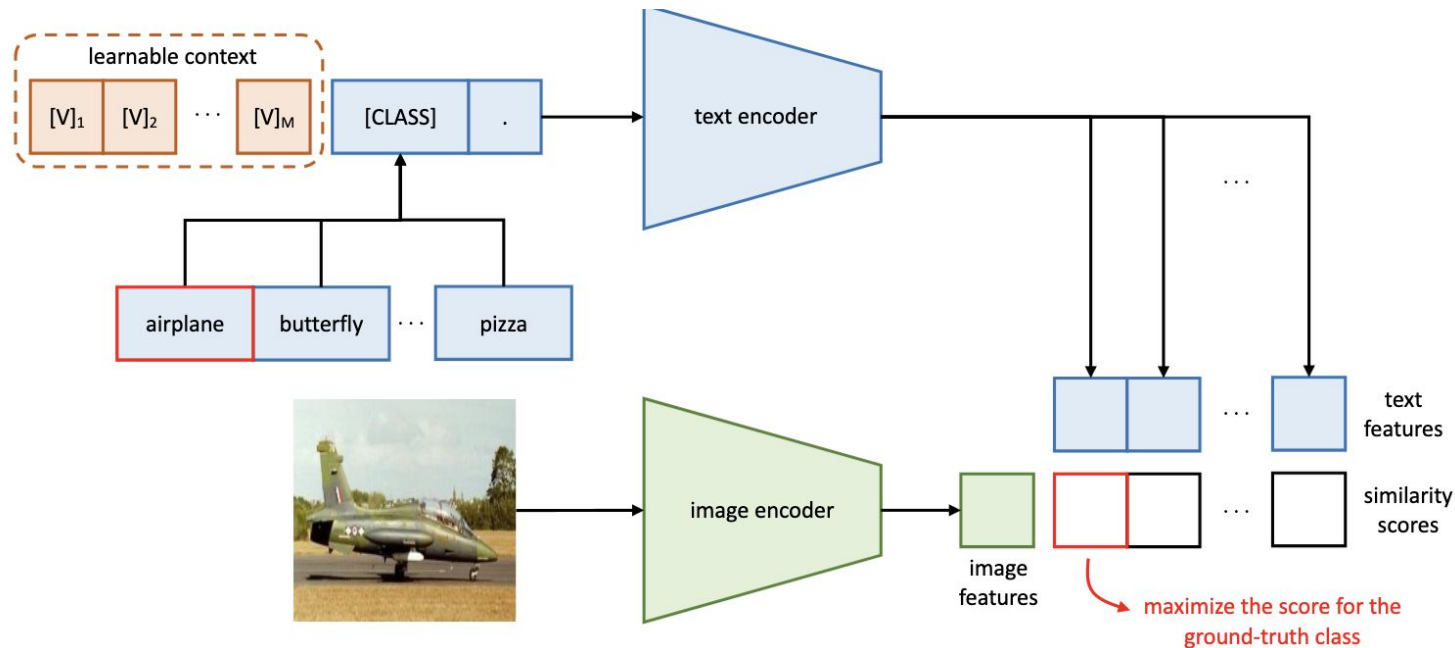
CLIP:

- Vision-language pre-training.
- Zero-shot classification on downstream datasets.

 CLIP – Radford, Alec, et al. ICML, 2021.

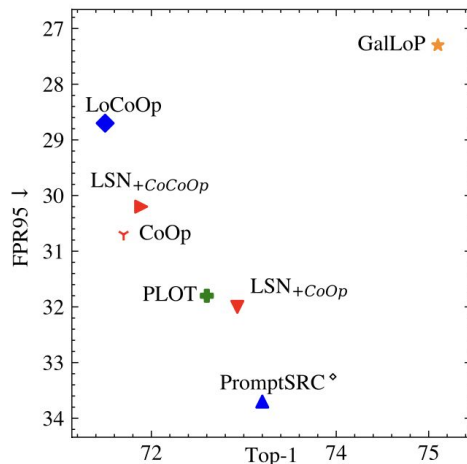
 ALIGN – Jia, Chao, et al. ICML, 2021.

Few shot adaptation with prompt learning.

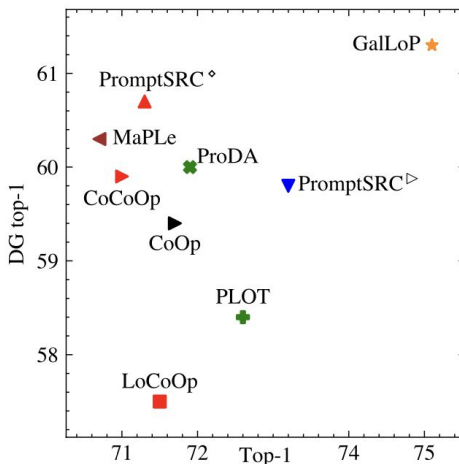


- Few shot adaptation of VLMs to downstream dataset.
- No need for “prompt engineering”.

CoOp – Zhou, Kaiyang, et al. IJCV, 2022.



(a) OOD detection



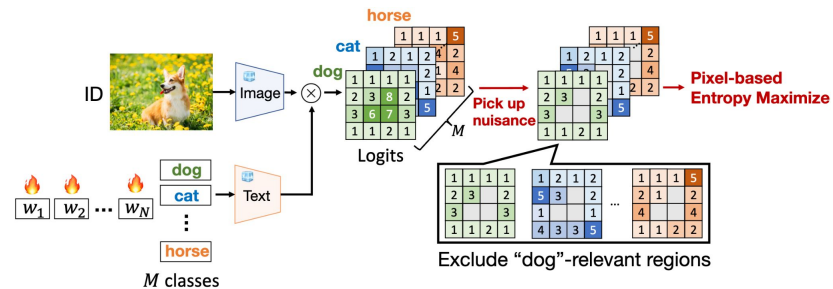
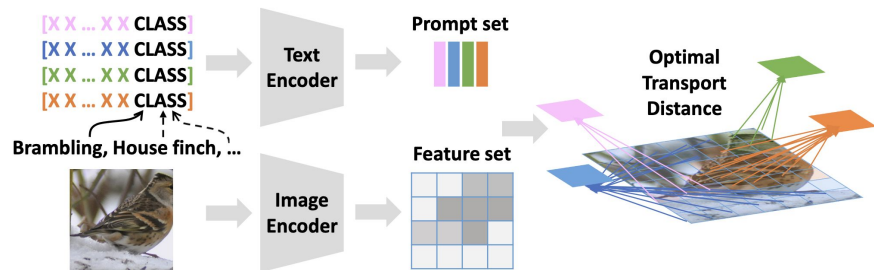
(b) Domain Generalization

- Prompt learning trade off top-1 vs. robustness.
- Ensembling and local features → robustness.

 LoCoOp – Miyai, Atsuyuki, et al. NeurIPS, 2023.

 PromptSRC – Khattak, Muhammad Uzair, et al. ICCV, 2023.

Prompt learning with local features.

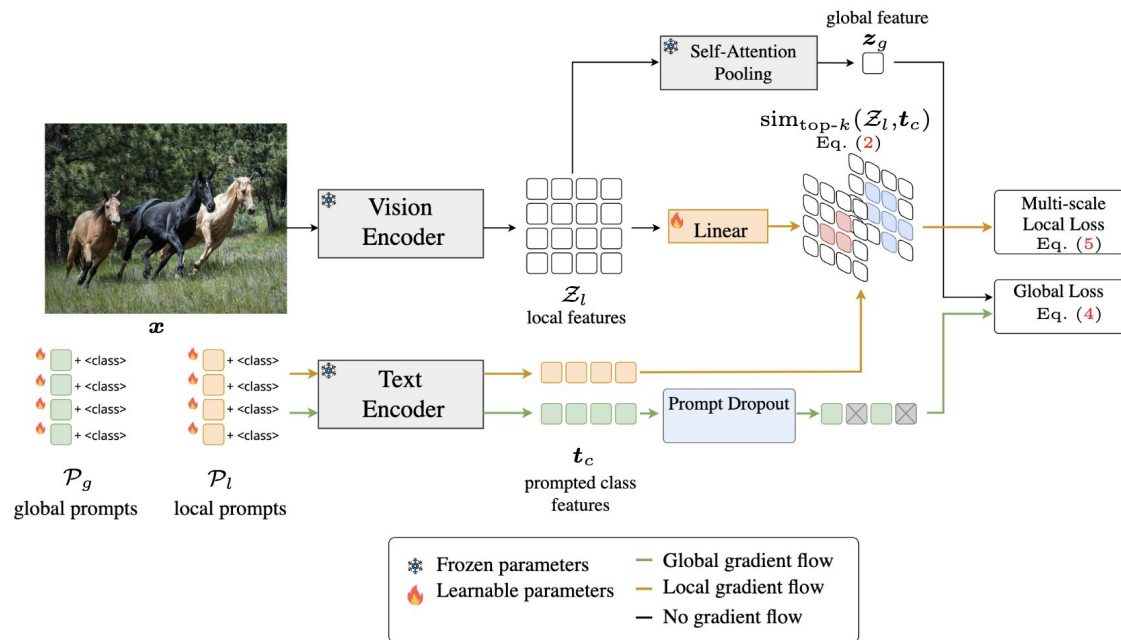


- Optimal transport for assignment.
- All local features $\rightarrow$ including irrelevant features.

- Local features for regularization.
- Lower top-1 accuracy

■ PLOT – Chen, Guangyi, et al. "Plot: Prompt learning with optimal transport for vision-language models." ICLR, 2023.

■ LoCoOp – Miyai, Atsuyuki, et al. "Locoop: Few-shot out-of-distribution detection via prompt learning." NeurIPS, 2023.



- Multiple global prompts w/ “prompt dropout”.
- Local alignment w/ sparsity and linear projection.
- Multiple local prompts w/ multiscale loss.

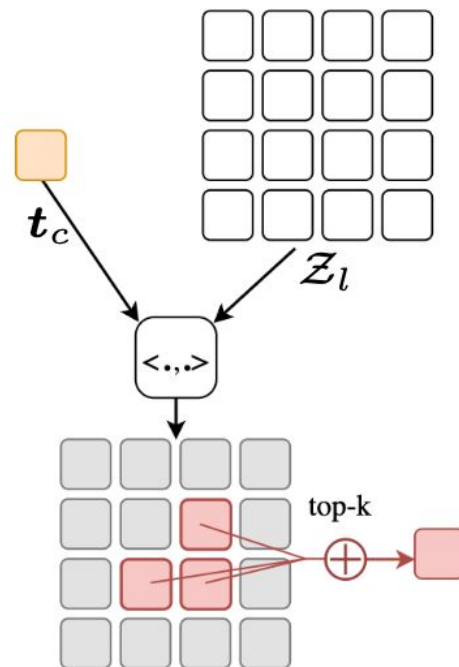
Prompt learning on local features.

$$\text{sim}_{\text{top-}k}(\mathcal{Z}_l, \mathbf{t}_c) := \frac{1}{k} \sum_{i=1}^L \mathbb{1}_{\text{top-}k}(i) \cdot \langle \mathbf{z}_i^l, \mathbf{t}_c \rangle$$

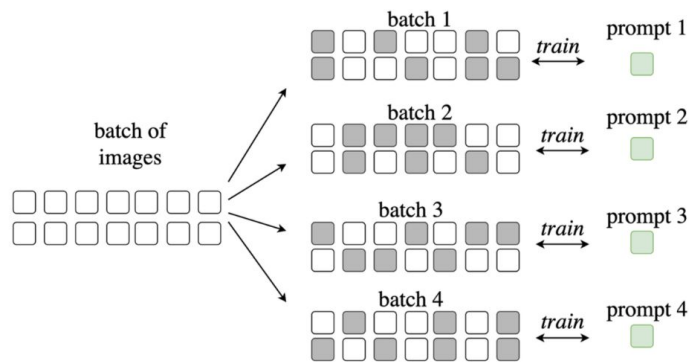
where $\mathbb{1}_{\text{top-}k}(i) = \begin{cases} 1 & \text{if } \text{rank}_i(\langle \mathbf{z}_i^l, \mathbf{t}_c \rangle) \leq k, \\ 0 & \text{otherwise.} \end{cases}$

Exploit local features for prompt learning:

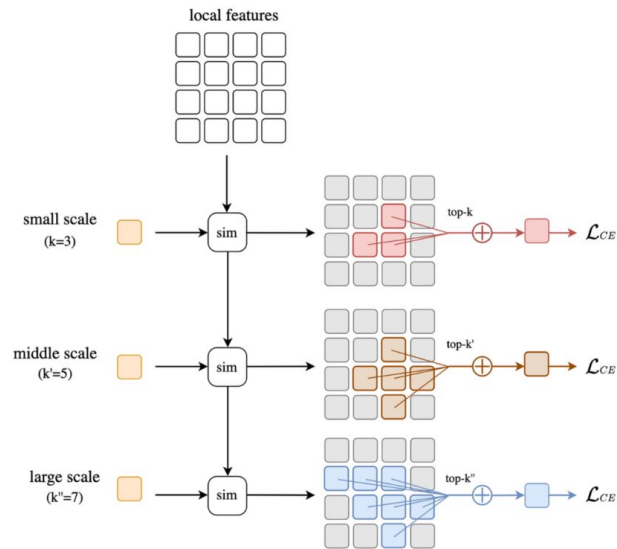
- Sparse local similarity.
- Learnable projection layer.



Diversity: prompt dropout & multiscale loss.



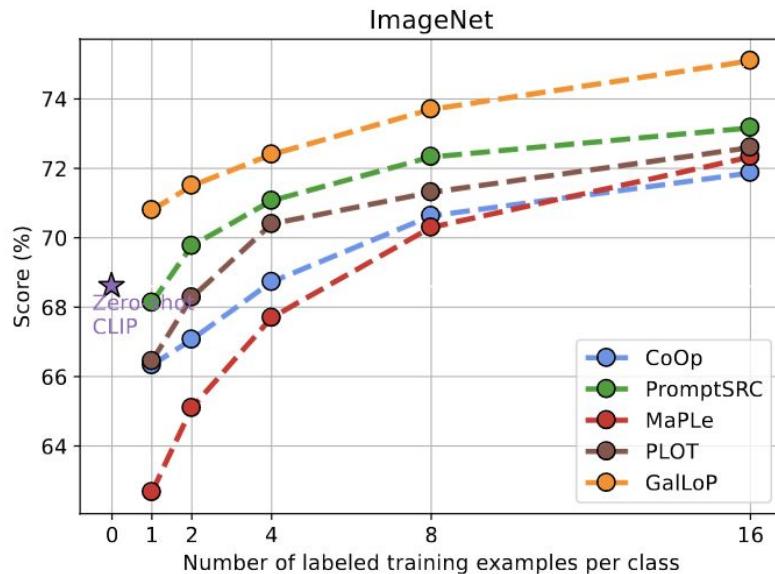
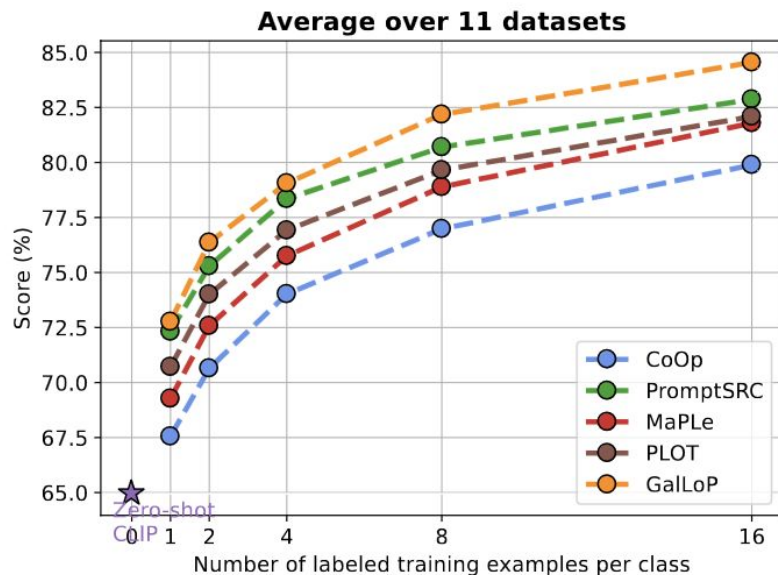
(a) Prompt dropout



(b) Multiscale loss

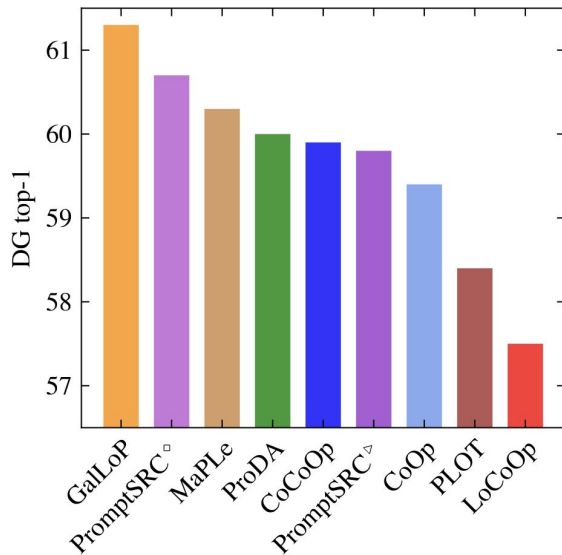
- Diversity by enforcing different gradient signal.

Few shots results.

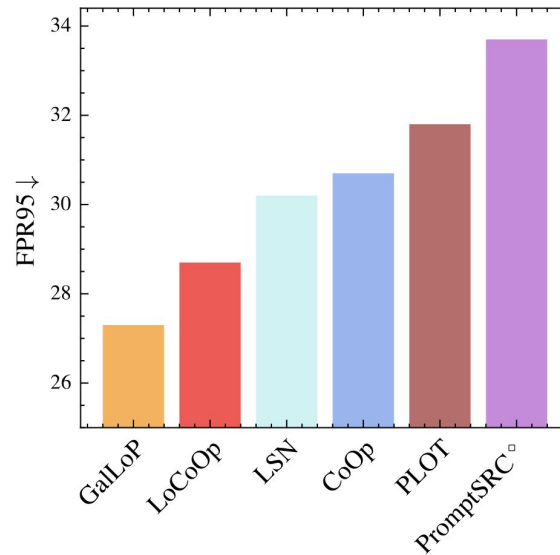


- Strong performances in **low shots** settings.
- Outperforms state-of-the-art prompt learning methods.

Robustness results.



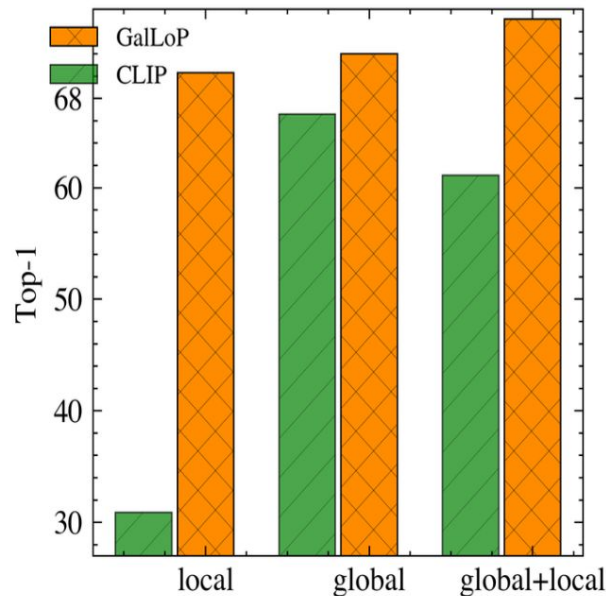
(a) Domain generalization.



(b) OOD detection.

- Strong domain generalization results (top-1).
- Outperforms dedicated OOD detection methods.

Combining global and local features

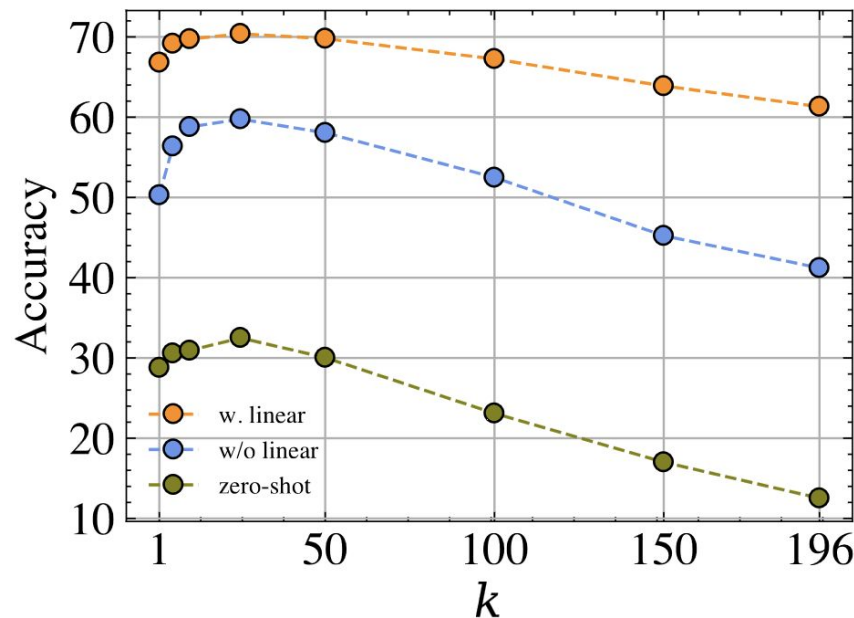


(c) Global-local

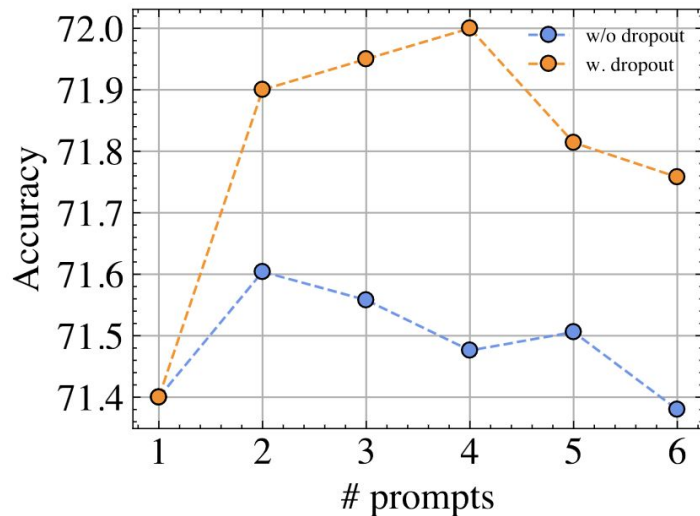
- Vanilla local features → no improvement.
- **Complementarity** of global / local features in GalLop.

The need for sparsity.

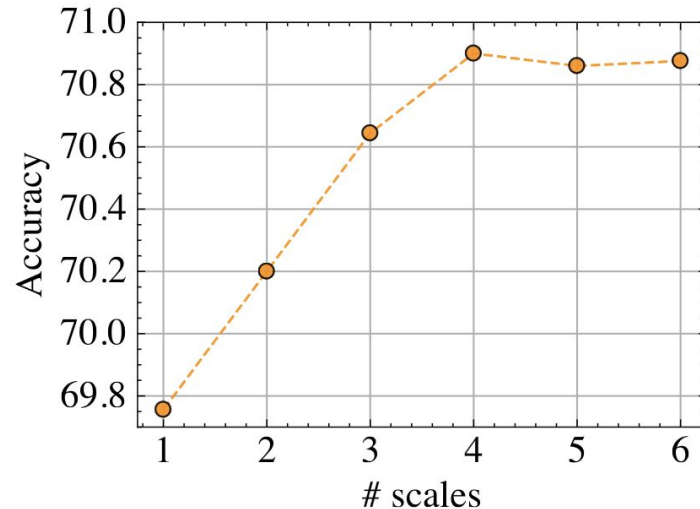
- Prompt learning → boosts performances.
- Sparsity → large gains in three regimes.
- Linear projection → better alignment.



Learning multiple prompts.



(a) Impact of prompt dropout when learning multiple global prompts.



(b) Impact of multiple scales when learning local prompts, with $k_1 = 10$, $\Delta_k = 10$.

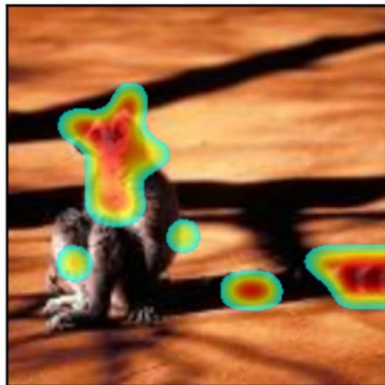
- Prompt dropout → gains when using multiple global prompts.
- Multiscale loss → boosts local performances.

Qualitative results.

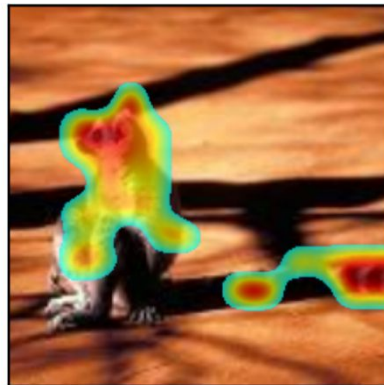
Scale 1, $k=10$



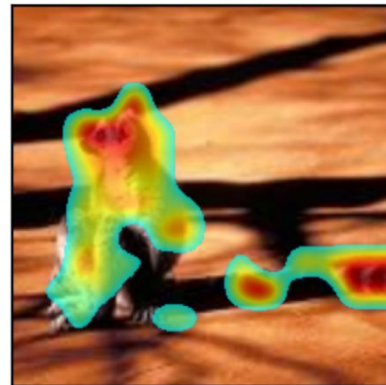
Scale 2, $k=20$



Scale 3, $k=30$



Scale 4, $k=40$



- Localization of objects.
- Few shot and weakly supervised segmentation.

- Exploit local Features: sparsity & learnable projection.
- Global + local prompt learning → strong top-1 + robustness.

impala (antelope)



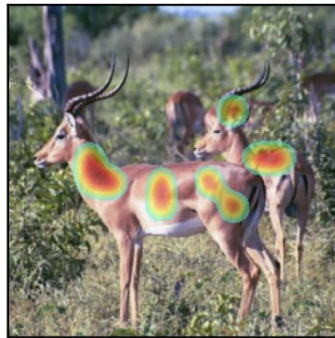
Ground truth

recreational vehicle



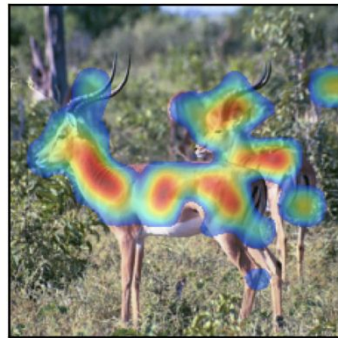
CLIP local

gazelle



GalLoP 1 scale (k=10)

gazelle

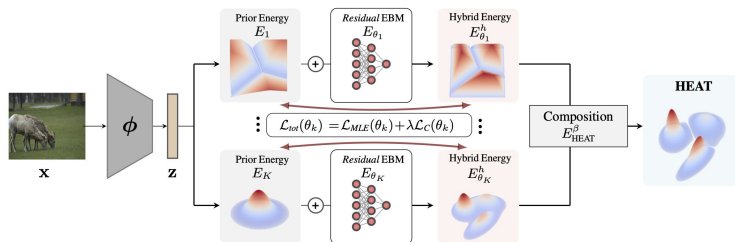


GalLoP 4 scales

Short term extensions.

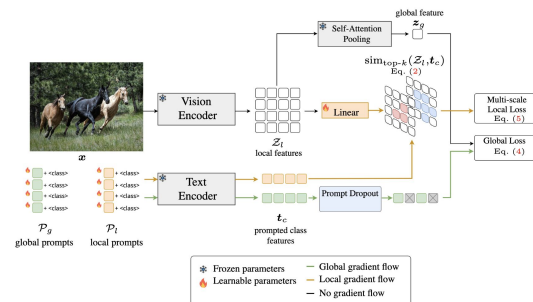
🔥 HEAT:

- Extend other scorers, e.g. kNN.
- Apply HEAT to dense prediction.
- Extension to other modalities, e.g. NLP, audio.



🐎 GalLoP:

- Integrate an adaptive sparsity ratio.
- Learn the linear projection on a bigger vision-language dataset.



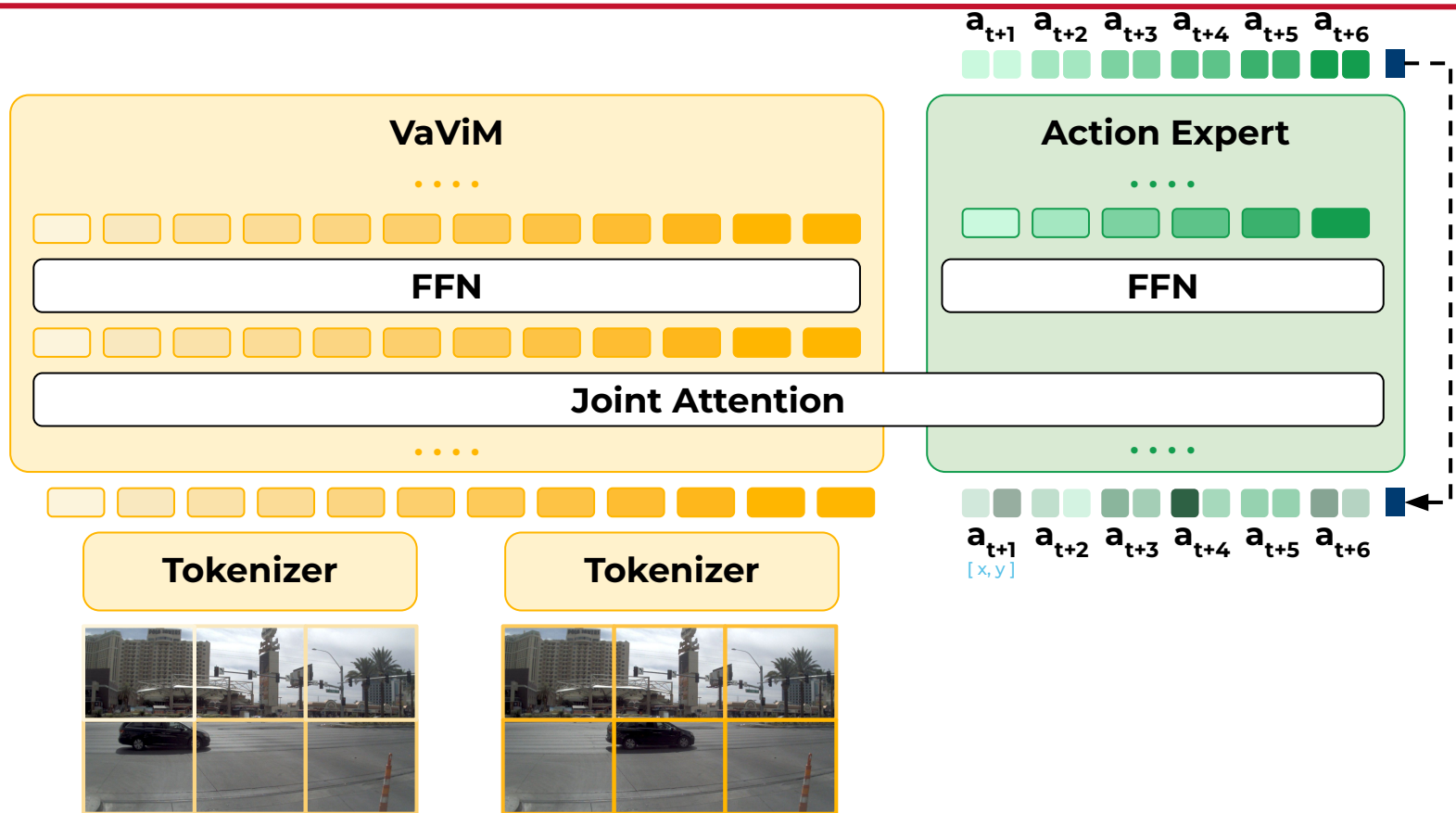
Current research interests.

le cnam
cedric EA4629

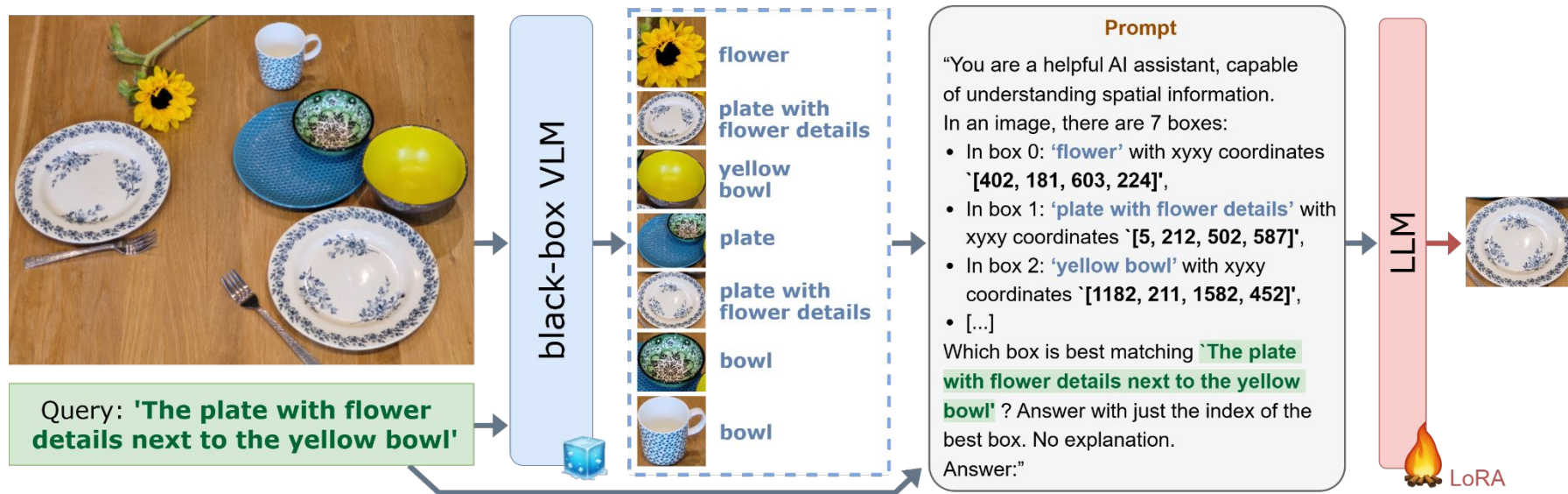
COEXYA
Connect skills, create more

valeo.ai

World model for autonomous driving



Vision language models



Cardiel, Amaia, et al. "LLM-wrapper: Black-Box Semantic-Aware Adaptation of Vision-Language Models for Referring Expression Comprehension." ICLR, 2025.

Thank you for your attention.

■ Elias Ramzi, et al. “Robust and Decomposable Average Precision for Image Retrieval.” *NeurIPS*, 2021 & *RFIAP*, 2022.

🔗 <https://github.com/elias-ramzi/ROADMAP>

■ EliasRamzi, et al. “Hierarchical Average Precision Training for Pertinent Image Retrieval.” *ECCV*, 2022.

🔗 <https://github.com/elias-ramzi/HAPPIER>

■ Marc Lafon, et al. “Hybrid Energy Based Model in the Feature Space for Out-of-Distribution Detection.” *ICML*, 2023.

🔗 <https://github.com/MarcLafon/heatood>

■ Elias Ramzi, et al. “Optimization of Rank Losses for Image Retrieval.” *TPAMI*, 2025.

🔗 <https://github.com/cvdfoundation/google-landmark>

■ Marc Lafon, et al. “GalLoP: Learning global and local prompts for vision-language models”, *ECCV*, 2024

🔗 <https://github.com/MarcLafon/gallop>

le cnam
cedric EA4629

COEXYA
Connect skills, create more

valeo.ai